

Operating System

Unit I OPERATING SYSTEMS OVERVIEW

Computer System Overview - Basic Elements, Instruction Execution, Interrupts, Memory Hierarchy, Cache Memory, Direct Memory Access, Multiprocessor and Multicore Organization. Operating system overview - objectives and functions, Evolution of Operating System - Computer System Organization - Operating System Structure and Operations - System Calls, System Programs, OS Generation and System Boot.

Unit II PROCESS MANAGEMENT

Processes-Process Concept, Process Scheduling, Operations on Processes, Interprocess Communication; Threads- Overview, Multicore Programming, Multithreading Models; Windows 7 - Process Synchronization - Critical Section Problem, Mutex Locks, Semaphores, Monitors; CPU Scheduling and Deadlocks.

Unit III STORAGE MANAGEMENT

Main Memory-Contiguous Memory Allocation, Segmentation, Paging, 32 and 64 bit architecture Examples; Virtual Memory- Demand Paging, Page Replacement, Allocation, Thrashing; Allocating Kernel Memory, OS Examples.

Unit IV I/O SYSTEMS

Mass Storage Structure- Overview, Disk Scheduling and Management; File System Storage-File Concepts, Directory and Disk Structure, Sharing and Protection; File System Implementation- File System Structure, Directory Structure, Allocation Methods, Free Space Management; I/O Systems.

Unit V CASE STUDY

Linux System Design Principles, Kernel Modules, Process Management, Scheduling, Memory Management, Input-Output Management, File System, Inter-process Communication; Mobile OS iOS and Android Architecture and SDK Framework, Media Layer, Services Layer, Core OS Layer, File System.

TEXT BOOK :

1. Abraham Silberschatz, Peter Baer Galvin and Greg Gagne, —Operating System Concepts||, 9th Edition, John Wiley and Sons Inc., 2012.

REFERENCES :

1. Ramaz Elmasri, A. Gil Carrick, David Levine, —Operating Systems – A Spiral Approach||, Tata McGraw Hill Edition, 2010.
2. Achyut S.Godbole, Atul Kahate, —Operating Systems||, McGraw Hill Education, 2016.
3. Andrew S. Tanenbaum, —Modern Operating Systems||, Second Edition, Pearson Education, 2004.
4. Gary Nutt, —Operating Systems||, Third Edition, Pearson Education, 2004.
5. Harvey M. Deitel, —Operating Systems||, Third Edition, Pearson Education, 2004.
6. Daniel P Bovet and Marco Cesati, —Understanding the Linux kernel||, 3rd edition, O'Reilly, 2005.

7. Neil Smyth, —iPhone iOS 4 Development Essentials – Xcode||, Fourth Edition, Payload media,2011.

COURSE OBJECTIVE

To understand the basic concepts, functions of OS .

1. Learn about process, threads & scheduling algorithms.
2. Understand the principles of concurrency & deadlock.
3. Learn various memory management schemes.
5. Learn the basics of Linux systems & mobile OS.

COURSE OUTCOMES

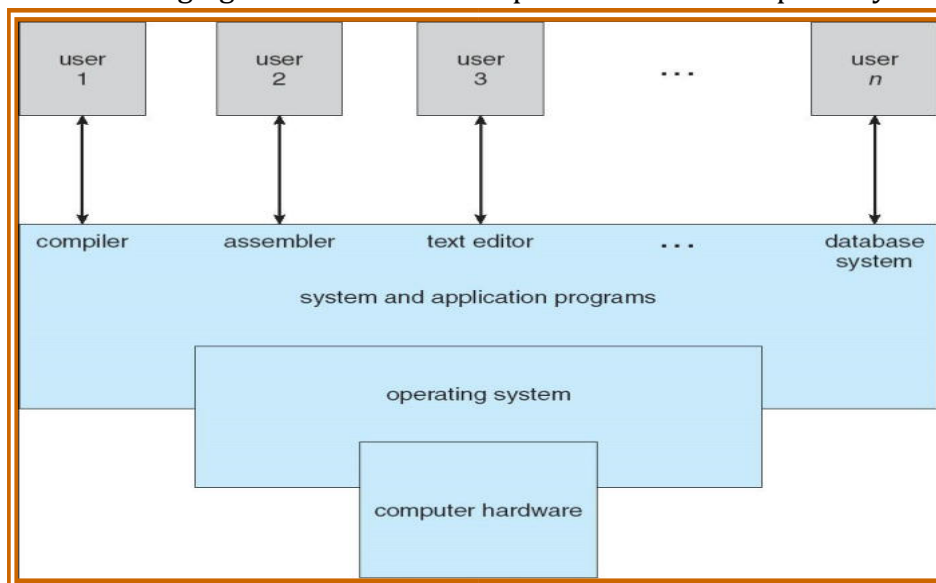
1. Understand the basic concepts and functions of Operating Systems
2. Delineate various threading models, process synchronization and deadlocks
3. Compare the performance of various CPU scheduling algorithms
4. Understand the basic concepts of memory management systems
5. Expound I/O management and file systems
6. Understand the model of Linux multifunction server and utilize local network services

UNIT-I

Introduction to OS & Their Classifications

- An OS is an intermediary between the user of the computer & the computer hardware.
 - It provides a basis for application program & acts as an intermediary between user of computer & computer hardware.
 - The purpose of an OS is to provide a environment in which the user can execute the program in a convenient & efficient manner.
 - OS is an important part of almost every computer systems.
 - A computer system can be roughly divided into four components
 - The Hardware
 - The OS
 - The application Program
 - The user
-
- The Hardware consists of memory, CPU, ALU, I/O devices, peripherals devices & storage devices.
 - The application program mainly consisted of word processors, spread sheets, compilers & web browsers defines the ways in which the resources are used to solve the problems of the users.
 - The OS controls & co-ordinates the use of hardware among various application program for various users.

The following figure shows the conceptual view of a computer system



Views of OS

1. User Views:- The user view of the computer depends on the interface used.

- Some users may use PCs. In this the system is designed so that only one user can utilize the resources and mostly for ease of use where the attention is mainly on performances and not on the resource utilization.
- Some users may use a terminal connected to a mainframe or minicomputers.
- Other users may access the same computer through other terminals. These users may share resources and exchange information. In this case the OS is designed to maximize resource utilization- so that all available CPU time, memory & I/O are used efficiently.
- Other users may sit at workstations, connected to the networks of other workstation and servers. In this case OS is designed to compromise between individual visibility & resource utilization.

2. System Views:-

- We can view system as resource allocator i.e. a computer system has many resources that may be used to solve a problem. The OS acts as a manager of these resources. The OS must decide how to allocate these resources to programs and the users so that it can operate the computer system efficiently and fairly.
- A different view of an OS is that it need to control various I/O devices & user programs i.e. an OS is a control program used to manage the execution of user program to prevent errors and improper use of the computer.
- Resources can be either CPU Time, memory space, file storage space, I/O devices and so on.

The OS must support the following tasks

- Provide the facility to create, modification of programs & data files using on editors.
- Access to compilers for translating the user program from high level language to machine language.
- Provide a loader program to move the compiled program code to computers memory for execution.
- Provides routines that handle the details of I/O programming.

I. Mainframe System:-

- Mainframe systems are mainly used for scientific & commercial applications.
- An OS may process its workload serially where the computer runs only one application or concurrently where computer runs many applications.

Batch Systems:-

- Early computers where physically large machines.
- The common I/P devices are card readers & tape drives.
- The common O/P devices are line printers, tape drives & card punches.
- The user do not interact directly with computers but we use to prepare a job with the program, data & some control information & submit it to the computer operator.
- The job was mainly in the form punched cards.
- At later time the O/P appeared and it consisted of result along with dump of memory and register content for debugging.
- The OS of these computers was very simple. Its major task was to transfer control from one job to the next. The OS was always resident in the memory. The processing of job was very slow. To improve the processing speed operators batched together the jobs with similar needs and processed it through the computers. This is called Batch Systems.
- In batch systems the CPU may be idle for some time because the speed of the mechanical devices slower compared to the electronic devices.
- Later improvement in technology and introduction of disks resulted in faster I/O devices.
- The introduction of disks allowed the OS to store all the jobs on the disk. The OS could perform the scheduling to use the resources and perform the task efficiently. ¶ The memory layout of simple batch system is shown below

OS
User program area

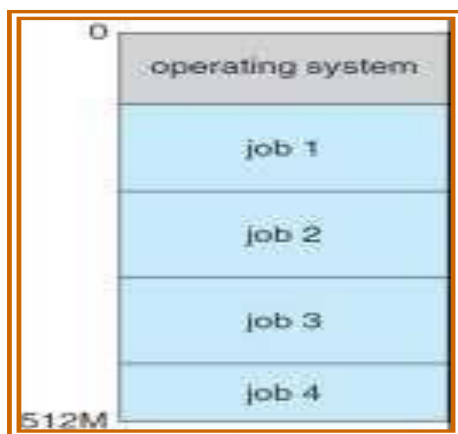
Disadvantages of Batch Systems:-

1. Turnaround time can be large from user.
2. Difficult to debug the program.
3. A job can enter into infinite loop.
4. A job could corrupt the monitor.
5. Due to lack of protection scheme, one job may affect the pending jobs.

Multi programmed System:-

- If there are two or more programs in the memory at the same time sharing the processor, this is referred as multi programmed OS.

- It increases the CPU utilization by organizing the jobs so that the CPU will always have one job to execute.
- Jobs entering the systems are kept in memory.
- OS picks the job from memory & it executes it.
- Having several jobs in the memory at the same time requires some form of memory management.
- Multi programmed systems monitors the state of all active program and system resources and ensures that CPU is never idle until there are no jobs.
- While executing a particular job, if the job has to wait for any task like I/O operation to be complete then the CPU will switch to some other jobs and starts executing it and when the first job finishes waiting the CPU will switch back to that.



- This will keep the CPU & I/O utilization busy.
- The following figure shows the memory layout of multi programmed OS

Time sharing Systems:-

- Time sharing system or multi tasking is logical extension of multi programming systems. The CPU executes multiple jobs by switching between them but the switching occurs so frequently that user can interact with each program while it is running.
- An interactive & hands on system provides direct communication between the user and the system. The user can give the instruction to the OS or program directly through key board or mouse and waits for immediate results.
- A time shared system allows multiple users to use the computer simultaneously. Since each action or commands are short in time shared systems only a small CPU time will be available for each of the user.
- A time shared systems uses CPU scheduling and multi programming to provide each user a small portion of time shared computers. When a process executes it will be

executing for a short time before it finishes or need to perform I/O. I/O is interactive i.e. O/P is to a display for the user and the I/O is from a keyboard, mouse etc.

- Since it has to maintain several jobs at a time, system should have memory management & protection.
- Time sharing systems are complex than the multi programmed systems. Since several jobs are kept in memory they need memory management and protection. To obtain less response time jobs are swapped in and out of main memory to disk. So disk will serve as backing store for main memory. This can be achieved by using a technique called virtual memory that allows for the execution of job i.e. not complete in memory.
- Time sharing system should also provide a file system & file system resides on collection of disks so this need disk management. It supports concurrent execution, job synchronization & communication.

I DESKTOP SYSTEMS:-

- PCs appeared in 1980s and during this they lacked the feature needed to protect an OS from user program & they even lack multi user nor multi tasking.
- The goals of those OS changed later with the time and new systems includes Microsoft Windows & Apple Macintosh.
- The Apple Macintosh OS ported to more advanced hardware & includes new features like virtual memory & multi tasking.
- Micro computers are developed for single user in 1980s & they can accommodate software with large capacity & greater speeds.
- MS-DOS is an example for micro computer OS & are used by commercial, educational, government enterprises.

II. Multi Processor Systems:-

- Multi processor systems include more than one processor in close communication.
- They share computer bus, the clock, m/y & peripheral devices.
- Two processes can run in parallel.
- Multi processor systems are of two types
 - a. Symmetric Multi processors (SMP)
 - b. Asymmetric Multi processors.
- In symmetric multi processing, each processors runs an identical copy of OS and they communicate with one another as needed. All the CPU shares the common memory.
- In asymmetric multi processing, each processors is assigned a specific task. It uses a master slave relationship. A master processor controls the system. The master processors schedules and allocates work to slave processors. The following figure shows asymmetric multi processors.

- SMP means all processors are peers i.e. no master slave relationship exists between processors. Each processor concurrently runs a copy of OS.
- The differences between symmetric & asymmetric multi-processing may be result of either H/w or S/w. Special H/w can differentiate the multiple processors or the S/w can be written to allow only master & multiple slaves.

Advantages of Multi Processor Systems:-

1. **Increased Throughput:-** By increasing the Number of processors we can get more work done in less time. When multiple process co operate on task, a certain amount of overhead is incurred in keeping all parts working correctly.
2. **Economy Of Scale:-** Multi processor system can save more money than multiple single processor, since they share peripherals, mass storage & power supplies. If many programs operate on same data, they will be stored on one disk & all processors can share them instead of maintaining data on several systems.
3. **Increased Reliability:-** If a program is distributed properly on several processors, than the failure of one processor will not halt the system but it only slows down.

III. Distributed Systems:-

- A distributed system is one in which H/w or S/w components located at the networked computers communicate & co ordinate their actions only by passing messages.
- A distributed systems looks to its user like an ordinary OS but runs on multiple, independent CPUs.
- Distributed systems depends on networking for their functionality which allows for communication so that distributed systems are able to share computational tasks and provides rich set of features to users.
- N/w may vary by the protocols used, distance between nodes & transport media.
Protocols->TCP/IP, ATM etc.
Network-> LAN, MAN, WAN etc.
Transport Media-> copper wires, optical fibers & wireless transmissions

Client-Server Systems:-

- Since PCs are faster, power full, cheaper etc. designers have shifted away from the centralized system architecture.
- User-interface functionality that used to be handled by centralized system is handled by PCs. So the centralized system today act as server program to satisfy the requests of client.

Server system can be classified as follows

- a. **Computer-Server System**:- Provides an interface to which client can send requests to perform some actions, in response to which they execute the action and send back result to the client.
- b. **File-Server Systems**:- Provides a file system interface where clients can create, update, read & delete files.

Peer-to-Peer Systems:-

- PC's are introduced in 1980's they are considered as standalone computers i.e. only one user can use it at a time.
- With wide spread use of internet PC's were connected to computer networks.
- With the introduction of the web in mid 1990's N/w connectivity became an essential component of a computer system.
- All modern PC's & workstation can run a web. OS also includes system software that enables the computer to access the web.
- In distributed systems or loosely coupled couple systems, the processor can communicate with one another through various communication lines like high speed buses or telephones lines.
- A N/w OS which has taken the concept of N/w & distributed system which provides features for file sharing across the N/w and also provides communication which allows different processors on different computers to share resources.

Advantages of Distributed Systems:-

1. Resource sharing.
2. Higher reliability.
3. Better price performance ratio.
4. Shorter response time.
5. Higher throughput.
6. Incremental growth

IV. Clustered Systems:-

- Like parallel systems the clustered systems will have multiple CPU but they are composed of two or more individual system coupled together.
- Clustered systems share storage & closely linked via LAN N/w.
- Clustering is usually done to provide high availability.
- Clustered systems are integrated with H/w & S/w. H/w clusters means sharing of high performance disk. S/w clusters are in the form of unified control of a computer system in a cluster.

- A layer of S/w cluster runs on the cluster nodes. Each node can monitor one or more of the others. If the monitored M/c fails the monitoring M/c take ownership of its storage and restart the application that were running on failed M/c.
- Clustered systems can be categorized into two groups
 - a. Asymmetric Clustering &
 - b. Symmetric clustering.
- In asymmetric clustering one M/c is in hot standby mode while others are running the application. The hot standby M/c does nothing but it monitors the active server. If the server fails the hot standby M/c becomes the active server.
- In symmetric mode two or more hosts are running the Application & they monitor each other. This mode is more efficient since it uses all the available H/w.
- Parallel clustering and clustering over a LAN is also available in clustering. Parallel clustering allows multiple hosts to access the same data on shared storage.
- Clustering provides better reliability than the multi processor systems.
- It provides all the key advantages of a distributed systems.
- Clustering technology is changing & include global clusters in which M/c could be anywhere in the world.

V. Real- Time Systems :-

- Real time system is one which were originally used to control autonomous systems like satellites, robots, hydroelectric dams etc.
- Real time system is one that must react to I/p & responds to them quickly.
- A real time system should not be late in response to one event.
- A real time should have well defined time constraints.

Real time systems are of two types

a. Hard Real Time Systems

b. Soft Real Time Systems

- A hard real time system guarantees that the critical tasks to be completed on time. This goal requires that all delays in the system be bounded from the retrieval of stored data to time that it takes the OS to finish the request.
- In soft real time system is a less restrictive one where a critical real time task gets priority over other tasks & retains the property until it completes. Soft real time system is achievable goal that can be mixed with other type of systems. They have limited utility than hard real time systems.
- Soft real time systems are used in area of multimedia, virtual reality & advanced scientific projects. It cannot be used in robotics or industrial controls due to lack of deadline support.

- Real time OS uses priority scheduling algorithm to meet the response requirement of a real time application.
- Soft real time requires two conditions to implement, CPU scheduling must be priority based & dispatch latency should be small.
- The primary objective of file management in real time systems is usually speed of access, rather than efficient utilization of secondary storage.

VI. Computing Environment:-

Different types of computing environments are:-

- a. Traditional Computing.
- b. Web Based Computing.
- c. Embedded Computing.

Traditional Computing:- Typical office environment uses traditional computing. Normal PC is used in traditional computing environment. N/w computers are essential terminals that understand web based computing. In domestic application most of the user had a single computer with internet connection. Cost of accessing internet is high. Web Based Computing has increased the emphasis on N/w. Web based computing uses PC, handheld PDA & cell phones. One of the feature of this type is load balancing. In load balancing, N/w connection is distributed among a pool of similar servers.

Embedded computing uses real time OS. Application of embedded computing is car engines, manufacturing robots, microwave ovens. This type of system provides limited features.

System Components :-

Modern OS supports all system components. The system components are,

- Process Management.
- Main M/y Management.
- File Management.
- Secondary Storage Management.
- I/O System management.
- Networking.
- Protection System.
- Command Interpreter System.

Process Management:-

- A process is a program in execution.
- A process abstraction is a fundamental OS mechanism for the management of concurrent program execution.
- The OS responds by creating process.

- Process requires certain resources like CPU time, M/y, I/O devices. These resources are allocated to the process when it created or while it is running.
- When process terminates the process reclaims all the reusable resources.
- Process refers to the execution of M/c instructions.
- A program by itself is not a process but is a passive entity.

The OS is responsible for the following activities of the process management, ?

Creating & destroying of the user & system process .

- Allocating H/w resources among the processes.
- Controlling the progress of the process.
- Provides mechanism for process communication.
- Provides mechanism for deadlock handling.

Main Memory Management:-

- Main M/y is the centre to the operation of the modern computer.
- Main M/y is the array of bytes ranging from hundreds of thousands to billions. Each byte will have their own address.
- The central processor reads the instruction from main M/y during instruction fetch cycle & it both reads & writes the data during the data-fetch cycle. The I/O operation reads and writes data in main M/y.
- The main M/y is generally a large storage device in which a CPU can address & access directly.
- When a program is to be executed it must be loaded into memory & mapped to absolute address. When it is executing it access the data & instruction from M/y by generating absolute address. When the program terminates all available M/y will be returned back.
- To improve the utilization of CPU & the response time several program will be kept in M/y.
- Several M/y management scheme are available & selection depends on the H/w design of the system.

The OS is responsible for the following activities.

- Keeping track of which part of the M/y is used & by whom.
- Deciding which process are to be loaded into M/y.
- Allocating & de allocating M/y space as needed.

File Management:-

- File management is one of the most visible component of an OS.
- Computer stores data on different types of physical media like Magnetic Disks, Magnetic tapes, optical disks etc.

- For convenient use of the computer system the OS provides uniform logical view of information storage.
- The OS maps file on to physical media & access these files via storage devices.
- A file is logical collection of information.
- File consists of both program & data. Data files may be numeric, alphabets or alphanumeric.
- Files can be organized into directories.

The OS is responsible for the following activities,

- Creating & deleting of files.
- ⇒ Creating & deleting directories.
- ⇒ Supporting primitives for manipulating files & directories.
- ⇒ Mapping files onto secondary storage.
- ⇒ Backing up files on stable storage media.

Secondary Storage management :-

- Is a mechanism where the computer system may store information in a way that it can be retrieved later.
- They are used to store both data & programs.
- The programs & data are stored in main memory.
- Since the size of the M/y is small & volatile Secondary storage devices is used.
- Magnetic disk is central importance of computer system.

The OS is responsible for the following activities, ⑦ Free space management.

- Storage allocation.
- Disk scheduling.

The entire speed of computer system depends on the speed of the disk sub system.

I/O System Management:-

- Each I/o device has a device handler that resides in separate process associated with that device.

The I/O management consists of,

- A M/y management component that include buffering,, caching & spooling.
- General device-driver interface.
- Drivers for specific H/w device.

Networking :-

- Networking enables users to share resources & speed up computations.
- The process communicates with one another through various communication lines like high speed buses or N/w.

Following parameters are considered while designing the N/w, 7 Topology of N/w.

- Type of N/w.
- Physical media.
- Communication protocol, 7 Routing algorithms.

Protection system:-

- Modern computer system supports many users & allows the concurrent execution of multiple processes organization rely on computers to store information. It necessary that the information & devices must be protected from unauthorized users or processors.
- The protection is a mechanism for controlling the access of program, processes or users to the resources defined by a computer system.
- Protection mechanism are implemented in OS to support various security policies.
- The goal of security system is to authenticate their access to any object.
- Protection can improve reliability by detecting latent errors at the interface B/w component sub system.
- Protection domains are extensions of H/w supervisor mode ability.

Command Interpreter System:-

- Command interpreter system between the user & the OS. It is a system program to the OS.
- Command interpreter is a special program in UNIX & MS DOS OS i.e. running when the user logs on.
- Many commands are given to the OS through control statements when the user logs on, a program that reads & interprets control statements is executed automatically. This program is sometimes called the control card interpreter or command line interpreter and is also called as shell.
- The command statements themselves deal with process creation & management, I/O handling, secondary storage management, main memory management, file system access, protection & N/w.

OPERATING SYSTEM SERVICES:-

An OS provides services for the execution of the programs and the users of such programs. The services provided by one OS may be different from other OS. OS makes the programming task easier. The common services provided by the OS are

1. Program Execution:- The OS must be able to load the program into memory & run that program. The program must end its execution either normally or abnormally.
2. I/O Operation:- A program running may require any I/O. This I/O may be a file or a specific device users can't control the I/O device directly so the OS must provide a means for controlling I/O devices.
3. File System Interface:- Program needs to read or write a file. The OS should provide permission for the creation or deletion of files by names.
4. Communication:- In certain situation one process may need to exchange information with another process. This communication may take place in two ways.
 - a. Between the processes executing on the same computer.
 - b. Between the processes executing on different computers that are connected by a network.

This communication can be implemented via shared memory or by OS.

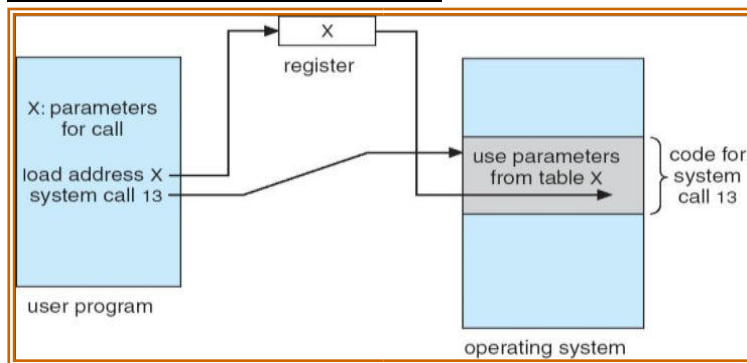
5. Error Detection:- Errors may occur in CPU, I/O devices or in M/y H/w. The OS constantly needs to be aware of possible errors. For each type of errors the OS should take appropriate actions to ensure correct & consistent computing. OS with multiple users provides the following services,
 - Resource Allocation:- When multiple users log onto the system or when multiple jobs are running, resources must be allocated to each of them. The OS manages different types of OS resources. Some resources may need some special allocation codes & others may have some general request & release code.
 - Accounting:- We need to keep track of which users use how many & what kind of resources. This record keeping may be used for accounting. This accounting data may be used for statistics or billing. It can also be used to improve system efficiency.
 - Protection:- Protection ensures that all the access to the system are controlled. Security starts with each user having authenticated to the system, usually by means of a password. External I/O devices must also be protected from invalid access. In multi process environment it is possible that one process may interface with the other or with the OS, so protection is required.

SYSTEM CALLS

- System provides interface between the process & the OS.

- The calls are generally available as assembly language instruction & certain system allow system calls to be made directly from a high level language program.
- Several language have been defined to replace assembly language program.
- A system call instruction generates an interrupt and allows OS to gain control of the processors.
- System calls occur in different ways depending on the computer. Some time more information is needed to identify the desired system call. The exact type & amount of information needed may vary according to the particular OS & call.

PASSING PARAMETERS TO OS



Three general methods are used to pass the parameters to the OS.

- The simplest approach is to pass the parameters in registers. In some there can be more parameters than register. In these the parameters are generally in a block or table in m/y and the address of the block is passed as parameters in register. This approach used by Linux.
- Parameters can also be placed or pushed onto stack by the program & popped off the stack by the OS.
- Some OS prefer the block or stack methods, because those approaches do not limit the number or length of parameters being passed.
- System calls may be grouped roughly into 5 categories
 1. Process control.
 2. File management.
 3. Device management.
 4. Information maintenance.
 5. Communication.

FILE MANAGEMENT

- System calls can be used to create & deleting of files. System calls may require the name of the files with attributes for creating & deleting of files.
- Other operation may involve the reading of the file, write & reposition the file after it is opened.

- Finally we need to close the file.
- For directories some set of operation are to be performed. Sometimes we require to reset some of the attributes on files & directories. The system call get file attribute & set file attribute are used for this type of operation.

DEVICE MANAGEMENT:-

- The system calls are also used for accessing devices.
- Many of the system calls used for files are also used for devices.
- In multi user environment the requirement are made to use the device. After using the device must be released using release system call the device is free to be used by another user. These function are similar to open & close system calls of files.
- Read, write & reposition system calls may be used with devices.
- MS-DOS & UNIX merge the I/O devices & the files to form file services structure. In file device structure I/O devices are identified by file names.

INFORMATION MAINTAINANCE:-

- Many system calls are used to transfer information between user program & OS. Example:- Most systems have the system calls to return the current time & date, number of current users, version number of OS, amount of free m/y or disk space & so on.
- In addition the OS keeps information about all its processes & there are system calls to access this information.

COMMUNICATION:-

There are two modes of communication,

1. **Message Passing Models:-**

- In this information is exchanged using inter-process communication facility provided by OS.
- Before communication the connection should be opened.
- The name of the other communicating party should be known, it can be on the same computer or it can be on another computer connected by a computer network.
- Each computer in a network may have a host name like IP name similarly each process can have a process name which can be translated into equivalent identifier by OS.

- The get host id & process id system call do this translation. These identifiers are then passed to the open & close connection system calls.
- The recipient process must give its permission for communication to take place with an accept connection call.

▪

Most processes receive the connection through special purpose system program dedicated for that purpose called daemons. The daemon on the server side is called server daemon & the daemon on the client side is called client daemon.

2. Shared Memory:-

- In this the processes uses the map m/y system calls to gain access to m/y owned by another process.
- The OS tries to prevent one process from accessing another process m/y.
- In shared m/y this restriction is eliminated and they exchange information by reading and writing data in shared areas. These areas are located by these processes and not under OS control.
- They should ensure that they are not writing to same m/y area.
- Both these types are commonly used in OS and some even implement both.
- Message passing is useful when small number of data need to be exchanged since no conflicts are to be avoided and it is easier to implement than in shared m/y. Shared m/y allows maximum speed and convenience of communication as it is done at m/y speed when within a computer.

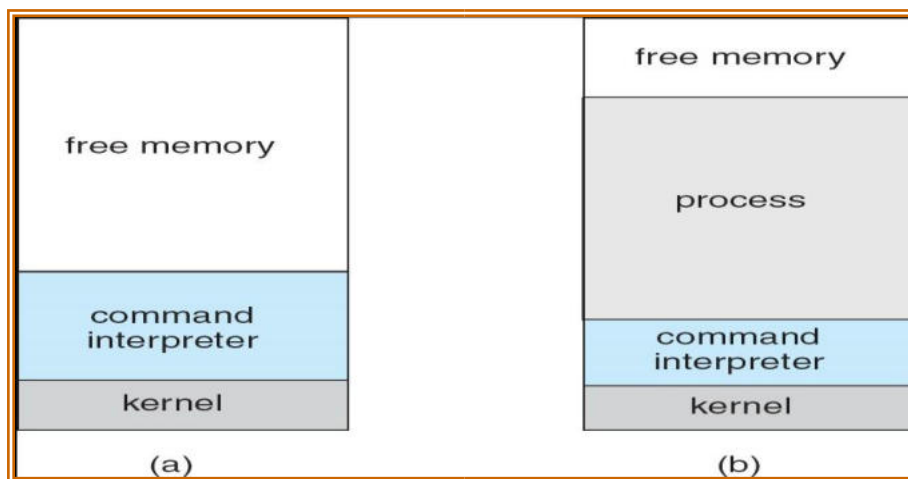
PROCESS CONTROL & JOB CONTROL

- A system call can be used to terminate the program either normally or abnormally. Reasons for abnormal termination are dump of m/y, error message generated etc.
- Debugger is mainly used to determine problem of the dump & returns back the dump to the OS.
- In normal or abnormal situations the OS must transfer the control to the command interpreter system.
- In batch system the command interpreter terminates the execution of job & continues with the next job.

- Some systems use control cards to indicate the special recovery action to be taken in case of errors.
- Normal & abnormal termination can be combined at some errors level. Error level is defined before & the command interpreter uses this error level to determine next action automatically.

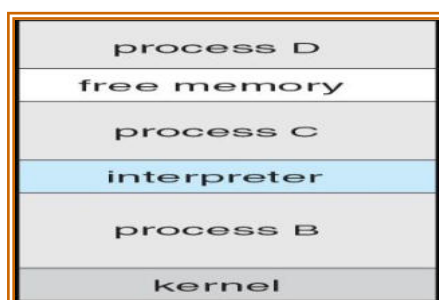
MS-DOS:-

MS-DOS is an example of single tasking system, which has command interpreter system i.e. invoked when the computer is started. To run a program MS-DOS uses simple method. It does not create a process when one process is running MS-DOS the program into m/y & gives the program as much as possible. It lacks the general multitasking capabilities.



BSD:-

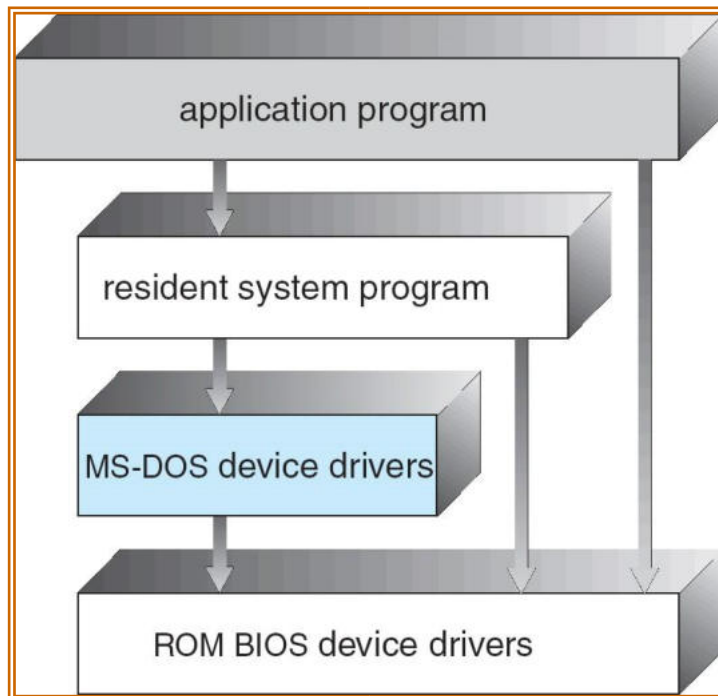
Free BSD is an example of multitasking system. In free BSD the command interpreter may continue running while other program is executing. FORK is used to create new process.



SYSTEM STRUCTURES

- Modern OS is large & complex.
- OS consists of different types of components.
- These components are interconnected & melded into kernel.

- For designing the system different types of structures are used. They are,
 - a. Simple structures.
 - b. Layered structured.
 - c. Micro kernels. **Simple Structures**
- Simple structure OS are small, simple & limited systems.
- The structure is not well defined
- MS-DOS is an example of simple structure OS. ? MS-DOS layer structure is shown below



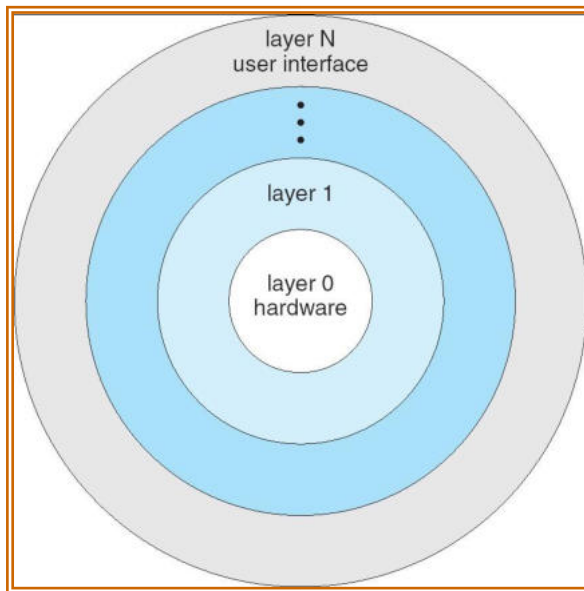
UNIX consisted of two separate modules

Kernel

The system programs.

- Kernel is further separated into series of interfaces & device drivers which were added & expanded as the UNIX evolved over years.
- The kernel also provides the CPU scheduling, file system, m/y management & other OS function through system calls.
- System calls define API to UNIX and system programs commonly available defines the user interface. The programmer and the user interface determines the context that the kernel must support.
- New versions of UNIX are designed to support more advanced H/w. the OS can be broken down into large number of smaller components which are more appropriate than the original MS-DOS.

Layered Approach



- In this OS is divided into number of layers, where one layer is built on the top of another layer. The bottom layer is hardware and higher layer is the user interface.
- An OS is an implementation of abstract object i.e. the encapsulation of data & operation to manipulate these data.
- The main advantage of layered approach is the modularity i.e. each layer uses the services & functions provided by the lower layer. This approach simplifies the debugging & verification. Once first layer is debugged the correct functionality is guaranteed while debugging the second layer. If an error is identified then it is a problem in that layer because the layer below it is already debugged.
- Each layer is designed with only the operations provided by the lower level layers.
- Each layer tries to hide some data structures, operations & hardware from the higher level layers.
- A problem with layered implementation is that they are less efficient than the other types.

Micro Kernels:-

- Micro kernel is a small OS which provides the foundation for modular extensions.
- The main function of the micro kernels is to provide communication facilities between the current program and various services that are running in user space.
- This approach was supposed to provide a high degree of flexibility and modularity.
- This benefits of this approach includes the ease of extending OS. All the new services are added to the user space & do not need the modification of kernel.
- This approach also provides more security & reliability.
- Most of the services will be running as user process rather than the kernel process.
- This was popularized by use in Mach OS.

UNIT II

PROCESS MANAGEMENT

Processes & Programs:-

- Process is a dynamic entity. A process is a sequence of instruction execution process exists in a limited span of time. Two or more process may execute the same program by using its own data & resources.

- A program is a static entity which is made up of program statement. Program contains the instruction. A program exists in a single space. A program does not execute by itself.

- A process generally consists of a process stack which consists of temporary data & data section which consists of global variables.

- It also contains program counter which represents the current activities.

- A process is more than the program code which is also called text section.

Process State:-

The process state consist of everything necessary to resume the process execution if it is somehow put aside temporarily. The process state consists of at least following:

- Code for the program.
- Program's static data.
- Program's dynamic data.
- Program's procedure call stack.
- Contents of general purpose registers.
- Contents of program counter (PC)
- Contents of program status word (PSW).
- Operating Systems resource in use.

Process operations

Process Creation

In general-purpose systems, some way is needed to create processes as needed during operation. There are four principal events led to processes creation.

- System initialization.
- Execution of a process Creation System calls by a running process.
- A user request to create a new process.
- Initialization of a batch job.

Foreground processes interact with users. Background processes that stay in background sleeping but suddenly springing to life to handle activity such as email, webpage, printing, and so on. Background processes are called daemons. This call creates an exact clone of the calling process.

A process may create a new process by some create process such as 'fork'. It choose to does so, creating process is called parent process and the created one is called the child processes. Only one parent is needed to create a child process. Note that unlike plants

and animals that use sexual representation, a process has only one parent. This creation of process (processes) yields a hierarchical structure of processes like one in the figure. Notice that each child has only one parent but each parent may have many children. After the fork, the two processes, the parent and the child, have the same memory image, the same environment strings and the same open files. After a process is created, both the parent and child have their own distinct address space. If either process changes a word in its address space, the change is not visible to the other process. Following are some reasons for creation of a process

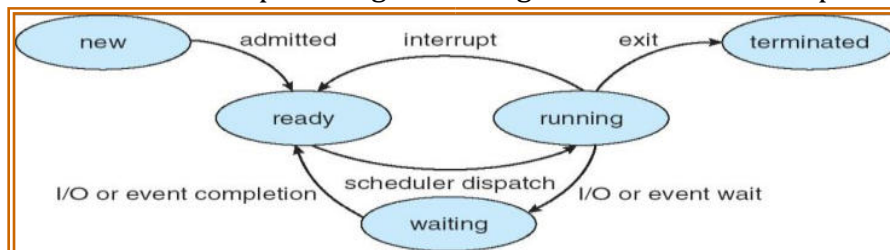
- User logs on.
- User starts a program.
- Operating systems creates process to provide service, e.g., to manage printer.
- Some program starts another process, e.g., Netscape calls xv to display a picture.

Process Termination

A process terminates when it finishes executing its last statement. Its resources are returned to the system, it is purged from any system lists or tables, and its process control block (PCB) is erased i.e., the PCB's memory space is returned to a free memory pool. The new process terminates the existing process, usually due to following reasons:

- **Normal Exist** Most processes terminates because they have done their job. This call is exist in UNIX.
- **Error Exist** When process discovers a fatal error. For example, a user tries to compile a program that does not exist.
- **Fatal Error** An error caused by process due to a bug in program for example, executing an illegal instruction, referring non-existing memory or dividing by zero.
- **Killed by another Process** A process executes a system call telling the Operating Systems to terminate some other process. In UNIX, this call is kill. In some systems when a process kills all processes it created are killed as well (UNIX does not work this way).

Process States : A process goes through a series of discrete process states.



- **New State** The process being created. **Terminated State** The process has finished execution.

Blocked (waiting) State When a process blocks, it does so because logically it cannot continue, typically because it is waiting for input that is not yet available. Formally, a process is said to be blocked if it is waiting for some event to happen (such as an I/O completion) before it can proceed. In this state a process is unable to run until some external event happens.

- **Running State** A process is said to be running if it currently has the CPU, that is, actually using the CPU at that particular instant.
- **Ready State** A process is said to be ready if it use a CPU if one were available. It is runnable but temporarily stopped to let another process run.

Logically, the 'Running' and 'Ready' states are similar. In both cases the process is willing to run, only in the case of 'Ready' state, there is temporarily no CPU available for it. The 'Blocked' state is different from the 'Running' and 'Ready' states in that the process cannot run, even if the CPU is available. **Process Control Block**

A process in an operating system is represented by a data structure known as a process control block (PCB) or process descriptor. The PCB contains important information about the specific process including

- The current state of the process i.e., whether it is ready, running, waiting, or whatever.
- Unique identification of the process in order to track "which is which" information.
- A pointer to parent process.
- Similarly, a pointer to child process (if it exists).
- The priority of process (a part of CPU scheduling information).
- Pointers to locate memory of processes.
 - A register save area.
 - The processor it is running on.

The PCB is a certain store that allows the operating systems to locate key information about a process. Thus, the PCB is the data structure that defines a process to the operating systems.

- The following figure shows the process control block.

The processes that are placed in main memory and are already waiting to execute are placed in a list called the ready queue. This is in the form of a linked list. The ready queue header contains a pointer to the first & final PCB in the list. Each PCB contains a pointer field that points to the next PCB in the ready queue.

Device Queue:-

The list of processes waiting for a particular I/O device is called the device queue. When the CPU is allocated to a process it may execute for some time & may quit or be interrupted or wait for the occurrence of a particular event like completion of an I/O request but the I/O may be busy with some other processes. In this case the process must wait for I/O. This will be placed in the device queue. Each device will have its own queue.

The process scheduling is represented using a queuing diagram. Queues are represented by the rectangular box & resources they need are represented by circles. It contains two queues: ready queue & device queues.

Once the process is assigned to the CPU and is executing the following events can occur,

- a. It can execute an I/O request and is placed in the I/O queue.
- b. The process can create a sub-process & wait for its termination.
- c. The process may be removed from the CPU as a result of an interrupt and can be put back into the ready queue.

Schedulers:-

The following are the different types of schedulers

1. **Long-term scheduler (or job scheduler)** – selects which processes should be brought into the ready queue.
2. **Short-term scheduler (or CPU scheduler)** – selects which process should be executed next and allocates the CPU.

3. Medium-term schedulers

-> Short-term scheduler is invoked very frequently (milliseconds) (must be fast)

-> Long-term scheduler is invoked very infrequently (seconds, minutes) (may be slow)

-> The long-term scheduler controls the *degree of multiprogramming* -> Processes can be described as either:

- I/O-bound process – spends more time doing I/O than computations, many short CPU bursts
- CPU-bound process – spends more time doing computations; few very long CPU bursts

Context Switch:-

1. When the CPU switches to another process, the system must save the state of the old process and load the saved state for the new process.
2. Context-switch time is overhead; the system does no useful work while switching.
3. Time dependent on hardware support

Cooperating Processes & Independent Processes

Independent process: one that is independent of the rest of the universe.

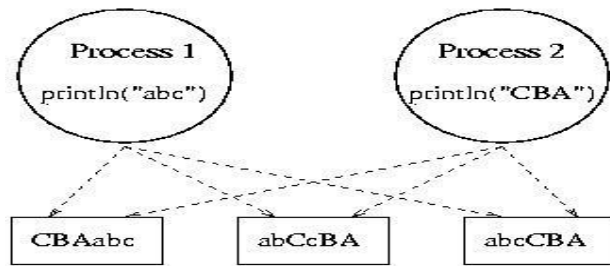
- Its state is not shared in any way by any other process.
- Deterministic: input state alone determines results.
- Reproducible.
- Can stop and restart with no bad effects (only time varies). Example: program that sums the integers from 1 to i (input).

There are many different ways in which a collection of independent processes might be executed on a processor:

- Uniprogramming: a single process is run to completion before anything else can be run on the processor.
- Multiprogramming: share one processor among several processes. If no shared state, then order of dispatching is irrelevant.
- Multiprocessing: if multiprogramming works, then it should also be ok to run processes in parallel on separate processors.
 - A given process runs on only one processor at a time.
 - A process may run on different processors at different times (move state, assume processors are identical).
 - Cannot distinguish multiprocessing from multiprogramming on a very fine grain.

Cooperating processes:

- Machine must model the social structures of the people that use it. People cooperate, so machine must support that cooperation. Cooperation means shared state, e.g. a single file system.
- Cooperating processes are those that share state. (May or may not actually be "cooperating")
- Behavior is nondeterministic: depends on relative execution sequence and cannot be predicted a priori.
- Behavior is irreproducible.
- Example: one process writes "ABC", another writes "CBA". Can get different outputs, cannot tell what comes from which. E.g. which process output first "C" in "ABCCBA"? Note the subtle state sharing that occurs here via the terminal. Not just anything can happen, though. For example, "AABBCC" cannot occur.



1. Independent process cannot affect or be affected by the execution of another process
2. Cooperating process can affect or be affected by the execution of another process
3. Advantages of process cooperation
 - Information sharing
 - Computation speed-up
 - Modularity
 - Convenience

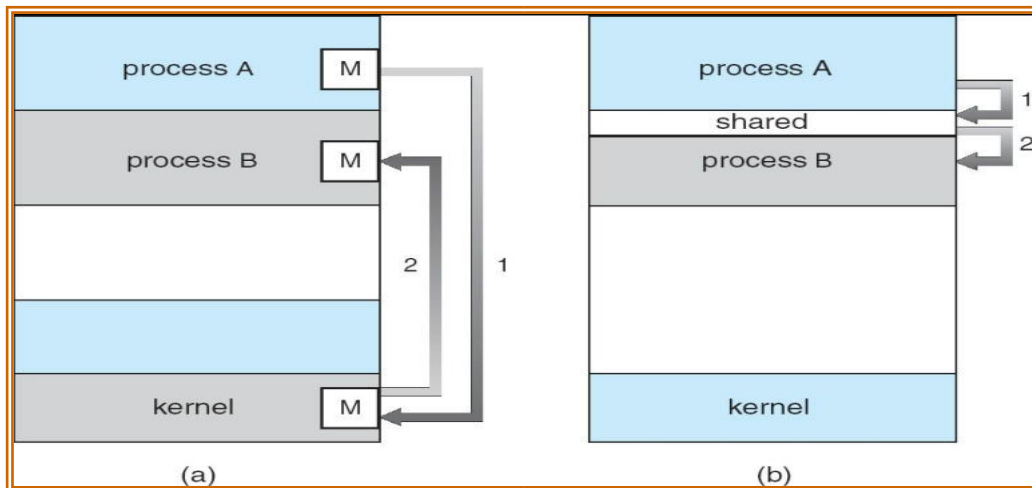
Interprocess Communication (IPC)

1. Mechanism for processes to communicate and to synchronize their actions
2. Message system – processes communicate with each other without resorting to shared variables
3. IPC facility provides two operations:
 - `send(message)` – message size fixed or variable
 - `receive(message)`
 1. If P and Q wish to communicate, they need to:
 - ☐ establish a *communication link* between them
 - ☐ exchange messages via send/receive
 2. Implementation of communication link
 - physical (e.g., shared memory, hardware bus)
 - logical (e.g., logical properties)

Communications Models

there are two types of communication models

1. Multi programming
2. Shared Memory



Direct Communication

1. Processes must name each other explicitly:
 - $\text{send}(P, \text{message})$ – send a message to process P
 - $\text{receive}(Q, \text{message})$ – receive a message from process Q
2. Properties of communication link
 - Links are established automatically
 - A link is associated with exactly one pair of communicating processes
 - Between each pair there exists exactly one link
 - The link may be unidirectional, but is usually bi-directional

Indirect Communication

1. Messages are directed and received from mailboxes (also referred to as ports)
 - Each mailbox has a unique id
 - Processes can communicate only if they share a mailbox
2. Properties of communication link
 - Link established only if processes share a common mailbox
 - A link may be associated with many processes
 - Each pair of processes may share several communication links
 - Link may be unidirectional or bi-direction
3. Operations
 - o create a new mailbox
 - o send and receive messages through mailbox
 - o destroy a mailbox
4. Primitives are defined as:
 - $\text{send}(A, \text{message})$ – send a message to mailbox A
 - $\text{receive}(A, \text{message})$ – receive a message from mailbox A
5. Mailbox sharing
 - $P1, P2$, and $P3$ share mailbox A

P1, sends; *P2* and *P3* receive Who gets the message?

6. Solutions

Allow a link to be associated with at most two processes

Allow only one process at a time to execute a receive operation

Allow the system to select arbitrarily the receiver. Sender is notified who the receiver was.

Synchronization

1. Message passing may be either blocking or non-blocking

2. Blocking is considered synchronous

->Blocking send has the sender block until the message is received.

->Blocking receive has the receiver block until a message is available.

3. Non-blocking is considered asynchronous

->Non-blocking send has the sender send the message and continue.

->Non-blocking receive has the receiver receive a valid message or null.

Buffering

->Queue of messages attached to the link; implemented in one of three ways

1. Zero capacity – 0 messages sender must wait for receiver (rendezvous) 2.

Bounded capacity – finite length of n messages

Sender must wait if link full

3. Unbounded capacity – infinite length sender never waits

THREADS:-

Despite of the fact that a thread must execute in process, the process and its associated threads are different concept. Processes are used to group resources together and threads are the entities scheduled for execution on the CPU.

A thread is a single sequence stream within in a process. Because threads have some of the properties of processes, they are sometimes called *lightweight processes*. In a process, threads allow multiple executions of streams. In many respect, threads are popular way to improve application through parallelism. The CPU switches rapidly back and forth among the threads giving illusion that the threads are running in parallel. Like a traditional process i.e., process with one thread, a thread can be in any of several states (Running, Blocked, Ready or Terminated). Each thread has its own stack. Since thread will generally call different procedures and thus a different execution history. This is why thread needs its own stack. An operating system that has thread facility, the basic unit of CPU utilization is a thread. A thread has or consists of a program counter (PC), a register set, and a stack space. Threads are not independent of one other like

processes as a result threads shares with other threads their code section, data section, OS resources also known as task, such as open files and signals.

Processes Vs Threads

As we mentioned earlier that in many respect threads operate in the same way as that of processes. Some of the similarities and differences are:

Similarities

- Like processes threads share CPU and only one thread active (running) at a time.
- Like processes, threads within a processes, threads within a processes execute sequentially.
- Like processes, thread can create children.
- And like process, if one thread is blocked, another thread can run.

Differences

- Unlike processes, threads are not independent of one another.
- Unlike processes, all threads can access every address in the task .
- Unlike processes, thread are design to assist one other. Note that processes might or might not assist one another because processes may originate from different users.

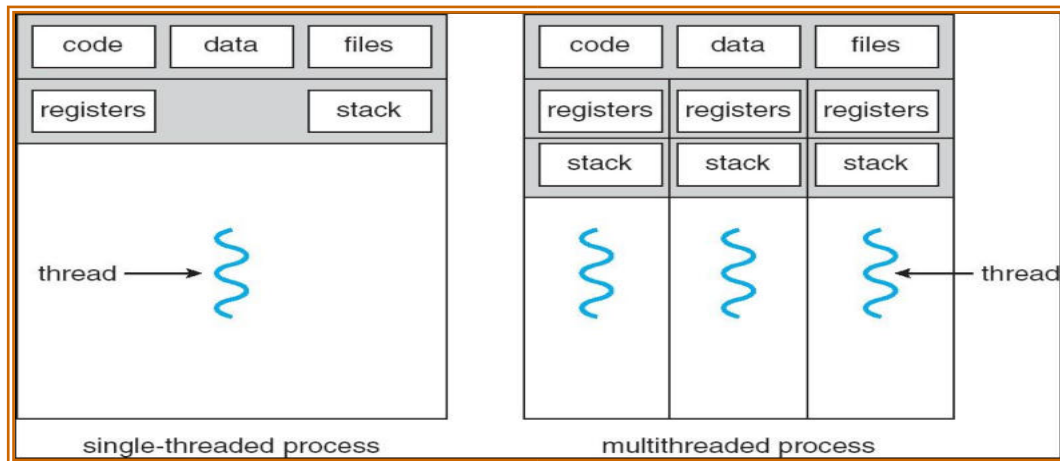
Why Threads?

Following are some reasons why we use threads in designing operating systems.

1. A process with multiple threads make a great server for example printer server.
2. Because threads can share common data, they do not need to use interprocess communication.
3. Because of the very nature, threads can take advantage of multiprocessors.
4. Responsiveness
5. Resource Sharing
6. Economy
7. Utilization of MP Architectures Threads are cheap in the sense that
 1. They only need a stack and storage for registers therefore, threads are cheap to create.
 2. Threads use very little resources of an operating system in which they are working. That is, threads do not need new address space, global data, program code or operating system resources.
 3. Context switching are fast when working with threads. The reason is that we only have to save and/or restore PC, SP and registers.

But this cheapness does not come free - the biggest drawback is that there is no protection between threads.

Single and Multithreaded Processes



User-Level Threads

1. Thread management done by user-level threads library
2. Three primary thread libraries:
 - > POSIX Pthreads
 - > Win32 threads
 - > Java threads

User-level threads implement in user-level libraries, rather than via systems calls, so thread switching does not need to call operating system and to cause interrupt to the kernel. In fact, the kernel knows nothing about user-level threads and manages them as if they were single-threaded processes.

Advantages:

The most obvious advantage of this technique is that a user-level threads package can be implemented on an Operating System that does not support threads. Some other advantages are

- User-level threads does not require modification to operating systems.
- Simple representation:
Each thread is represented simply by a PC, registers, stack and a small control block, all stored in the user process address space.
- Simple Management:
This simply means that creating a thread, switching between threads and synchronization between threads can all be done without intervention of the kernel.
- Fast and Efficient:
Thread switching is not much more expensive than a procedure call.

Disadvantages:

- There is a lack of coordination between threads and operating system kernel. Therefore, process as whole gets one time slice irrespective of whether process has one thread or 1000 threads within. It is up to each thread to relinquish control to other threads.
- User-level threads requires non-blocking systems call i.e., a multithreaded kernel. Otherwise, entire process will be blocked in the kernel, even if there are runnable threads left in the processes. For example, if one thread causes a page fault, the process blocks.

Kernel-Level Threads

1. Supported by the Kernel
2. Examples

->Windows XP/2000

->Solaris

->Linux

->Tru64 UNIX ->Mac OS X

In this method, the kernel knows about and manages the threads. No runtime system is needed in this case. Instead of thread table in each process, the kernel has a thread table that keeps track of all threads in the system. In addition, the kernel also maintains the traditional process table to keep track of processes. Operating Systems kernel provides system call to create and manage threads.

Advantages:

- Because kernel has full knowledge of all threads, Scheduler may decide to give more time to a process having large number of threads than process having small number of threads.
- Kernel-level threads are especially good for applications that frequently block.

Disadvantages:

- The kernel-level threads are slow and inefficient. For instance, threads operations are hundreds of times slower than that of user-level threads.
- Since kernel must manage and schedule threads as well as processes. It requires a full thread control block (TCB) for each thread to maintain information about threads. As a result there is significant overhead and increased in kernel complexity.

Advantages of Threads over Multiple Processes

it is much faster to switch between threads. In other words, it is relatively easier for a context switch using threads.

- **Sharing** Threads allow the sharing of a lot resources that cannot be shared in process, for example, sharing code section, data section, Operating System resources like open file etc.

Disadvantages of Threads over Multiprocesses

- **Blocking** The major disadvantage is that if the kernel is single threaded, a system call of one thread will block the whole process and CPU may be idle during the blocking period.
- **Security** Since there is, an extensive sharing among threads there is a potential problem of security. It is quite possible that one thread over writes the stack of another thread (or damaged shared data) although it is very unlikely since threads are meant to cooperate on a single task.

Application that Benefits from Threads

A proxy server satisfying the requests for a number of computers on a LAN would be benefited by a multi-threaded process. In general, any program that has to do more than one task at a time could benefit from multitasking. For example, a program that reads input, process it, and outputs could have three threads, one for each task.

Application that cannot Benefit from Threads

Any sequential process that cannot be divided into parallel task will not benefit from thread, as they would block until the previous one completes. For example, a program that displays the time of the day would not benefit from multiple threads.

Multithreading Models

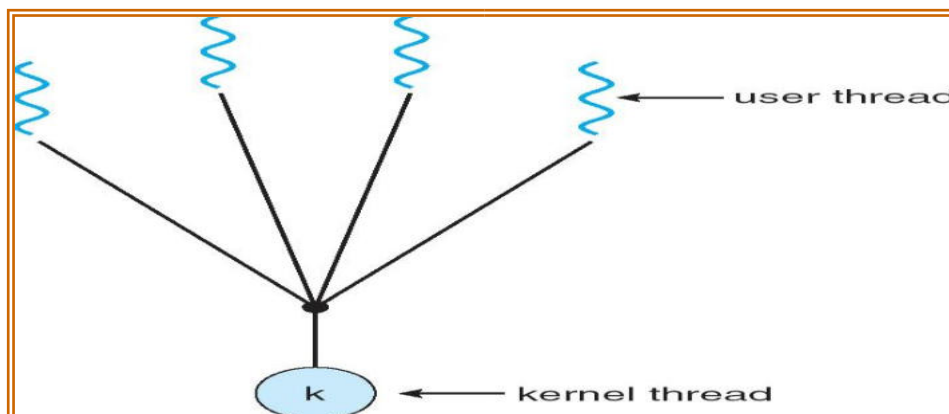
☐ Many-to-One ☐ One-to-One ☐ Many-to-Many

Many-to-One

Many user-level threads mapped to single kernel thread ->Examples:

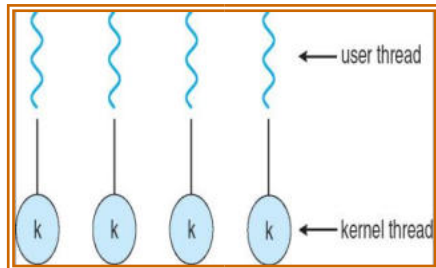
->Solaris Green Threads

->GNU Portable Threads



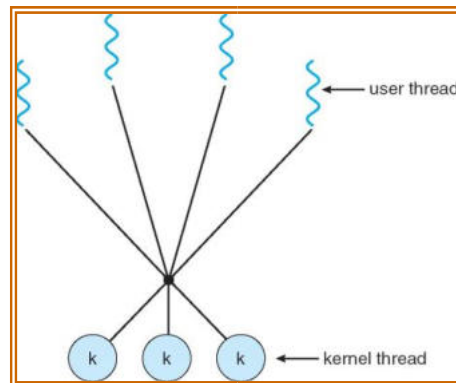
One-to-One

1. Each user-level thread maps to kernel thread
2. Examples
 - ❖ Windows NT/XP/2000
 - ❖ Linux
 - ❖ Solaris 9 and later



Many-to-Many Model

1. Allows many user level threads to be mapped to many kernel threads.
2. Allows the operating system to create a sufficient number of kernel threads.
3. Solaris prior to version 9.
4. Windows NT/2000 with the *ThreadFiber* package. Resources used in Thread Creation and Process Creation



When a new thread is created it shares its code section, data section and operating system resources like open files with other threads. But it is allocated its own stack, register set and a program counter.

The creation of a new process differs from that of a thread mainly in the fact that all the shared resources of a thread are needed explicitly for each process. So though two processes may be running the same piece of code they need to have their own copy of the code in the main memory to be able to run. Two processes also do not share other resources with each other. This makes the creation of a new process very costly compared to that of a new thread.

Thread Pools

1. Create a number of threads in a pool where they await work
2. Advantages:
 - ❖ Usually slightly faster to service a request with an existing thread than create a new thread
 - ❖ Allows the number of threads in the application(s) to be bound to the size of the pool

Context Switch

To give each process on a multiprogrammed machine a fair share of the CPU, a hardware clock generates interrupts periodically. This allows the operating system to schedule all processes in main memory (using scheduling algorithm) to run on the CPU at equal intervals. Each time a clock interrupt occurs, the interrupt handler checks how much time the current running process has used. If it has used up its entire time slice, then the CPU scheduling algorithm (in kernel) picks a different process to run. Each switch of the CPU from one process to another is called a context switch. **Major Steps of Context**

Switching

- The values of the CPU registers are saved in the process table of the process that was running just before the clock interrupt occurred.
- The registers are loaded from the process picked by the CPU scheduler to run next. In a multiprogrammed uniprocessor computing system, context switches occur frequently enough that all processes appear to be running concurrently. If a process has more than one thread, the Operating System can use the context switching technique to schedule the threads so they appear to execute in parallel. This is the case if threads are implemented at the kernel level. Threads can also be implemented entirely at the user level in run-time libraries. Since in this case no thread scheduling is provided by the Operating System, it is the responsibility of the programmer to yield the CPU frequently enough in each thread so all threads in the process can make progress.

Action of Kernel to Context Switch Among Threads

The threads share a lot of resources with other peer threads belonging to the same process. So a context switch among threads for the same process is easy. It involves switch of register set, the program counter and the stack. It is relatively easy for the kernel to accomplish this task.

Action of kernel to Context Switch Among Processes

Context switches among processes are expensive. Before a process can be switched its process control block (PCB) must be saved by the operating system. The PCB consists of the following information:

- The process state.
- The program counter, PC.
- The values of the different registers.
- The CPU scheduling information for the process.
- Memory management information regarding the process.
- Possible accounting information for this process.
- I/O status information of the process.

When the PCB of the currently executing process is saved the operating system loads the PCB of the next process that has to be run on CPU. This is a heavy task and it takes a lot of time.

CPU/Process Scheduling:-

The assignment of physical processors to processes allows processors to accomplish work. The problem of determining when processors should be assigned and to which processes is called processor scheduling or CPU scheduling.

When more than one process is runnable, the operating system must decide which one first. The part of the operating system concerned with this decision is called the scheduler, and algorithm it uses is called the scheduling algorithm.

CPU Scheduler

- a. Selects from among the processes in memory that are ready to execute, and allocates the CPU to one of them
 - b. CPU scheduling decisions may take place when a process:
 1. Switches from running to waiting state
 2. Switches from running to ready state
 3. Switches from waiting to ready
 4. Terminates
- ❖ Scheduling under 1 and 4 is *nonpreemptive*
 - ❖ All other scheduling is *preemptive*

Dispatcher

1. Dispatcher module gives control of the CPU to the process selected by the shortterm scheduler; this involves:
 - switching context
 - switching to user mode
 - jumping to the proper location in the user program to restart that program
2. Dispatch latency – time it takes for the dispatcher to stop one process and start another running.

Scheduling Criteria

1. CPU utilization – keep the CPU as busy as possible
2. Throughput – # of processes that complete their execution per time unit
3. Turnaround time – amount of time to execute a particular process
4. Waiting time – amount of time a process has been waiting in the ready queue
5. Response time – amount of time it takes from when a request was submitted until the first response is produced, **not** output (for time-sharing environment)

General Goals

Fairness

Fairness is important under all circumstances. A scheduler makes sure that each process gets its fair share of the CPU and no process can suffer indefinite postponement. Note that giving equivalent or equal time is not fair. Think of *safety control* and *payroll* at a nuclear plant. **Policy Enforcement**

The scheduler has to make sure that system's policy is enforced. For example, if the local policy is safety then the *safety control processes* must be able to run whenever they want to, even if it means delay in *payroll processes*.

Efficiency

Scheduler should keep the system (or in particular CPU) busy cent percent of the time when possible. If the CPU and all the Input/Output devices can be kept running all the time, more work gets done per second than if some components are idle.

Response Time

A scheduler should minimize the response time for interactive user.

Turnaround

A scheduler should minimize the time batch users must wait for an output.

Throughput

A scheduler should maximize the number of jobs processed per unit time. A little thought will show that some of these goals are contradictory. It can be shown that any scheduling algorithm that favors some class of jobs hurts another class of jobs.

The amount of CPU time available is finite, after all.

Preemptive Vs Nonpreemptive Scheduling

The Scheduling algorithms can be divided into two categories with respect to how they deal with clock interrupts.

Nonpreemptive Scheduling

A scheduling discipline is nonpreemptive if, once a process has been given the CPU, the CPU cannot be taken away from that process.

Following are some characteristics of nonpreemptive scheduling

1. In nonpreemptive system, short jobs are made to wait by longer jobs but the overall treatment of all processes is fair.
2. In nonpreemptive system, response times are more predictable because incoming high priority jobs can not displace waiting jobs.
3. In nonpreemptive scheduling, a scheduler executes jobs in the following two situations.
 - a. When a process switches from running state to the waiting state.
 - b. When a process terminates.

Preemptive Scheduling

A scheduling discipline is preemptive if, once a process has been given the CPU can taken away.

The strategy of allowing processes that are logically runnable to be temporarily suspended is called Preemptive Scheduling and it is contrast to the "run to completion" method.

Scheduling Algorithms

CPU Scheduling deals with the problem of deciding which of the processes in the ready queue is to be allocated the CPU.

Following are some scheduling algorithms we will study

❖ FCFS Scheduling.

▪

❖ Round Robin Scheduling.

❖ SJF Scheduling.

❖ SRT Scheduling.

❖ Priority Scheduling.

❖ Multilevel Queue Scheduling.

❖ Multilevel Feedback Queue Scheduling.

First-Come-First-Served (FCFS) Scheduling Other names of this algorithm are:

- First-In-First-Out (FIFO)
- Run-to-Completion
- Run-Until-Done

Perhaps, First-Come-First-Served algorithm is the simplest scheduling algorithm is the simplest scheduling algorithm. Processes are dispatched according to their arrival time on the ready queue. Being a nonpreemptive discipline, once a process has a CPU, it runs to completion. The FCFS scheduling is fair in the formal sense or human sense of fairness but it is unfair in the sense that long jobs make short jobs wait and unimportant jobs make important jobs wait.

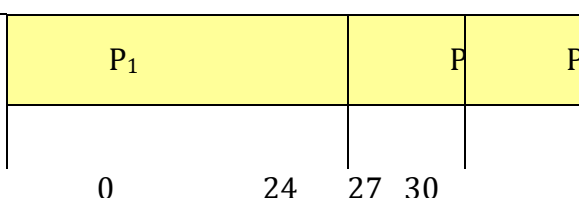
FCFS is more predictable than most of other schemes since it offers time. FCFS scheme is not useful in scheduling interactive users because it cannot guarantee good response time. The code for FCFS scheduling is simple to write and understand. One of the major drawback of this scheme is that the average time is often quite long.

The First-Come-First-Served algorithm is rarely used as a master scheme in modern operating systems but it is often embedded within other schemes.

Example:-

Process	Burst Time
<i>P1</i>	24
<i>P2</i>	3
<i>P3</i>	3

Suppose that the processes arrive in the order: *P1* , *P2* , *P3* The Gantt Chart for the schedule is:



Waiting time for $P1 = 0$; $P2 = 24$; $P3 = 27$

Average waiting time: $(0 + 24 + 27)/3 = 17$

Suppose that the processes arrive in the order $P2, P3, P1$. The Gantt chart for the schedule is:



Waiting time for $P1 = 6$; $P2 = 0$; $P3 = 3$

Average waiting time: $(6 + 0 + 3)/3 = 3$

Much better than previous case

Convoy effect short process behind long process

Round Robin Scheduling

One of the oldest, simplest, fairest and most widely used algorithm is round robin (RR).

In the round robin scheduling, processes are dispatched in a FIFO manner but are given a limited amount of CPU time called a time-slice or a quantum.

If a process does not complete before its CPU-time expires, the CPU is preempted and given to the next process waiting in a queue. The preempted process is then placed at the back of the ready list.

Round Robin Scheduling is preemptive (at the end of time-slice) therefore it is effective in time-sharing environments in which the system needs to guarantee reasonable response times for interactive users.

The only interesting issue with round robin scheme is the length of the quantum. Setting the quantum too short causes too many context switches and lower the CPU efficiency. On the other hand, setting the quantum too long may cause poor response time and approximates FCFS.

In any event, the average waiting time under round robin scheduling is often quite long.

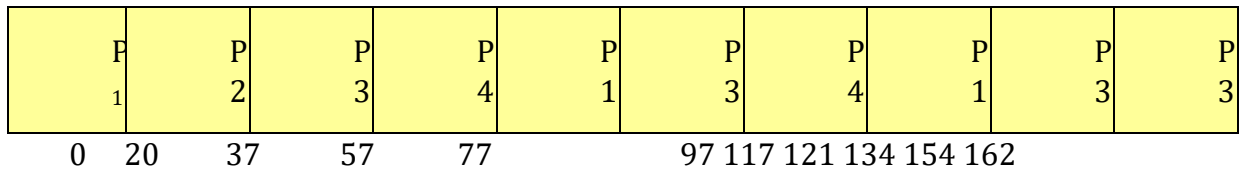
1. Each process gets a small unit of CPU time (*time quantum*), usually 10-100 milliseconds. After this time has elapsed, the process is preempted and added to the end of the ready queue.
2. If there are n processes in the ready queue and the time quantum is q , then each process gets $1/n$ of the CPU time in chunks of at most q time units at once. No process waits more than $(n-1)q$ time units.
3. Performance
 - > q large $\Rightarrow$ FIFO
 - > q small $\Rightarrow$ q must be large with respect to context switch, otherwise overhead is too high.

Example:-

<u>Process</u>	<u>Burst</u>
	<u>Time</u>

P1 53
P2 17
P3 68
P4 24

The Gantt chart is:



->Typically, higher average turnaround than SJF, but better *response*

C. Shortest-Job-First (SJF) Scheduling

Other name of this algorithm is Shortest-Process-Next (SPN).

Shortest-Job-First (SJF) is a non-preemptive discipline in which waiting job (or process) with the smallest estimated run-time-to-completion is run next. In other words, when CPU is available, it is assigned to the process that has smallest next CPU burst.

The SJF scheduling is especially appropriate for batch jobs for which the run times are known in advance. Since the SJF scheduling algorithm gives the minimum average time for a given set of processes, it is probably optimal.

The SJF algorithm favors short jobs (or processors) at the expense of longer ones. The obvious problem with SJF scheme is that it requires precise knowledge of how long a job or process will run, and this information is not usually available. The best SJF algorithm can do is to rely on user estimates of run times.

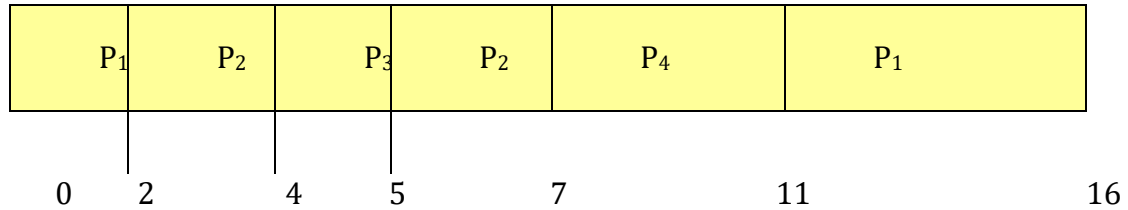
In the production environment where the same jobs run regularly, it may be possible to provide reasonable estimate of run time, based on the past performance of the process.

But in the development environment users rarely know how their program will execute. Like FCFS, SJF is non preemptive therefore, it is not useful in timesharing environment in which reasonable response time must be guaranteed.

1. Associate with each process the length of its next CPU burst. Use these lengths to schedule the process with the shortest time
2. Two schemes:
 - nonpreemptive – once CPU given to the process it cannot be preempted until completes its CPU burst
 - preemptive – if a new process arrives with CPU burst length less than remaining time of current executing process, preempt. This scheme is know as the ShortestRemaining-Time-First (SRTF)
3. SJF is optimal – gives minimum average waiting time for a given set of processes

Process	Arrival Time	Burst Time
P1	0.0	7
P2	2.0	4
P3	4.0	1
P4	5.0	4

->SJF (preemptive)



->Average waiting time = $(9 + 1 + 0 + 2)/4 = 3$

D. Shortest-Remaining-Time (SRT) Scheduling

- The SRT is the preemptive counterpart of SJF and useful in time-sharing environment.
- In SRT scheduling, the process with the smallest estimated run-time to completion is run next, including new arrivals.
- In SJF scheme, once a job begin executing, it run to completion.
- In SJF scheme, a running process may be preempted by a new arrival process with shortest estimated run-time.
- The algorithm SRT has higher overhead than its counterpart SJF.
- The SRT must keep track of the elapsed time of the running process and must handle occasional preemptions.
- In this scheme, arrival of small processes will run almost immediately. However, longer jobs have even longer mean waiting time.

E. Priority Scheduling

1. A priority number (integer) is associated with each process
2. The CPU is allocated to the process with the highest priority (smallest integer ▯ highest priority) ->Preemptive
- >nonpreemptive
3. SJF is a priority scheduling where priority is the predicted next CPU burst time
4. Problem ▯ Starvation – low priority processes may never execute
5. Solution ▯ Aging – as time progresses increase the priority of the process

The basic idea is straightforward: each process is assigned a priority, and priority is allowed to run. Equal-Priority processes are scheduled in FCFS order. The shortest-JobFirst (SJF) algorithm is a special case of general priority scheduling algorithm.

An SJF algorithm is simply a priority algorithm where the priority is the inverse of the (predicted) next CPU burst. That is, the longer the CPU burst, the lower the priority and vice versa.

Priority can be defined either internally or externally. Internally defined priorities use some measurable quantities or qualities to compute priority of a process. Examples of

Internal priorities are

- Time limits.
- Memory requirements.
- File requirements,
for example, number of open files.
- CPU Vs I/O requirements.

Externally defined priorities are set by criteria that are external to operating system such as

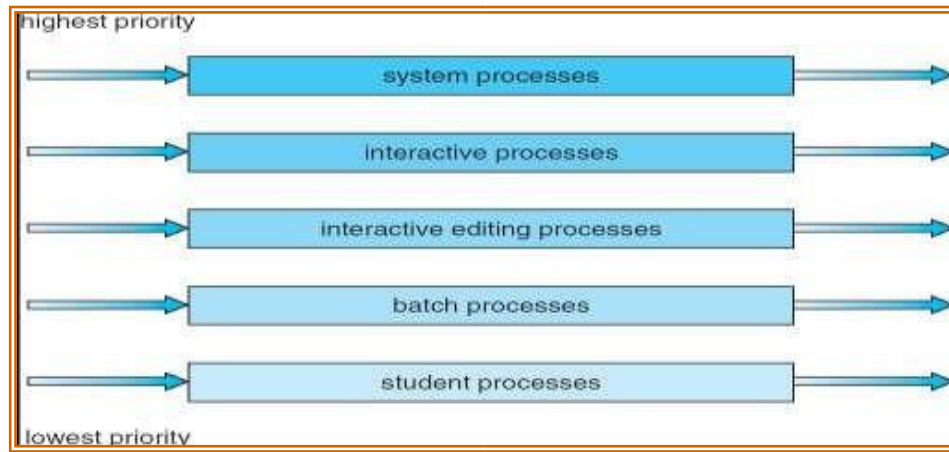
- The importance of process.
- Type or amount of funds being paid for computer use.
- The department sponsoring the work.
- Politics.

Priority scheduling can be either preemptive or non preemptive

- A preemptive priority algorithm will preemptive the CPU if the priority of the newly arrival process is higher than the priority of the currently running process.
- A non-preemptive priority algorithm will simply put the new process at the head of the ready queue.

A major problem with priority scheduling is indefinite blocking or starvation. A solution to the problem of indefinite blockage of the low-priority process is *aging*. Aging is a technique of gradually increasing the priority of processes that wait in the system for a long period of time.

F. Multilevel Queue Scheduling



A multilevel queue scheduling algorithm partitions the ready queue in several separate queues, for instance

In a multilevel queue scheduling processes are permanently assigned to one queues. The processes are permanently assigned to one another, based on some property of the process, such as

- Memory size
- Process priority
- Process type

Algorithm choose the process from the occupied queue that has the highest priority, and run that process either

- Preemptive or
- Non-preemptively

Each queue has its own scheduling algorithm or policy.

Possibility I

If each queue has absolute priority over lower-priority queues then no process in the queue could run unless the queue for the highest-priority processes were all empty. For example, in the above figure no process in the batch queue could run unless the queues for system processes, interactive processes, and interactive editing processes will all empty.

Possibility II

If there is a time slice between the queues then each queue gets a certain amount of

CPU times, which it can then schedule among the processes in its queue. For instance; 80% of the CPU time to foreground queue using RR.

20% of the CPU time to background queue using FCFS.

Since processes do not move between queue so, this policy has the advantage of low scheduling overhead, but it is inflexible.

G. Multilevel Feedback Queue Scheduling

Multilevel feedback queue-scheduling algorithm allows a process to move between queues. It uses many ready queues and associate a different priority with each queue. The Algorithm chooses to process with highest priority from the occupied queue and run that process either preemptively or unpreemptively. If the process uses too much CPU time it will moved to a lower-priority queue. Similarly, a process that wait too long in the lower-priority queue may be moved to a higher-priority queue may be moved to a highest-priority queue. Note that this form of aging prevents starvation.

- A process entering the ready queue is placed in queue 0.
- If it does not finish within 8 milliseconds time, it is moved to the tail of queue 1.
- If it does not complete, it is preempted and placed into queue 2.
- Processes in queue 2 run on a FCFS basis, only when 2 run on a FCFS basis queue, only when queue 0 and queue 1 are empty. Example:-

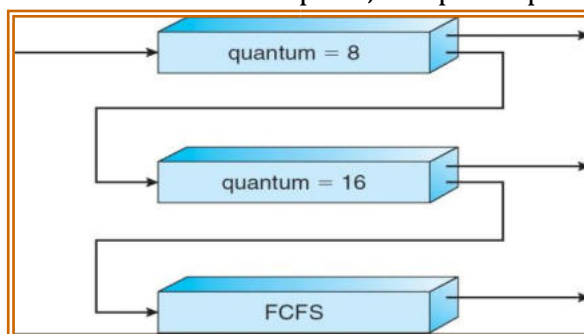
1. Three queues:

- ❖ Q0 – RR with time quantum 8 milliseconds
- ❖ Q1 – RR time quantum 16 milliseconds
- ❖ Q2 – FCFS

2. Scheduling

- ❖ A new job enters queue Q0 which is served FCFS. When it gains CPU, job receives 8 milliseconds. If it does not finish in 8 milliseconds, job is moved to queue Q1.
- ❖ At Q1 job is again served FCFS and receives 16 additional milliseconds.

If it still does not complete, it is preempted and moved to queue Q2.

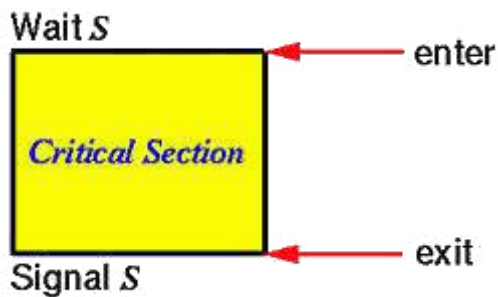


Process Synchronization & Deadlocks Interprocess Communication

Since processes frequently needs to communicate with other processes therefore, there is a need for a well-structured communication, without using interrupts, among processes.

Race Conditions

In operating systems, processes that are working together share some common storage (main memory, file etc.) that each process can read and write. When two or more



The key to preventing trouble involving shared storage is find some way to prohibit more than one process from reading and writing the shared data simultaneously. That part of the program where the shared memory is accessed is called the *Critical Section*. To avoid race conditions and flawed results, one must identify codes in *Critical Sections* in each thread. The characteristic properties of the code that form a *Critical Section* are

- Codes that reference one or more variables in a $\{ \text{read-update-write} \}$ fashion while any of those variables is possibly being altered by another thread.
- Codes that alter one or more variables that are possibly being referenced in $\{ \text{readupdate-write} \}$ fashion by another thread.
- Codes use a data structure while any part of it is possibly being altered by another thread.
- Codes alter any part of a data structure while it is possibly in use by another thread.

Here, the important point is that when one process is executing shared modifiable data in its critical section, no other process is to be allowed to execute in its critical section.

Thus, the execution of critical sections by the processes is mutually exclusive in time.

B. Mutual Exclusion

A way of making sure that if one process is using a shared modifiable data, the other processes will be excluded from doing the same thing.

Formally, while one process executes the shared variable, all other processes desiring to do so at the same time moment should be kept waiting; when that process has finished

executing the shared variable, one of the processes waiting; while that process has finished executing the shared variable, one of the processes waiting to do so should be allowed to proceed. In this fashion, each process executing the shared data (variables) excludes all others from doing so simultaneously. This is called Mutual Exclusion. Note that mutual exclusion needs to be enforced only when processes access shared modifiable data - when processes are performing operations that do not conflict with one another they should be allowed to proceed concurrently.

Mutual Exclusion Conditions

If we could arrange matters such that no two processes were ever in their critical sections simultaneously, we could avoid race conditions. We need four conditions to

hold to have a good solution for the critical section problem (mutual exclusion).

- No two processes may at the same moment inside their critical sections.
- No assumptions are made about relative speeds of processes or number of CPUs.
- No process should outside its critical section should block other processes.
- No process should wait arbitrary long to enter its critical section.

C. Proposals for Achieving Mutual Exclusion

The mutual exclusion problem is to devise a pre-protocol (or entry protocol) and a postprotocol (or exist protocol) to keep two or more threads from being in their critical sections at the same time. Tanenbaum examine proposals for critical-section problem or mutual exclusion problem.

Problem

When one process is updating shared modifiable data in its critical section, no other process should allowed to enter in its critical section.

Proposal 1 -Disabling Interrupts (Hardware Solution)

Each process disables all interrupts just after entering in its critical section and reenable all interrupts just before leaving critical section. With interrupts turned off the CPU could not be switched to other process. Hence, no other process will enter its critical and mutual exclusion achieved. **Conclusion**

Disabling interrupts is sometimes a useful interrupts is sometimes a useful technique within the kernel of an operating system, but it is not appropriate as a general mutual exclusion mechanism for users process. The reason is that it is unwise to give user process the power to turn off interrupts.

Proposal 2 - Lock Variable (Software Solution)

In this solution, we consider a single, shared, (lock) variable, initially 0. When a process wants to enter in its critical section, it first test the lock. If lock is 0, the process first sets it to 1 and then enters the critical section. If the lock is already 1, the process just waits until (lock) variable becomes 0. Thus, a 0 means that no process in

its critical section, and 1 means hold your horses - some process is in its critical section.

Conclusion

The flaw in this proposal can be best explained by example. Suppose process A sees that the lock is 0. Before it can set the lock to 1 another process B is scheduled, runs, and sets the lock to 1. When the process A runs again, it will also set the lock to 1, and two processes will be in their critical section simultaneously.

Proposal 3 - Strict Alteration

In this proposed solution, the integer variable 'turn' keeps track of whose turn is to enter the critical section. Initially, process A inspect turn, finds it to be 0, and enters in its critical section. Process B also finds it to be 0 and sits in a loop continually testing 'turn' to see when it becomes 1. Continuously testing a variable waiting for some value to appear is called the *Busy-Waiting*.

Conclusion

Taking turns is not a good idea when one of the processes is much slower than the other. Suppose process 0 finishes its critical section quickly, so both processes are now in their noncritical section. This situation violates above mentioned condition 3.

Using Systems calls 'sleep' and 'wakeup'

Basically, what above mentioned solution do is this: when a processes wants to enter in its critical section , it checks to see if then entry is allowed. If it is not, the process goes into tight loop and waits (i.e., start busy waiting) until it is allowed to enter. This approach waste CPU-time.

Now look at some interprocess communication primitives is the pair of steep-wakeup.

- Sleep

- o It is a system call that causes the caller to block, that is, be suspended until some other process wakes it up.
 - Wakeup
- o It is a system call that wakes up the process.

Both 'sleep' and 'wakeup' system calls have one parameter that represents a memory address used to match up 'sleeps' and 'wakeups'.

The Bounded Buffer Producers and Consumers

The bounded buffer producers and consumers assumes that there is a fixed buffer size i.e., a finite number of slots are available.

Statement

To suspend the producers when the buffer is full, to suspend the consumers when the buffer is empty, and to make sure that only one process at a time manipulates a buffer so there are no race conditions or lost updates.

As an example how sleep-wakeup system calls are used, consider the producer consumer problem also known as bounded buffer problem.

Two processes share a common, fixed-size (bounded) buffer. The producer puts information into the buffer and the consumer takes information out. Trouble arises when

1. The producer wants to put a new data in the buffer, but buffer is already full. Solution: Producer goes to sleep and to be awakened when the consumer has removed data.
2. The consumer wants to remove data the buffer but buffer is already empty. Solution: Consumer goes to sleep until the producer puts some data in buffer and wakes consumer up.

Conclusion

This approaches also leads to same race conditions we have seen in earlier approaches. Race condition can occur due to the fact that access to 'count' is unconstrained. The essence of the problem is that a wakeup call, sent to a process that is not sleeping, is lost.

D. Semaphores

E.W. Dijkstra (1965) abstracted the key notion of mutual exclusion in his concepts of semaphores.

Definition

A semaphore is a protected variable whose value can be accessed and altered only by the operations P and V and initialization operation called 'Semaphoiinitisize'.

Binary Semaphores can assume only the value 0 or the value 1 counting semaphores also called general semaphores can assume only nonnegative values.

The P (or wait or sleep or down) operation on semaphores S, written as P(S) or wait (S), operates as follows:

```
P(S): IF S > 0  
THEN S := S - 1  
ELSE (wait on S)
```

The V (or signal or wakeup or up) operation on semaphore S, written as V(S) or signal (S), operates as follows:

```
V(S): IF (one or more process are waiting on S)  
THEN (let one of these processes proceed) ELSE S := S + 1
```

Operations P and V are done as single, indivisible, atomic action. It is guaranteed that once a semaphore operations has started, no other process can access the semaphore until operation has completed. Mutual exclusion on the semaphore, S, is enforced within **P(S)** and **V(S)**.

If several processes attempt a P(S) simultaneously, only process will be allowed to proceed. The other processes will be kept waiting, but the implementation of P and V guarantees that processes will not suffer indefinite postponement.

Semaphores solve the lost-wakeup problem.

Semaphore as General Synchronization Tool

1. Counting semaphore – integer value can range over an unrestricted domain.
2. Binary semaphore – integer value can range only between 0 and 1; can be simpler to implement Also known as mutex locks.
3. Can implement a counting semaphore S as a binary semaphore.
4. Provides mutual exclusion
 - Semaphore S; // initialized to 1
 - wait (S);

Critical Section signal (S);

Semaphore Implementation

1. Must guarantee that no two processes can execute wait () and signal () on the same semaphore at the same time
2. Thus, implementation becomes the critical section problem where the wait and signal code are placed in the critical section.
 - Could now have busy waiting in critical section implementation
 - But implementation code is short
 - Little busy waiting if critical section rarely occupied
3. Note that applications may spend lots of time in critical sections and therefore this is not a good solution.

Semaphore Implementation with no Busy waiting

1. With each semaphore there is an associated waiting queue. Each entry in a waiting queue has two data items:

value (of type integer)

pointer to next record in the list

2. Two operations:

- block – place the process invoking the operation on the appropriate waiting queue.
- wakeup – remove one of processes in the waiting queue and place it in the ready queue.

->Implementation of wait:

```
wait (S){                                value--;                                if (value < 0) {  
    add this process to waiting queue    block(); }  
}
```

->Implementation of signal:

```
Signal (S){                                value++;                                if (value <= 0) {  
    remove a process P from the waiting queue    wakeup(P); }  
}
```

Synchronization Hardware

1. Many systems provide hardware support for critical section code
2. Uniprocessors – could disable interrupts
 - Currently running code would execute without preemption
 - Generally too inefficient on multiprocessor systems
3. Modern machines provide special atomic hardware instructions ->Atomic = non-interruptable

☐ Either test memory word and set value

☐ Or swap contents of two memory words

Classical Problems of Synchronization

1. Bounded-Buffer Problem
2. Readers and Writers Problem
3. Dining-Philosophers Problem

Bounded-Buffer Problem

1. N buffers, each can hold one item
2. Semaphore mutex initialized to the value 1
3. Semaphore full initialized to the value 0
4. Semaphore empty initialized to the value N .
5. The structure of the producer process

```
// produce an item    while (true) {  
                        wait (empty);    wait (mutex);
```

```

// add the item to the buffer          signal (mutex);
signal (full);
}
6. The structure of the consumer process      while (true) {          wait (full);
    wait (mutex);
    // remove an item from buffer
    signal (mutex);          signal (empty);
    // consume the removed item
}

```

Readers-Writers Problem

1. A data set is shared among a number of concurrent processes $\Rightarrow$ Readers – only read the data set; they do not perform any updates $\Rightarrow$ Writers – can both read and write.
2. Problem – allow multiple readers to read at the same time. Only one single writer can access the shared data at the same time.

3. Shared Data

- Data set
- Semaphore mutex initialized to 1. $\Rightarrow$ Semaphore wrt initialized to 1.
- Integer readcount initialized to 0.

```

4. The structure of a writer process      while (true) {          wait (wrt) ;
// writing is performed          signal (wrt) ;
}
5. The structure of a reader process
while (true) {          wait (mutex) ;          readcount ++ ;
if (readcount == 1) wait (wrt) ;          signal (mutex)          // reading is performed
    wait (mutex) ;          readcount -- ;
if (readcount == 0) signal (wrt) ;
signal (mutex) ;
}

```

Dining-Philosophers Problem



1. Shared data
 - Bowl of rice (data set)
 - Semaphore chopstick [5] initialized to 1
2. The structure of Philosopher i :


```
While (true) {
```

```

wait ( chopstick[i] );
wait ( chopstick[ (i + 1) % 5] );
// eat    signal ( chopstick[i] );
signal ( chopstick[ (i + 1) % 5] );
// think
}

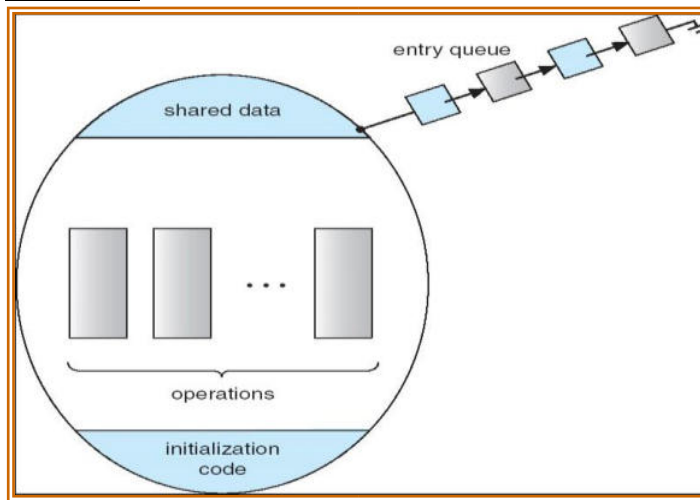
```

Problems with Semaphores

1. Correct use of semaphore operations:

- signal $\emptyset$ mutex $\succ$ wait $\emptyset$ mutex $\succ$
- wait $\emptyset$ mutex $\succ$... wait $\emptyset$ mutex $\succ$
- Omitting of wait (mutex) or signal (mutex) (or both)

Monitors



1. high-level abstraction that provides a convenient and effective mechanism for process synchronization
2. Only one process may be active within the monitor at a time
 - a. monitor monitor-name
 - i. {
 - b. // shared variable declarations
 - c. procedure P $\emptyset$... $\succ$ { }
 - i. ...
 - d. procedure Pn $\emptyset$... $\succ$ {.....}
3. Initialization code $\emptyset$ $\succ$ { ... }
- a. ...
- b. }
4. }

5. Solution to Dining Philosophers monitor DP

```

{
enum { THINKING; HUNGRY, EATING) state [5];    condition self [5];

void pickup (int i) {      state[i] = HUNGRY;      test(i);
if (state[i] != EATING) self [i].wait;
}
void putdown (int i) {
state[i] = THINKING;
// test left and right neighbors      test((i + 4) % 5);      test((i + 1) % 5);
}
}

```

```

void test (int i) {
if ( (state[(i + 4) % 5] != EATING) &&          (state[i] == HUNGRY) &&
(state[(i + 1) % 5] != EATING) ) {
state[i] = EATING ;          self[i].signal () ;
}
}
initialization_code() {          for (int i = 0; i < 5; i++)
state[i] = THINKING;
}
}

```

->Each philosopher *I* invokes the operations pickup() and putdown() in the following sequence:

```

dp.pickup (i)          EAT
dp.putdown (i)

```

Monitor Implementation Using Semaphores

- Variables
 - semaphore mutex; // (initially = 1)
 - semaphore next; // (initially = 0)
 - int next-count = 0;
- 2.Each procedure *F* will be replaced by
- wait(mutex);
- ...
- $\square$ body of *F*;
- ...
- if (next-count > 0) o signal(next)
- else o signal(mutex);
- Mutual exclusion within a monitor is ensured.

- For each condition variable x , we have:
 - semaphore $x\text{-sem}$; // (initially = 0)
 - int $x\text{-count} = 0$;
- The operation $x.\text{wait}$ can be implemented as:
 - $x\text{-count}++$;
 - if ($x\text{-count} > 0$)
- signal(x);
- else
- signal($x\text{-sem}$);
- wait($x\text{-sem}$);
- $x\text{-count}--$;
- The operation $x.\text{signal}$ can be implemented as:
 - if ($x\text{-count} > 0$) {
 - next-count++;
 - signal($x\text{-sem}$);
 - wait(next);
 - next-count--;
 - }

Producer-Consumer Problem Using Semaphores

The Solution to producer-consumer problem uses three semaphores, namely, full, empty and mutex.

The semaphore 'full' is used for counting the number of slots in the buffer that are full. The 'empty' for counting the number of slots that are empty and semaphore 'mutex' to make sure that the producer and consumer do not access modifiable shared section of the buffer simultaneously.

Initialization

- Set full buffer slots to 0.
i.e., semaphore Full = 0.
- Set empty buffer slots to N.
i.e., semaphore empty = N.
- For control access to critical section set mutex to 1.
i.e., semaphore mutex = 1.

Producer ()

WHILE (true) produce-Item ();

P (empty); P (mutex); enter-Item ()

V (mutex)

V (full);

Consumer ()

WHILE (true)

P (full) P (mutex); remove-Item (); V (mutex); V (empty);
consume-Item (Item)

- When processes request a resource and if the resources are not available at that time the process enters into waiting state. Waiting process may not change its state because the resources they are requested are held by other process. This situation is called deadlock.
- The situation where the process waiting for the resource i.e., not available is called deadlock.

System Model:-

➔ A system may consist of finite number of resources and is distributed among number of processes. These resources are partitioned into several instances each with identical instances.

➔ A process must request a resource before using it and it must release the resource after using it. It can request any number of resources to carry out a designated task. The amount of resource requested may not exceed the total number of resources available. ☐ A process may utilize the resources in only the following sequences:-

1. Request:- If the request is not granted immediately then the requesting process must wait it can acquire the resources.
2. Use:- The process can operate on the resource.
3. Release:- The process releases the resource after using it.

➔ Deadlock may involve different types of resources.

For eg:- Consider a system with one printer and one tape drive. If a process P_i currently holds a printer and a process P_j holds the tape drive. If process P_i request a tape drive and process P_j request a printer then a deadlock occurs.

Multithread programs are good candidates for deadlock because they compete for shared resources.

Deadlock Characterization:-

Necessary Conditions:-

A deadlock situation can occur if the following 4 conditions occur simultaneously in a system:-

1. Mutual Exclusion:-

Only one process must hold the resource at a time. If any other process requests for the resource, the requesting process must be delayed until the resource has been released.

2. Hold and Wait:-

A process must be holding at least one resource and waiting to acquire additional resources that are currently being held by the other process.

3. No Preemption:-

Resources can't be preempted i.e., only the process holding the resources must release it after the process has completed its task.

4. Circular Wait:-

A set $\{P_1, P_2, \dots, P_n\}$ of waiting process must exist such that P_1 is waiting for a resource i.e., held by P_2 , P_2 is waiting for a resource i.e., held by P_3 . P_{n-1} is waiting for resource held by process P_n and P_n is waiting for the resource i.e., held by P_1 .

All the four conditions must hold for a deadlock to occur.

Resource Allocation Graph:-

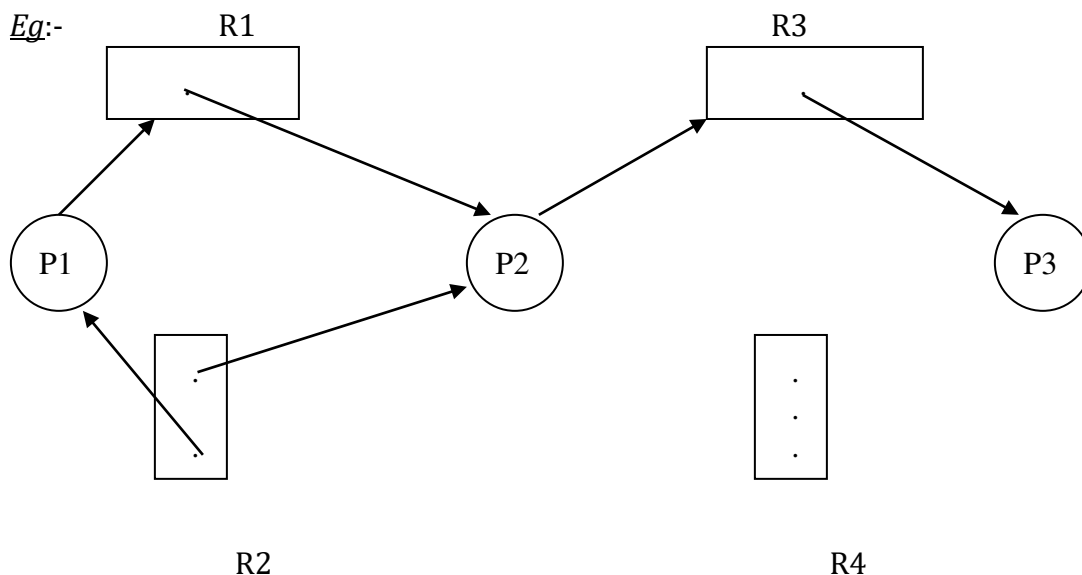
→ Deadlocks are described by using a directed graph called system resource allocation graph. The graph consists of set of vertices (v) and set of edges (e).

→ The set of vertices (v) can be described into two different types of nodes $P = \{P_1, P_2, \dots, P_n\}$ i.e., set consisting of all active processes and $R = \{R_1, R_2, \dots, R_n\}$ i.e., set consisting of all resource types in the system.

→ A directed edge from process P_i to resource type R_j denoted by $P_i \rightarrow R_j$ indicates that P_i requested an instance of resource R_j and is waiting. This edge is called Request edge.

→ A directed edge $R_i \rightarrow P_j$ signifies that resource R_j is held by process P_i . This is called Assignment edge.

Eg:-



→ If the graph contains no cycle, then no process in the system is in a deadlock. If the graph contains a cycle then a deadlock may exist.

→ If each resource type has exactly one instance then a cycle implies that a deadlock has occurred. If each resource has several instances then a cycle does not necessarily imply that a deadlock has occurred.

Methods for Handling Deadlocks:-

There are three ways to deal with deadlock problem

We can use a protocol to prevent deadlocks ensuring that the system will never enter into the deadlock state.

We allow a system to enter into deadlock state, detect it and recover from it.

We ignore the problem and pretend that the deadlock never occur in the system.

This is used by most OS including UNIX.

- ➔ To ensure that the deadlock never occur the system can use either deadlock avoidance or a deadlock prevention.
- ➔ Deadlock prevention is a set of method for ensuring that at least one of the necessary conditions does not occur.
- ➔ Deadlock avoidance requires the OS is given advance information about which resource a process will request and use during its lifetime.
- ➔ If a system does not use either deadlock avoidance or deadlock prevention then a deadlock situation may occur. During this it can provide an algorithm that examines the state of the system to determine whether a deadlock has occurred and algorithm to recover from deadlock.
- ➔ Undetected deadlock will result in deterioration of the system performance.

Deadlock Prevention:-

- For a deadlock to occur each of the four necessary conditions must hold. If at least one of the there condition does not hold then we can prevent occurrence of deadlock.

1. Mutual Exclusion:-

This holds for non-sharable resources.

Eg:- A printer can be used by only one process at a time.

Mutual exclusion is not possible in sharable resources and thus they cannot be involved in deadlock. Read-only files are good examples for sharable resources. A process never waits for accessing a sharable resource. So we cannot prevent deadlock by denying the mutual exclusion condition in non-sharable resources.

2. Hold and Wait:-

This condition can be eliminated by forcing a process to release all its resources held by it when it request a resource i.e., not available.

- One protocol can be used is that each process is allocated with all of its resources before its start execution.

Eg:- consider a process that copies the data from a tape drive to the disk, sorts the file and then prints the results to a printer. If all the resources are allocated at the beginning then the tape drive, disk files and printer are assigned to the process. The main problem with this is it leads to low resource utilization because it requires printer at the last and is allocated with it from the beginning so that no other process can use it.

- Another protocol that can be used is to allow a process to request a resource when the process has none. i.e., the process is allocated with tape drive and disk file. It performs the required operation and releases both. Then the process once again request for disk file and the printer and the problem and with this is starvation is possible.

3. No Preemption:-

To ensure that this condition never occurs the resources must be preempted. The following protocol can be used.

- If a process is holding some resource and request another resource that cannot be immediately allocated to it, then all the resources currently held by the requesting process are preempted and added to the list of resources for which other processes may be waiting. The process will be restarted only when it regains the old resources and the new resources that it is requesting.
- When a process request resources, we check whether they are available or not. If they are available we allocate them else we check that whether they are allocated to some other waiting process. If so we preempt the resources from the waiting process and allocate them to the requesting process. The requesting process must wait.

4. Circular Wait:-

The fourth and the final condition for deadlock is the circular wait condition. One way to ensure that this condition never, is to impose ordering on all resource types and each process requests resource in an increasing order.

Let $R = \{R_1, R_2, \dots, R_n\}$ be the set of resource types. We assign each resource type with a unique integer value. This will allows us to compare two resources and determine whether one precedes the other in ordering. *Eg:-* we can define a one to one function

$F: R \rightarrow N$ as follows :- $F(\text{disk drive})=5$

$F(\text{printer})=12$

$F(\text{tape drive})=1$

- Deadlock can be prevented by using the following protocol:-
- Each process can request the resource in increasing order. A process can request any number of instances of resource type say R_i and it can request instances of resource type R_j only $F(R_j) > F(R_i)$.
- Alternatively when a process requests an instance of resource type R_j , it has released any resource R_i such that $F(R_i) \geq F(R_j)$.
- If these two protocol are used then the circular wait can't hold.

Deadlock Avoidance:-

- ➔ Deadlock prevention algorithm may lead to low device utilization and reduces system throughput.
- ➔ Avoiding deadlocks requires additional information about how resources are to be requested. With the knowledge of the complete sequences of requests and releases we can decide for each requests whether or not the process should wait.
- ➔ For each requests it requires to check the resources currently available, resources that are currently allocated to each processes future requests and release of each process to decide whether the current requests can be satisfied or must wait to avoid future possible deadlock.
- ➔ A deadlock avoidance algorithm dynamically examines the resources allocation state to ensure that a circular wait condition never exists. The resource allocation state is defined by the number of available and allocated resources and the maximum demand of each process.

Safe State:-

- ➔ A state is a safe state in which there exists at least one order in which all the process will run completely without resulting in a deadlock.
- ➔ A system is in safe state if there exists a safe sequence.
- ➔ A sequence of processes $\langle P_1, P_2, \dots, P_n \rangle$ is a safe sequence for the current allocation state if for each P_i the resources that P_i can request can be satisfied by the currently available resources.
- ➔ If the resources that P_i requests are not currently available then P_i can obtain all of its needed resource to complete its designated task.
- ➔ A safe state is not a deadlock state.
- ➔ Whenever a process request a resource i.e., currently available, the system must decide whether resources can be allocated immediately or whether the process must wait. The request is granted only if the allocation leaves the system in safe state.
- ➔ In this, if a process requests a resource i.e., currently available it must still have to wait. Thus resource utilization may be lower than it would be without a deadlock avoidance algorithm.

Resource Allocation Graph Algorithm:-

- ➔ This algorithm is used only if we have one instance of a resource type. In addition to the request edge and the assignment edge a new edge called claim edge is used. *For eg:-* A claim edge $P_i \rightarrow R_j$ indicates that process P_i may request R_j in future. The claim edge is represented by a dotted line.
- ➔ When a process P_i requests the resource R_j , the claim edge is converted to a request edge.
- ➔ When resource R_j is released by process P_i , the assignment edge $R_j \rightarrow P_i$ is replaced by the claim edge $P_i \rightarrow R_j$.
- ➔ When a process P_i requests resource R_j the request is granted only if converting the request edge $P_i \rightarrow R_j$ to as assignment edge $R_j \rightarrow P_i$ do not result in a cycle. Cycle detection algorithm is used to detect

the cycle. If there are no cycles then the allocation of the resource to process leave the system in safe state

Banker's Algorithm:-

→ This algorithm is applicable to the system with multiple instances of each resource types, but this is less efficient than the resource allocation graph algorithm.

→ When a new process enters the system it must declare the maximum number of resources that it may need. This number may not exceed the total number of resources in the system. The system must determine that whether the allocation of the resources will leave the system in a safe state or not. If it is so resources are allocated else it should wait until the process release enough resources.

→ Several data structures are used to implement the banker's algorithm. Let n be the number of processes in the system and m be the number of resource types. We need the following data structures:-

•**Available:-** A vector of length m indicates the number of available resources. If $Available[i]=k$, then k instances of resource type R_j is available.

•**Max:-** An $n \times m$ matrix defines the maximum demand of each process if $Max[i,j]=k$, then P_i may request at most k instances of resource type R_j .

•**Allocation:-** An $n \times m$ matrix defines the number of resources of each type currently allocated to each process. If $Allocation[i,j]=k$, then P_i is currently k instances of resource type R_j .

•**Need:-** An $n \times m$ matrix indicates the remaining resources need of each process. If $Need[i,j]=k$, then P_i may need k more instances of resource type R_j to complete its task. So $Need[i,j]=Max[i,j]-Allocation[i,j]$

Safety Algorithm:-

→ This algorithm is used to find out whether or not a system is in safe state or not.

Step 1. Let work and finish be two vectors of length M and N respectively.

Initialize work = available and

Finish[i]=false for $i=1, 2, 3, \dots, n$

Step 2. Find i such that both

Finish[i]=false

Need $i \leq \text{work}$

If no such i exist then go to step 4

Step 3. $\text{Work} = \text{work} + \text{Allocation}$

$\text{Finish}[i] = \text{true}$

Go to step 2

Step 4. If $\text{finish}[i] = \text{true}$ for all i , then the system is in safe state.

This algorithm may require an order of $m \cdot n \cdot n$ operation to decide whether a state is safe.

Resource Request Algorithm:-

Let $\text{Request}(i)$ be the request vector of process P_i . If $\text{Request}(i)[j] = k$, then process P_i wants k instances of the resource type R_j . When a request for resources is made by process P_i the following actions are taken.

- If $\text{Request}(i) \leq \text{Need}(i)$ go to step 2 otherwise raise an error condition since the process has exceeded its maximum claim.
- If $\text{Request}(i) \leq \text{Available}$ go to step 3 otherwise P_i must wait. Since the resources are not available.
- If the system want to allocate the requested resources to process P_i then modify the state as follows.
 $\text{Available} = \text{Available} - \text{Request}(i)$
 $\text{Allocation}(i) = \text{Allocation}(i) + \text{Request}(i)$
 $\text{Need}(i) = \text{Need}(i) - \text{Request}(i)$
- If the resulting resource allocation state is safe, the transaction is complete and P_i is allocated its resources. If the new state is unsafe then P_i must wait for $\text{Request}(i)$ and old resource allocation state is restored.

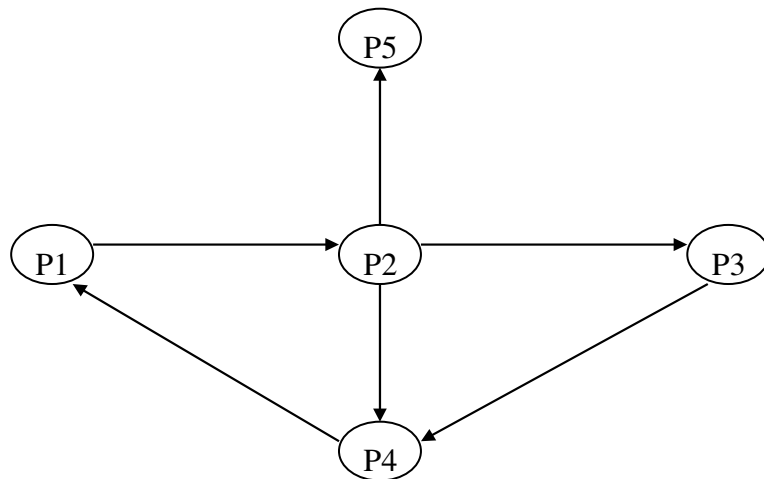
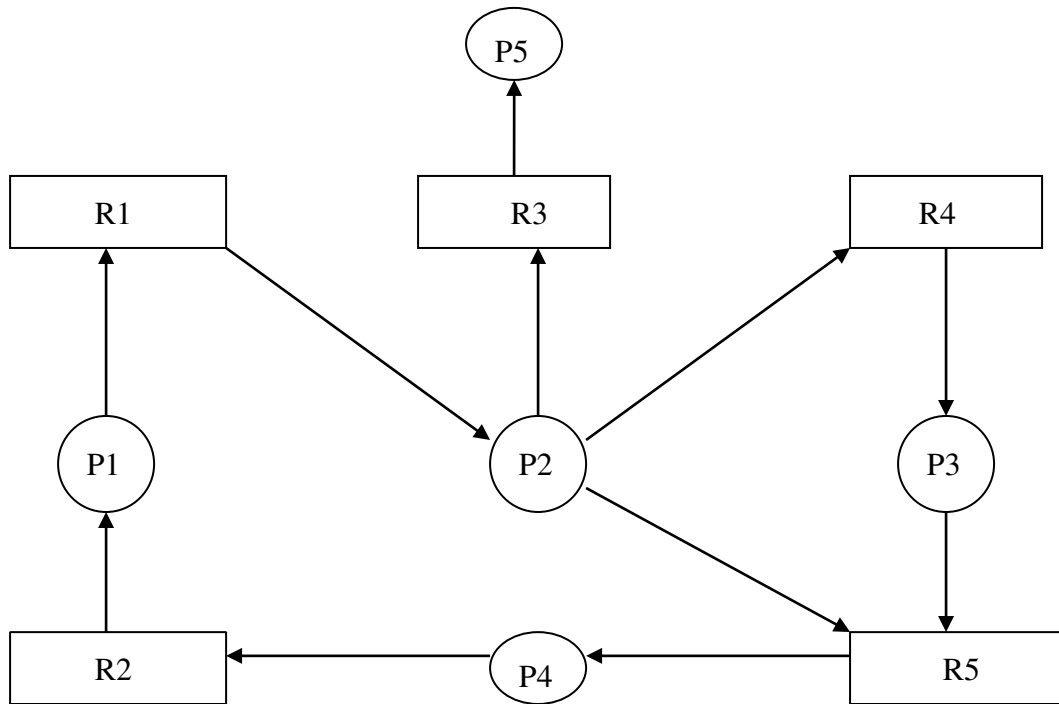
Deadlock Detection:-

If a system does not employ either deadlock prevention or a deadlock avoidance algorithm then a deadlock situation may occur. In this environment the system may provide

- An algorithm that examines the state of the system to determine whether a deadlock has occurred.

- An algorithm to recover from the deadlock.

Single Instances of each Resource Type:-



Several Instances of a Resource Types:-

- The wait for graph is applicable to only a single instance of a resource type. The following algorithm applies if there are several instances of a resource type. The following data structures are used:-

Available:-Is a vector of length m indicating the number of available resources of each type .

Allocation:-Is an $m \times n$ matrix which defines the number of resources of each type currently allocated to each process.

Request:-Is an $m \times n$ matrix indicating the current request of each process. If $\text{request}[i,j]=k$ then P_i is requesting k more instances of resources type R_j .

Step 1. let work and finish be vectors of length m and n respectively. Initialize Work = available/expression

For $i=1,2,\dots,n$ if allocation $P_i \neq 0$ then $\text{Finish}[i]=0$ else

$\text{Finish}[i]=\text{true}$

Step 2. Find an index(i) such that both

$\text{Finish}[i] = \text{false}$

$\text{Request}(i) \leq \text{work}$

If no such i exist go to step 4.

Step 3. $\text{Work} = \text{work} + \text{Allocation}(i)$

$\text{Finish}[i] = \text{true}$ Go to step 2.

Step 4. If $\text{Finish}[i] = \text{false}$ for some i where $m \geq i \geq 1$.

When a system is in a deadlock state.

This algorithm needs an order of $m \times n$ square operations to detect whether the system is in deadlock state or not.

Example Problem:-

1. For the following snapshot of the system find the safe sequence ϕ using Banker's algorithm).

Process	Allocation			Max			Available		
	R1	R2	R3	R1	R2	R3	R1	R2	R3
P1	0	1	0	7	5	3	3	3	2
P2	2	0	0	3	2	2			
P3	3	0	2	9	0	2			
P4	2	1	1	2	2	2			
P5	0	0	2	4	3	2			

Calculate the need of each process?

To find safe sequence?

2. Consider the following snapshot of the system and answer the following questions using Banker's algorithm?

a. Find the need of the allocation?

b. Is the system is in safe state?

c. If the process P1 request (0,4,2,0) resources can the request be granted immediately?

Process		Allocation		Max				Available							
A	B	C	D	A	B	C	D	A	B	C	D	P1	0	0	1
1	2	1	5	2								2	0	0	0
								0	P2	1		0	0	0	1
0															
	P3	1	3	5	4	2		3	5	6					
	P4	0	6	3	2	0		6	5	2					
	P5	0	0	1	4	0		6	5	6					

3. The operating system contains three resources. The numbers of instances of each resource type are (7, 7, 10). The current allocation state is given below.

a. Is the current allocation is safe?

b. find need?

c. Can the request made by the process P1(1,1,0) can be granted?

Process		Allocation			Max		
R1 R2 R3		R1	R2	R3			
P1		2	2	3	3	6	8
P2		2	0	3	4	3	3
P3		1	2	4	3	4	4

4. Explain different methods to recover from deadlock?

5. Write advantage and disadvantage of deadlock avoidance and deadlock prevention?

UNIT III

STORAGE MANAGEMENT

- Memory management is concerned with managing the primary memory.
- Memory consists of array of bytes or words each with their own address.
- The instructions are fetched from the memory by the cpu based on the value program counter.

Functions of memory management:-

- Keeping track of status of each memory location..
- Determining the allocation policy.
- Memory allocation technique.
- De-allocation technique.

Address Binding:-

- Programs are stored on the secondary storage disks as binary executable files.
- When the programs are to be executed they are brought in to the main memory and placed within a process.
- The collection of processes on the disk waiting to enter the main memory forms the input queue.
- One of the processes which are to be executed is fetched from the queue and placed in the main memory.
- During the execution it fetches instruction and data from main memory. After the process terminates it returns back the memory space.
- During execution the process will go through different steps and in each step the address is represented in different ways.
- In source program the address is symbolic.
- The compiler converts the symbolic address to re-locatable address.
- The loader will convert this re-locatable address to absolute address. Binding of instructions and data can be done at any step along the way:-

Compile time:-

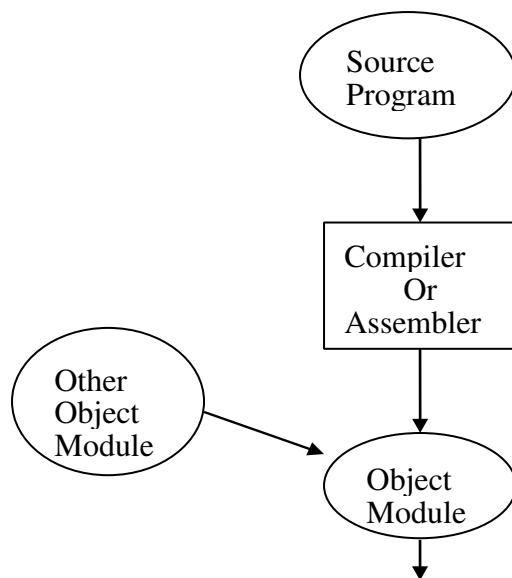
If we know whether the process resides in memory then absolute code can be generated. If the static address changes then it is necessary to re-compile the code from the beginning.

Load time:-

If the compiler doesn't know whether the process resides in memory then it generates the re-locatable code. In this the binding is delayed until the load time.

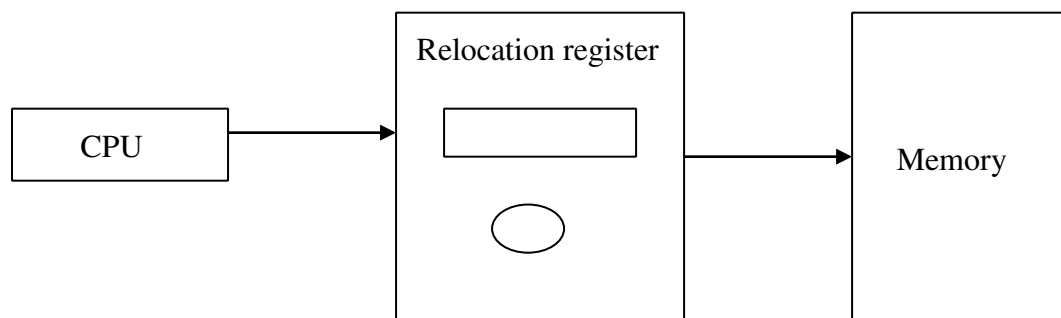
Execution time:-

If the process is moved during its execution from one memory segment to another then the binding is delayed until run time. Special hardware is used for this. Most of the general purpose operating system uses this method.



Logical versus physical address:-

- The address generated by the CPU is called logical address or virtual address.
- The address seen by the memory unit i.e., the one loaded in to the memory register is called the physical address.
- Compile time and load time address binding methods generate some logical and physical address.
- The execution time addressing binding generate different logical and physical address.
- Set of logical address space generated by the programs is the logical address space.
- Set of physical address corresponding to these logical addresses is the physical address space.
- The mapping of virtual address to physical address during run time is done by the hardware device called memory management unit (MMU).
- The base register is also called re-location register.
- Value of the re-location register is added to every address generated by the user process at the time it is sent to memory.



MMU

Dynamic re-location using a re-location registers

- The above figure shows that dynamic re-location which implies mapping from virtual addresses space to physical address space and is performed by the hardware at run time.
- Re-location is performed by the hardware and is invisible to the user dynamic relocation makes it possible to move a partially executed process from one area of memory to another without affecting.

Dynamic Loading:-

- For a process to be executed it should be loaded in to the physical memory. The size of the process is limited to the size of the physical memory.
- Dynamic loading is used to obtain better memory utilization.
- In dynamic loading the routine or procedure will not be loaded until it is called.
- Whenever a routine is called, the calling routine first checks whether the called routine is already loaded or not. If it is not loaded it cause the loader to load the desired program in to the memory and updates the programs address table to indicate the change and control is passed to newly called routine.

Advantage:-

- Gives better memory utilization.
- Unused routine is never loaded.
- Do not need special operating system support.
- This method is useful when large amount of codes are needed to handle in frequently occurring cases.

Dynamic linking and Shared libraries:-

Some operating system supports only the static linking.

In dynamic linking only the main program is loaded in to the memory. If the main program requests a procedure, the procedure is loaded and the link is established at the time of references. This linking is postponed until the execution time.

With dynamic linking a `†stub†` is used in the image of each library referenced routine. A `†stub†` is a piece of code which is used to indicate how to locate the appropriate memory resident library routine or how to load library if the routine is not already present.

When `†stub†` is executed it checks whether the routine is present in memory or not. If not it loads the routine into the memory.

This feature can be used to update libraries i.e., library is replaced by a new version and all the programs can make use of this library.

More than one version of the library can be loaded in memory at a time and each program uses its version of the library. Only the programs that are compiled with the new version are affected by the changes incorporated in it. Other programs linked before new version is installed will continue using older libraries. This type of system is called `†shared library†`.

Overlays:-

The size of the process is limited to the size of physical memory. If the size is more than the size of physical memory then a technique called overlays is used.

The idea is to load only those instructions and data that are needed at any given time. When other instructions are needed, they are loaded into memory space that was previously occupied by the instructions that are no longer needed.

Eg:-

Consider a 2-pass assembler where pass-1 generates a symbol table and pass-2 generates a machine code.

Assume that the sizes of components are as follows:

Pass-1 = 70k

Pass-2 = 80k

Symbol table = 20k

Common routine = 30k

To load everything at once, it requires 200k of memory. Suppose if 150k of memory is available, we can't run all the components at the same time.

Thus we define 2 overlays, overlay A which consists of symbol table, common routine and pass-1 and overlay B which consists of symbol table, common routine and pass-2.

We add an overlay driver and start overlay A in memory. When we finish pass-1 we jump to overlay driver, then the control is transferred to pass-2.

Thus we can run assembler in 150k of memory.

The code for overlays A and B are kept on disk as absolute memory images. Special re-location and linking algorithms are needed to construct the overlays. They can be implemented using simple file structures.

Swapping:-

➔ Swapping is a technique of temporarily removing inactive programs from the memory of the system.

➔ A process can be swapped temporarily out of the memory to a backing store and then brought back in to the memory for continuing the execution. This process is called swapping.

Eg:- In a multi-programming environment with a round robin CPU scheduling whenever the time quantum expires then the process that has just finished is swapped out and a new process swaps in to the memory for execution.

➔ A variation of swap is priority based scheduling. When a low priority is executing and if a high priority process arrives then a low priority will be swapped out and high priority is allowed for execution. This process is also called as Roll out and Roll in.

➔ Normally the process which is swapped out will be swapped back to the same memory space that is occupied previously. This depends upon address binding.

➔ If the binding is done at load time, then the process is moved to same memory location.

➔ If the binding is done at run time, then the process is moved to different memory location.

This is because the physical address is computed during run time.

➔ Swapping requires backing store and it should be large enough to accommodate the copies of all memory images.

➔ The system maintains a ready queue consisting of all the processes whose memory images are on the backing store or in memory that are ready to run.

➔ Swapping is constant by other factors:-

- To swap a process, it should be completely idle.
- A process may be waiting for an i/o operation. If the i/o is asynchronously accessing the user memory for i/o buffers, then the process cannot be swapped.

Contiguous Memory Allocation:-

• One of the simplest method for memory allocation is to divide memory in to several fixed partition. Each partition contains exactly one process. The degree of multiprogramming depends on the number of partitions.

• In multiple partition method, when a partition is free, process is selected from the input queue and is loaded in to free partition of memory.

• When process terminates, the memory partition becomes available for another process.

• Batch OS uses the fixed size partition scheme.

• The OS keeps a table indicating which part of the memory is free and is occupied.

• When the process enters the system it will be loaded in to the input queue. The OS keeps track of the memory requirement of each process and the amount of memory available and determines which process to allocate the memory.

• When a process requests, the OS searches for large hole for this process, hole is a large block of free memory available.

- If the hole is too large it is split in to two. One part is allocated to the requesting process and other is returned to the set of holes.

- The set of holes are searched to determine which hole is best to allocate. There are three strategies to select a free hole:-

- First bit:- Allocates first hole that is big enough. This algorithm scans memory from the beginning and selects the first available block that is large enough to hold the process.

- Best bit:- It chooses the hole i.e., closest in size to the request. It allocates the smallest hole i.e., big enough to hold the process.

- Worst fit:- It allocates the largest hole to the process request. It searches for the largest hole in the entire list.

➔ First fit and best fit are the most popular algorithms for dynamic memory allocation. First fit is generally faster. Best fit searches for the entire list to find the smallest hole i.e., large enough. Worst fit reduces the rate of production of smallest holes.

➔ All these algorithms suffer from fragmentation.

Memory Protection:-

- Memory protection means protecting the OS from user process and protecting process from one another.

- Memory protection is provided by using a re-location register, with a limit register.

- Re-location register contains the values of smallest physical address and limit register contains range of logical addresses. (Re-location = 100040 and limit = 74600).

- The logical address must be less than the limit register, the MMU maps the logical address dynamically by adding the value in re-location register.

- When the CPU scheduler selects a process for execution, the dispatcher loads the re-location and limit register with correct values as a part of context switch.

- Since every address generated by the CPU is checked against these register we can protect the OS and other users programs and data from being modified.

Fragmentation:-

☑ Memory fragmentation can be of two types:-

- Internal Fragmentation

- External Fragmentation

➔ In Internal Fragmentation there is wasted space internal to a portion due to the fact that block of data loaded is smaller than the partition.

Eg:- If there is a block of 50kb and if the process requests 40kb and if the block is allocated to the process then there will be 10kb of memory left.

➔ External Fragmentation exists when there is enough memory space exists to satisfy the request, but it not contiguous i.e., storage is fragmented in to large number of small holes.

- ➔ External Fragmentation may be either minor or a major problem.
- ➔ One solution for over-coming external fragmentation is compaction. The goal is to move all the free memory together to form a large block. Compaction is not possible always. If the re-location is static and is done at load time then compaction is not possible. Compaction is possible if the re-location is dynamic and done at execution time.
- ➔ Another possible solution to the external fragmentation problem is to permit the logical address space of a process to be non-contiguous, thus allowing the process to be allocated physical memory whenever the latter is available.

Paging:-

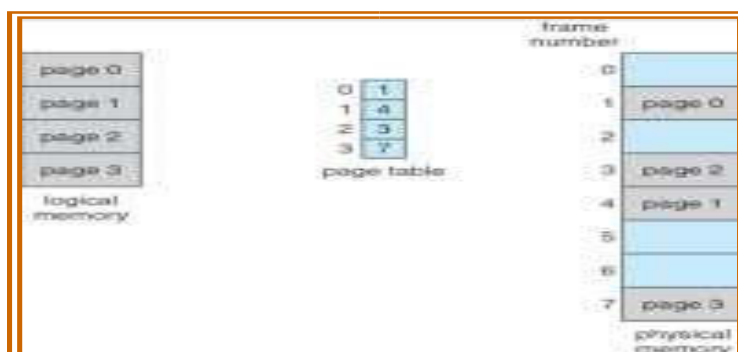
- Paging is a memory management scheme that permits the physical address space of a process to be non-contiguous. Support for paging is handled by hardware.
- It is used to avoid external fragmentation.
- Paging avoids the considerable problem of fitting the varying sized memory chunks on to the backing store.
- When some code or data residing in main memory need to be swapped out, space must be found on backing store.

Basic Method:-

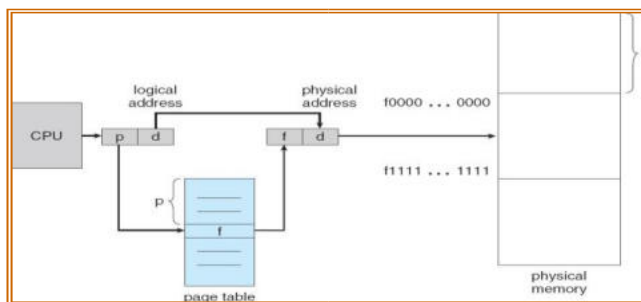
- Physical memory is broken in to fixed sized blocks called frames (f).
- Logical memory is broken in to blocks of same size called pages (p).
- When a process is to be executed its pages are loaded in to available frames from backing store.
- The backing store is also divided in to fixed-sized blocks of same size as memory frames.
- The following figure shows paging hardware:-
- Logical address generated by the CPU is divided in to two parts: page number (p) and page offset (d).
- The page number (p) is used as index to the page table. The page table contains base address of each page in physical memory. This base address is combined with the page offset to define the physical memory i.e., sent to the memory unit.

- The page size is defined by hardware. The size of a power of 2, varying between 512 bytes and per page.

the
10Mb



- If the size of logical address space is 2^m address unit and page size is 2^n , then high order $m-n$ designates the page number and n low order bits represents page offset.



Eg:- To show how to map logical memory in to physical memory consider a page size of 4 bytes and physical memory of 32 bytes (8 pages).

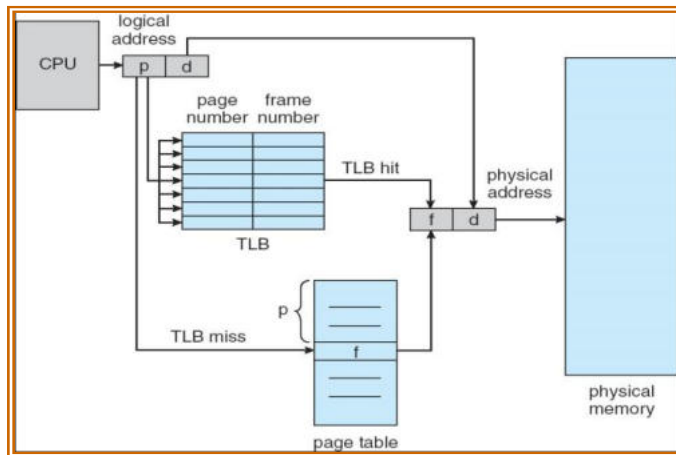
- Logical address 0 is page 0 and offset 0. Page 0 is in frame 5. The logical address 0 maps to physical address 20. $[(5 \times 4) + 0]$.
- Logical address 3 is page 0 and offset 3 maps to physical address 23 $[(5 \times 4) + 3]$.
- Logical address 4 is page 1 and offset 0 and page 1 is mapped to frame 6. So logical address 4 maps to physical address 24 $[(6 \times 4) + 0]$.
- Logical address 13 is page 3 and offset 1 and page 3 is mapped to frame 2. So logical address 13 maps to physical address 9 $[(2 \times 4) + 1]$.

Hardware Support for Paging:-

The hardware implementation of the page table can be done in several ways:-

1. The simplest method is that the page table is implemented as a set of dedicated registers. These registers must be built with very high speed logic for making paging address translation. Every accessed memory must go through paging map. The use of registers for page table is satisfactory if the page table is small.

2. If the page table is large then the use of registers is not visible. So the page table is kept in the main memory and a page table base register [PTBR] points to the page table. Changing the page table requires only one register which reduces the context switching type. The problem with this approach is the time required to access memory location. To access a location [i] first we have to index the page table using PTBR offset.



3. It gives the frame number which is combined with the page offset to produce the actual address. Thus we need two memory accesses for a byte.

4. The only solution is to use special, fast, lookup hardware cache called translation look aside buffer [TLB] or associative register.

- TLB is built with associative register with high speed memory. Each register contains two paths a key and a value.

- When an associative register is presented with an item, it is compared with all the key values, if found the corresponding value field is return and searching is fast. TLB is used with the page table as follows:- TLB contains only few page table entries.

- When a logical address is generated by the CPU, its page number along with the frame number is added to TLB. If the page number is found its frame memory is used to access the actual memory.

- If the page number is not in the TLB (TLB miss) the memory reference to the page table is made. When the frame number is obtained use can use it to access the memory.

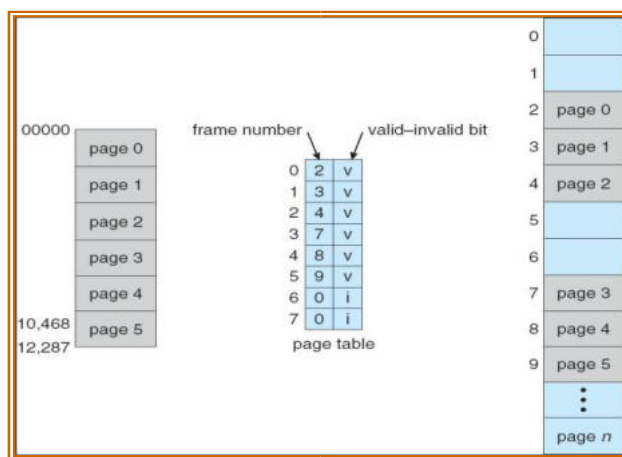
- If the TLB is full of entries the OS must select anyone for replacement.

- Each time a new page table is selected the TLB must be flushed [erased] to ensure that next executing process do not use wrong information.

- The percentage of time that a page number is found in the TLB is called HIT ratio.

Protection:-

- Memory protection in paged environment is done by protection bits that are associated with each frame these bits are kept in page table.
- One bit can define a page to be read-write or read-only.
- To find the correct frame number every reference to the memory should go through page table. At the same time physical address is computed.
- The protection bits can be checked to verify that no writers are made to read-only page.
- Any attempt to write in to read-only page causes a hardware trap to the OS.
- This approach can be used to provide protection to read-only, read-write or execute only pages.
- One more bit is generally added to each entry in the page table: a valid-invalid bit.



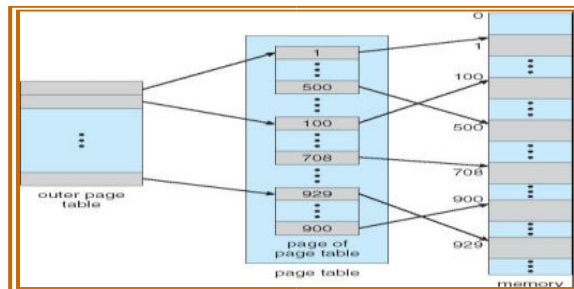
- A valid bit indicates that associated page is in the processes logical address space and thus it is a legal or valid page.
- If the bit is invalid, it indicates the page is not in the processes logical addressed space and illegal. Illegal addresses are trapped by using the valid-invalid bit.
- The OS sets this bit for each page to allow or disallow accesses to that page.

Structure of the Page Table:-

a. Hierarchical paging:-

- ➔ Recent computer system support a large logical address space from 2^{32} to 2^{64} . In this system the page table becomes large. So it is very difficult to allocate contiguous main memory for page table. One simple solution to this problem is to divide page table in to smaller pieces. There are several ways to accomplish this division.
- ➔ One way is to use two-level paging algorithm in which the page table itself is also paged.
Eg:- In a 32 bit machine with page size of 4kb. A logical address is divided in to a page number consisting of 20 bits and a page offset of 12 bit. The page table is further divided

since the page table is paged, the page number is further divided in to 10 bit page number and a 10 bit offset. So the logical address is



b. **Hashed page table** :-

➔ Hashed page table handles the address space larger than 32 bit. The virtual page number is used as hashed value. Linked list is used in the hash table which contains a list of elements that hash to the same location.

➔ Each element in the hash table contains the following three fields:-

- Virtual page number
- Mapped page frame value
- Pointer to the next element in the linked list

Working:-

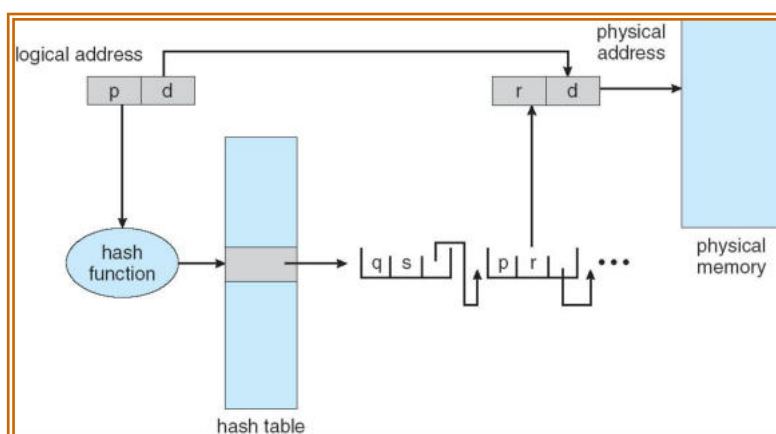
- Virtual page number is taken from virtual address.
- Virtual page number is hashed in to hash table.
- Virtual page number is compared with the first element of linked list.
- Both the values are matched, that value is (page frame) used for calculating the physical address.

- If not match then entire linked list is searched for matching virtual page number.

- Clustered pages are

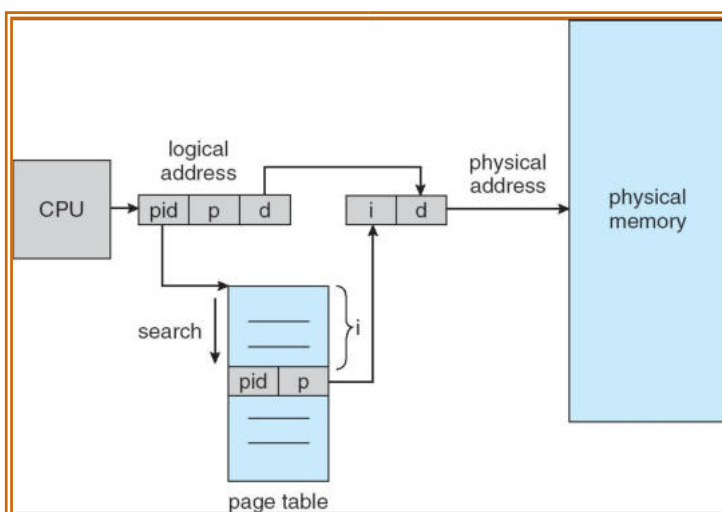
similar
is that
to

to hash table but one difference
each entity in the hash table refer
several pages.



c. Inverted Page Tables:-

- Since the address spaces have grown to 64 bits, the traditional page tables become a problem. Even with two level page tables. The table can be too large to handle.
- An inverted page table has only entry for each page in memory.
- Each entry consisted of virtual address of the page stored in that read-only location with information about the process that owns that page.
- Each virtual address in the Inverted page table consists of triple <process-id , page number , offset >.
- The inverted page table entry is a pair <process-id , page number>. When a memory reference is made, the part of virtual address i.e., <process-id , page number> is presented in to memory sub-system.
- The inverted page table is searched for a match.
- If a match is found at entry I then the physical address <i , offset> is generated. If no match is found then an illegal address access has been attempted.
- This scheme decreases the amount of memory needed to store each page table, it increases the amount of time needed to search the table when a page reference occurs. If the whole table is to be searched it takes too long.



Advantage:-

- Eliminates fragmentation.

- Support high degree of multiprogramming.
- Increases memory and processor utilization.
- Compaction overhead required for the re-locatable partition scheme is also eliminated.

Disadvantage:-

- Page address mapping hardware increases the cost of the computer.
- Memory must be used to store the various tables like page tables, memory map table etc.
- Some memory will still be unused if the number of available block is not sufficient for the address space of the jobs to be run.

Shared Pages:-

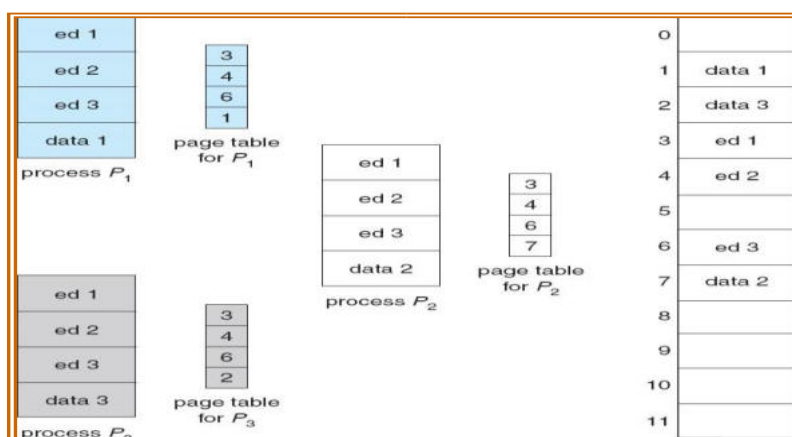
➔ Another advantage of paging is the possibility of sharing common code. This is useful in time-sharing environment.

Eg:- Consider a system with 40 users, each executing a text editor. If the text editor is of 150k and data space is 50k, we need 8000k for 40 users. If the code is reentrant it can be shared. Consider the following figure

➔ If the code is reentrant then it never changes during execution. Thus two or more processes can execute same code at the same time. Each process has its own copy of registers and the data of two processes will vary.

➔ Only one copy of the editor is kept in physical memory. Each users page table maps to same physical copy of editor but date pages are mapped to different frames.

➔ So to support 40 users we need only one copy of editor (150k) plus 40 copies of 50k of data space i.e., only 2150k instead of 8000k.



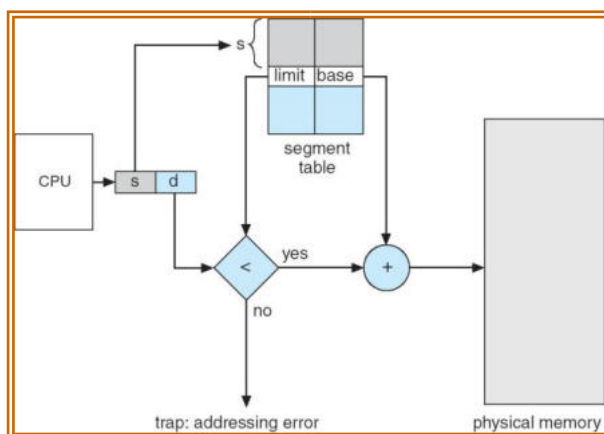
Segmentation:-

Basic method:-

- Most users do not think memory as a linear array of bytes rather the users think memory as a collection of variable sized segments which are dedicated to a particular use such as code, data, stack, heap etc.
- A logical address is a collection of segments. Each segment has a name and length. The address specifies both the segment name and the offset within the segments.
- The user specifies address by using two quantities: a segment name and an offset.
- For simplicity the segments are numbered and referred by a segment number. So the logical address consists of <segment number, offset>.

Hardware support:-

- We must define an implementation to map 2D user defined address in to 1D physical address.
- This mapping is affected by a segment table. Each entry in the segment table has a segment base and segment limit.
- The segment base contains the starting physical address where the segment resides and limit specifies the length of the segment.



The use of segment table is shown in the above figure:-

- Logical address consists of two parts: segment number s and an offset d to that segment.
- The segment number is used as an index to segment table.
- The offset d must be in between 0 and limit, if not an error is reported to OS.
- If legal the offset is added to the base to generate the actual physical address.
- The segment table is an array of base limit register pairs.

Protection and Sharing:-

- A particular advantage of segmentation is the association of protection with the segments.
- The memory mapping hardware will check the protection bits associated with each segment table entry to prevent illegal access to memory like attempts to write in to read-only segment.

- Another advantage of segmentation involves the sharing of code or data. Each process has a segment table associated with it. Segments are shared when the entries in the segment tables of two different processes points to same physical location.
- Sharing occurs at the segment table. Any information can be shared at the segment level. Several segments can be shared so a program consisting of several segments can be shared.
- We can also share parts of a program.

Advantages:-

- Eliminates fragmentation.
- Provides virtual growth.
- Allows dynamic segment growth.
- Assist dynamic linking.
- Segmentation is visible.

Differences between segmentation and paging:-

Segmentation:-

- Program is divided in to variable sized segments.
- User is responsible for dividing the program in to segments.
- Segmentation is slower than paging.
- Visible to user.
- Eliminates internal fragmentation.
- Suffers from external fragmentation.
- Process or user segment number, offset to calculate absolute address.

Paging:-

- Programs are divided in to fixed size pages.
- Division is performed by the OS.
- Paging is faster than segmentation.
- Invisible to user.
- Suffers from internal fragmentation.
- No external fragmentation.
- Process or user page number, offset to calculate absolute address.

Virtual memory

Virtual memory is a technique that allows for the execution of partially loaded process.

There are many advantages of this:-

- A program will not be limited by the amount of physical memory that is available user can able to write in to large virtual space.

- Since each program takes less amount of physical memory, more than one program could be run at the same time which can increase the throughput and CPU utilization.
 - Less i/o operation is needed to swap or load user program in to memory. So each user program could run faster.
- ➔ Virtual memory is the separation of users logical memory from physical memory. This separation allows an extremely large virtual memory to be provided when there is less physical memory.
- ➔ Separating logical memory from physical memory also allows files and memory to be shared by several different processes through page sharing.
- ➔ Virtual memory is implemented using Demand Paging.

Demand Paging:-

- A demand paging is similar to paging system with swapping when we want to execute a process we swap the process in to memory otherwise it will not be loaded in to memory.
- A swapper manipulates the entire processes, whereas a pager manipulates individual pages of the process.

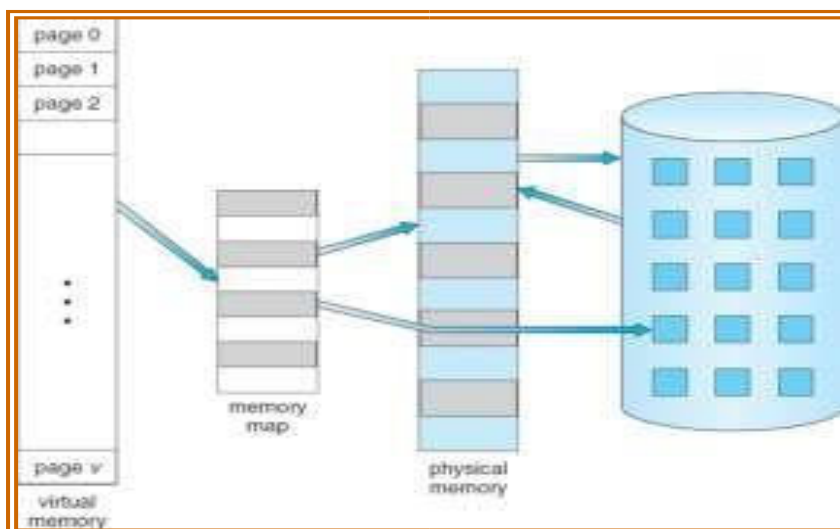
Basic concept:-

- Instead of swapping the whole process the pager swaps only the necessary pages in to memory. Thus it avoids reading unused pages and decreases the swap time and amount of physical memory needed.

The valid-invalid bit scheme can be used to distinguish between the pages that are on the disk and that are in memory.

- If the bit is valid then the page is both legal and is in memory.
- If the bit is invalid then either page is not valid or is valid but is currently on the disk.

Marking a page as invalid will have no effect if the processes never access to that page. Suppose

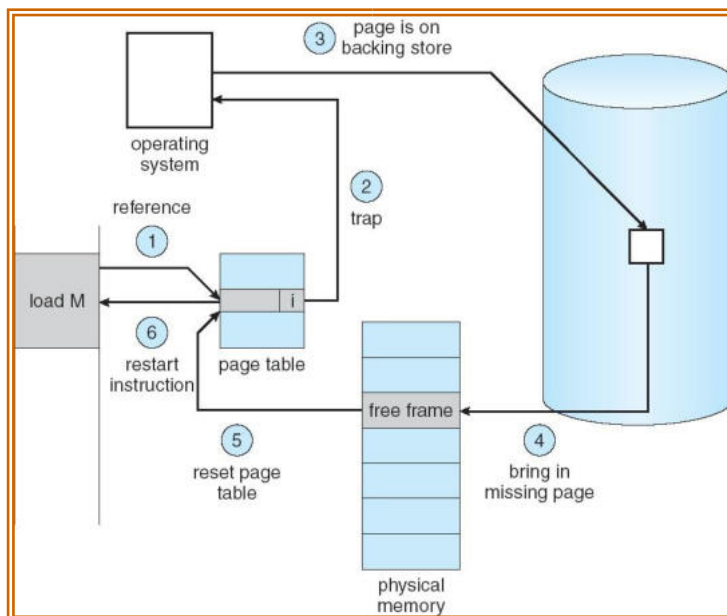


if it access the page which is marked invalid, causes a page fault trap. This may result in failure of OS to bring the desired page in to memory.

The step for handling page fault is straight forward and is given below:-

- We check the internal table of the process to determine whether the reference made is valid or invalid.
- If invalid terminate the process,. If valid, then the page is not yet loaded and we now page it in.
- We find a free frame.
- We schedule disk operation to read the desired page in to newly allocated frame.
- When disk read is complete, we modify the internal table kept with the process to indicate that the page is now in memory.
- We restart the instruction which was interrupted by illegal address trap. The process can now access the page.

In extreme cases, we start the process without pages in memory. When the OS points to the instruction of process it generates a page fault. After this page is brought in to memory the process continues to execute, faulting as necessary until every demand paging i.e., it never brings the page in to memory until it is required.



Hardware support:-

For demand paging the same hardware is required as paging and swapping.

- **Page table**:- Has the ability to mark an entry invalid through valid-invalid bit.
- **Secondary memory**:- This holds the pages that are not present in main memory.

It is a high speed disk.

Performance of demand paging:-

Demand paging can have significant effect on the performance of the computer system.

Let P be the probability of the page fault ($0 \leq P \leq 1$) Effective access time = $(1-P) * ma + P * \text{page fault}$.

Where P = page fault and ma = memory access time.

Effective access time is directly proportional to page fault rate. It is important to keep page fault rate low in demand paging.

A page fault causes the following sequence to occur:-
1. Trap to the OS.

- Save the user registers and process state.
- Determine that the interrupt was a page fault.
- Checks the page references were legal and determine the location of page on disk.
- Issue a read from disk to a free frame.
- If waiting, allocate the CPU to some other user.
- Interrupt from the disk.
- Save the registers and process states of other users.
- Determine that the interrupt was from the disk.
- Correct the page table and other table to show that the desired page is now in memory.
- Wait for the CPU to be allocated to this process again.
- Restore the user register process state and new page table then resume the interrupted instruction.

Comparison of demand paging with segmentation:-

Segmentation:-

- Segment may of different size.
- Segment can be shared.
- Allows for dynamic growth of segments.
- Segment map table indicate the address of each segment in memory.
- Segments are allocated to the program while compilation.

Demand Paging:-

- Pages are of same size.
- Pages can't be shared.
- Page size is fixed.
- Page table keeps track of pages in memory.
- Pages are allocated in memory on demand.

Process creation:-

a. Copy-on-write:-

Demand paging is used when reading a file from disk in to memory. Fork () is used to create a process and it initially bypass the demand paging using a technique called page sharing. Page sharing provides rapid speed for process creation and reduces the number of pages allocated to the newly created process.

Copy-on-write technique initially allows the parent and the child to share the same pages. These pages are marked as copy-on-write pages i.e., if either process writes to a shared page, a copy of shared page is created.

Eg:- If a child process try to modify a page containing portions of the stack; the OS recognizes them as a copy-on-write page and create a copy of this page and maps it on to the address space of the child process. So the child process will modify its copied page and not the page belonging to parent.

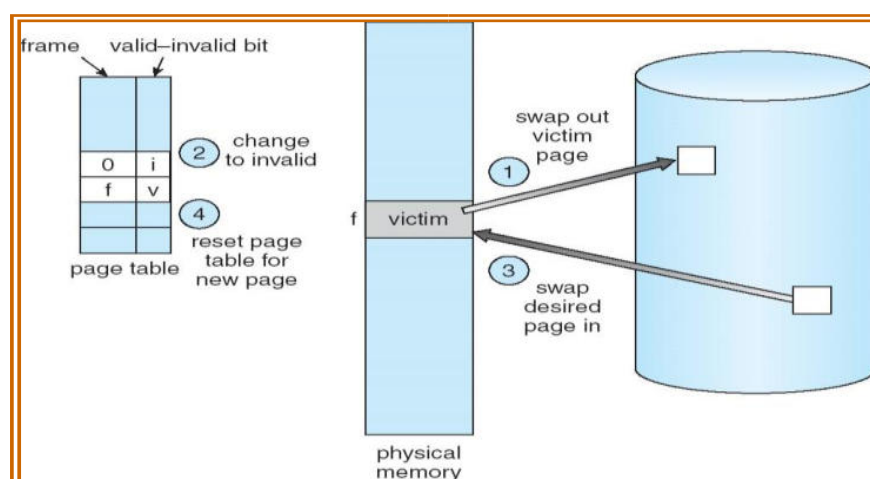
The new pages are obtained from the pool of free pages.

b. Memory Mapping:-

Standard system calls i.e., open (), read () and write () is used for sequential read of a file. Virtual memory is used for this. In memory mapping a file allows a part of the virtual address space to be logically associated with a file. Memory mapping a file is possible by mapping a disk block to page in memory.

Page Replacement

- ➔ Demand paging shares the I/O by not loading the pages that are never used.
- ➔ Demand paging also improves the degree of multiprogramming by allowing more process to run at the same time.



Working of Page Replacement Algorithm

- ➔ Page replacement policy deals with the solution of pages in memory to be replaced by a new page that must be brought in. When a user process is executing a page fault occurs.

- ➔ The hardware traps to the operating system, which checks the internal table to see that this is a page fault and not an illegal memory access.
- ➔ The operating system determines where the derived page is residing on the disk, and this finds that there are no free frames on the list of free frames.
- ➔ When all the frames are in main memory, it is necessary to bring a new page to satisfy the page fault, replacement policy is concerned with selecting a page currently in memory to be replaced.
- ➔ The page i.e. to be removed should be the page i.e. least likely to be referenced in future.

1. Find the location of derived page on the disk.
2. Find a free frame
 - If there is a free frame, use it.
 - Otherwise, use a replacement algorithm to select a victim.
 - Write the victim page to the disk; change the page and frame tables accordingly.
3. Read the desired page into the free frame; change the page and frame tables.
4. Restart the user process.

Victim Page

The page that is swapped out of physical memory is called victim page.

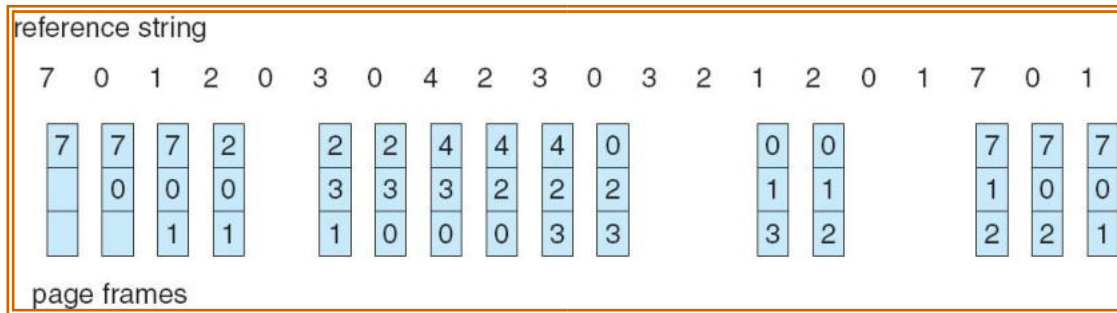
- If no frames are free, the two page transforms come (out and one in) are read. This will see the effective access time.
- Each page or frame may have a dirty (modify) bit associated with the hardware. The modify bit for a page is set by the hardware whenever any word or byte in the page is written into, indicating that the page has been modified.
- When we select the page for replacement, we check its modify bit. If the bit is set, then the page is modified since it was read from the disk.
- If the bit was not set, the page has not been modified since it was read into memory. Therefore, if the copy of the page has not been modified we can avoid writing the memory page to the disk, if it is already there. Some pages cannot be modified.
- We must solve two major problems to implement demand paging: we must develop a frame allocation algorithm and a page replacement algorithm. If we have multiple processors in memory, we must decide how many frames to allocate and page replacement is needed.

Page replacement Algorithms FIFO Algorithm:

- This is the simplest page replacement algorithm. A FIFO replacement algorithm associates each page the time when that page was brought into memory.
- When a Page is to be replaced the oldest one is selected.

- We replace the queue at the head of the queue. When a page is brought into memory, we insert it at the tail of the queue.

Example: Consider the following references string with frames initially empty.



- The first three references (7,0,1) cases page faults and are brought into the

empty frames.

- The next references 2 replaces page 7 because the page 7 was brought in first.
- Since 0 is the next references and 0 is already in memory e has no page faults.
- The next references 3 results in page 0 being replaced so that the next references to 0 causer page fault.

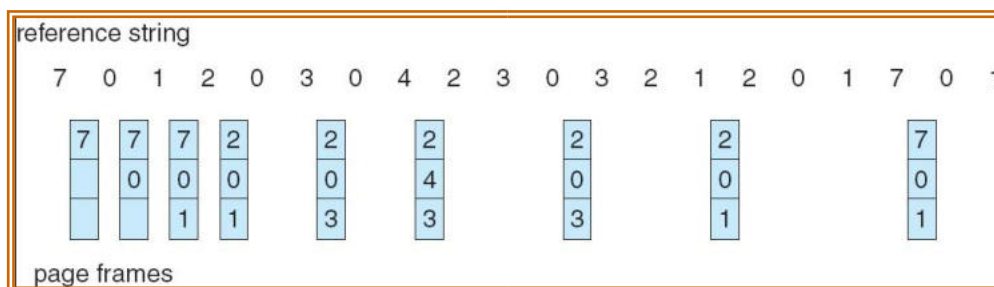
This will continue till the end of string.

There are 15 faults all together.

Belady's Anamoly

For some page replacement algorithm, the page fault may increase as the number of allocated frames increases. FIFO replacement algorithm may face this problem.

Optimal Algorithm



- Optimal page replacement algorithm is mainly to solve the problem of Belady's Anamoly.

- Optimal page replacement algorithm has the lowest page fault rate of all algorithms.
- An optimal page replacement algorithm exists and has been called OPT.
- The working is simple †Replace the page that will not be used for the longest period of time

Example: consider the following reference string

- The first three references cause faults that fill the three empty frames.
- The references to page 2 replaces page 7, because 7 will not be used until reference 18.
- The page 0 will be used at 5 and page 1 at 14.
- With only 9 page faults, optimal replacement is much better than a FIFO, which had 15 faults.

This algorithm is difficult to implement because it requires future knowledge of reference strings.

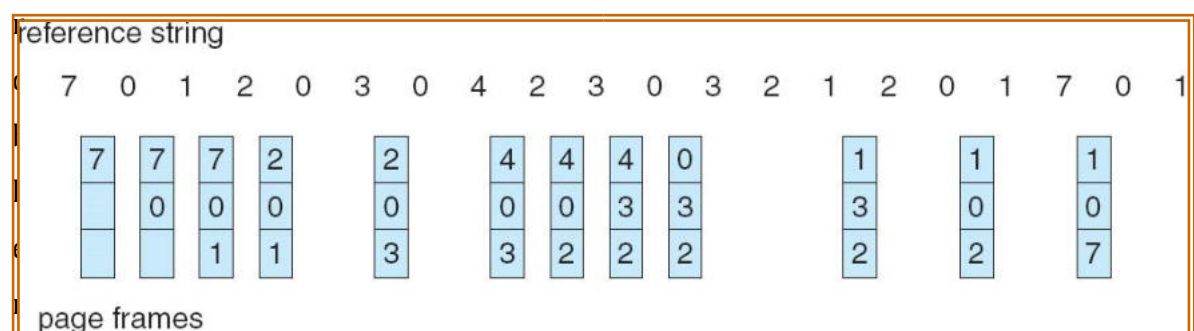
Least Recently Used (LRU) Algorithm

If the optimal algorithm is not feasible, an approximation to the optimal algorithm is possible.

The main difference b/w OPTS and FIFO is that;

- FIFO algorithm uses the time when the pages was built in and OPT uses the time when a page is to be used.
- The LRU algorithm replaces the pages that have not been used for longest period of time.
- The LRU associated its pages with the time of that pages last use.
 - This strategy is the optimal page replacement algorithm looking backward in time rather than forward.
- The first 5 faults are similar to optimal replacement.
- When reference to page 4 occurs, LRU sees that of the three frames, page 2 as used least recently. The most recently used page is page 0 and just before page 3 was used. The LRU policy is often used as a page replacement algorithm and considered to be good. The main

Ex: consider the following reference string



to how to implement LRU is the LRU requires additional h/w assistance.

Two implementations are possible:

Counters: In this we associate each page table entry a time-of-use field, and add to the cpu a logical clock or counter. The clock is incremented for each memory reference. When a reference to a page is made, the contents of the clock register are copied to the time-of-use field in the page table entry for that page.

In this way we have the time of last reference to each page we replace the page with smallest time value. The time must also be maintained when page tables are changed.

Stack: Another approach to implement LRU replacement is to keep a stack of page numbers when a page is referenced it is removed from the stack and put on to the top of stack. In this way the top of stack is always the most recently used page and the bottom in least recently used page. Since the entries are removed from the stack it is best implement by a doubly linked list. With a head and tail pointer.

Neither optimal replacement nor LRU replacement suffers from Belady's Anamoly. These are called stack algorithms.

LRU Approximation

- An LRU page replacement algorithm should update the page removal status information after every page reference updating is done by software, cost increases.
- But hardware LRU mechanism tend to degrade execution performance at the same time, then substantially increases the cost. For this reason, simple and efficient algorithm that approximation the LRU have been developed. With h/w support the reference bit was used. A reference bit associate with each memory block and this bit automatically set to 1 by the h/w whenever the page is referenced. The single reference bit per clock can be used to approximate LRU removal.
- The page removal s/w periodically resets the reference bit to 0, write the execution of the users job causes some reference bit to be set to 1.
- If the reference bit is 0 then the page has not been referenced since the last time the reference bit was set to 0.

Count Based Page Replacement

There is many other algorithms that can be used for page replacement, we can keep a counter of the number of references that has made to a page.

a) LFU (least frequently used) :

This causes the page with the smallest count to be replaced. The reason for this selection is that actively used page should have a large reference count.

This algorithm suffers from the situation in which a page is used heavily during the initial phase of a process but never used again. Since it was used heavily, it has a large count and remains in memory even though it is no longer needed.

b) Most Frequently Used(MFU) :

This is based on the principle that the page with the smallest count was probably just brought in and has yet to be used.

Allocation of Frames

- The allocation policy in a virtual memory controls the operating system decision regarding the amount of real memory to be allocated to each active process.
- In a paging system if more real pages are allocated, it reduces the page fault frequency and improved turnaround throughput.
- If too few pages are allocated to a process its page fault frequency and turnaround times may deteriorate to unacceptable levels.
- The minimum number of frames per process is defined by the architecture, and the maximum number of frames. This scheme is called equal allocation.
- With multiple processes competing for frames, we can classify page replacement into two broad categories

a) Local Replacement: requires that each process selects frames from only its own sets of allocated frame.

b). Global Replacement: allows a process to select frame from the set of all frames. Even if the frame is currently allocated to some other process, one process can take a frame from another.

In local replacement the number of frames allocated to a process do not change but with global replacement number of frames allocated to a process do not change global replacement results in greater system throughput.

Other consideration

There is much other consideration for the selection of a replacement algorithm and allocation policy.

1) Preparing: This is an attempt to present high level of initial paging. This strategy is to bring into memory all the pages at one time.

2) TLB Reach: The TLB reach refers to the amount of memory accessible from the TLB and is simply the no of entries multiplied by page size.

3) Page Size: following parameters are considered

a) page size us always power of 2 (from 512 to 16k)

b) Internal fragmentation is reduced by a small page size.

c) A large page size reduces the number of pages needed.

4) Invented Page table: This will reduce the amount of primary memory i.e. needed to track virtual to physical address translations.

5) Program Structure: Careful selection of data structure can increase the locality and hence lower the page fault rate and the number of pages in working state.

6) Real time Processing: Real time system almost never has virtual memory. Virtual memory is the antithesis of real time computing, because it can introduce unexpected long term delay in the execution of a process.

Thrashing

➔ If the number of frames allocated to a low-priority process falls below the minimum number required by the computer architecture then we suspend the process execution.

➔ A process is thrashing if it is spending more time in paging than executing.

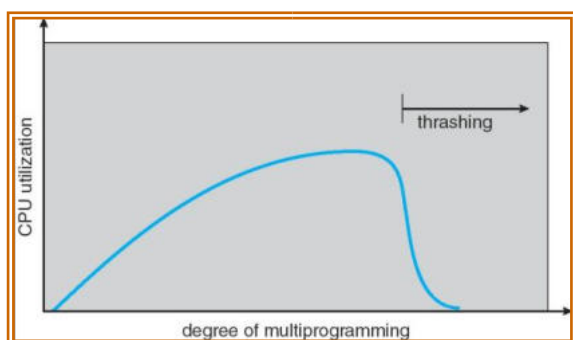
➔ If the processes do not have enough number of frames, it will quickly page fault. During this it must replace some page that is not currently in use. Consequently it quickly faults again and again. The process continues to fault, replacing pages for which it then faults and brings back. This high paging activity is called thrashing. The phenomenon of excessively moving pages back and forth b/w memory and secondary has been called thrashing.

Cause of Thrashing

- Thrashing results in severe performance problem.
- The operating system monitors the cpu utilization is low. We increase the degree of multi programming by introducing new process to the system.
- A global page replacement algorithm replaces pages with no regards to the process to which they belong.

The figure shows the thrashing

➔ As the degree of multi programming increases, more slowly until a maximum is reached. If the degree of multi programming is increased further thrashing sets in and the cpu utilization drops sharply.



➔ At this point, to increase CPU utilization and stop thrashing, we must increase the degree of multi programming. We can limit the effect of thrashing by using a local replacement algorithm. To prevent thrashing, we must provide a process as many frames as it needs.

Locality of Reference:

- As the process executes it moves from locality to locality.
- A locality is a set of pages that are actively used.
- A program may consist of several different localities, which may overlap.
- Locality is caused by loops in code that find to reference arrays and other data structures by indices.

The ordered list of page number accessed by a program is called reference string. Locality is of two types

- 1) spatial locality
- 2) temporal locality

Working set model

Working set model algorithm uses the current memory requirements to determine the number of page frames to allocate to the process, an informal definition is

†the collection of pages that a process is working with and which must be resident if the process to avoid thrashing†.

The idea is to use the recent needs of a process to predict its future reader.

The working set is an approximation of programs locality.

Ex: given a sequence of memory reference, if the working set window size to memory references, then working set at time t1 is {1,2,5,6,7} and at t2 is changed to {3,4}

• At any given time, all pages referenced by a process in its last 4 seconds of execution are considered to comprise its working set.

- A process will never execute until its working set is resident in main memory.
- Pages outside the working set can be discarded at any movement.

Working sets are not enough and we must also introduce balance set.

a) If the sum of the working sets of all the run able process is greater than the size of memory the refuse some process for a while.

b) Divide the run able process into two groups, active and inactive. The collection of active set is called the balance set. When a process is made active its working set is loaded.

c) Some algorithm must be provided for moving process into and out of the balance set.

As a working set is changed, corresponding change is made to the balance set. Working set prevents thrashing by keeping the degree of multi programming as high as possible. Thus it optimizes the CPU utilization. The main disadvantage of this is keeping track of the working set.

File System Interface

- A file is a collection of similar records.
- The data can't be written on to the secondary storage unless they are within a file.
- Files represent both the program and the data. Data can be numeric, alphanumeric, alphabetic or binary.
- Many different types of information can be stored on a file ---Source program, object programs, executable programs, numeric data, payroll recorder, graphic images, sound recordings and so on.
- A file has a certain defined structures according to its type:-
- Text file:- Text file is a sequence of characters organized in to lines.
- Object file:- Object file is a sequence of bytes organized in to blocks understandable by the systems linker.
- Executable file:- Executable file is a series of code section that the loader can bring in to memory and execute.
- Source File:- Source file is a sequence of subroutine and function, each of which are further organized as declaration followed by executable statements.

UNIT IV FILE SYSTEMS AND IO SYSTEMS

File Attributes:-

o File attributes varies from one OS to other. The common file attributes are:

- Name:- The symbolic file name is the only information kept in human readable form.
- Identifier:- The unique tag, usually a number, identifies the file within the file system. It is the non-readable name for a file.
- Type:- This information is needed for those systems that supports different types.
- Location:- This information is a pointer to a device and to the location of the file on that device.
- Size:- The current size of the file and possibly the maximum allowed size are included in this attribute.
- Protection:- Access control information determines who can do reading, writing, execute and so on.
- Time, data and User Identification:- This information must be kept for creation, last modification and last use. These data are useful for protection, security and usage monitoring.

File Operation:-

File is an abstract data type. To define a file we need to consider the operation that can be performed on the file. Basic operations of files are:-

1. Creating a file:- Two steps are necessary to create a file. First space in the file system for file is found. Second an entry for the new file must be made in the directory. The directory entry records the name of the file and the location in the file system.
2. Writing a file:- System call is mainly used for writing in to the file. System call specify the name of the file and the information i.e., to be written on to the file. Given the name the system search the entire directory for the file. The system must keep a write pointer to the location in the file where the next write to be taken place.
3. Reading a file:- To read a file system call is used. It requires the name of the file and the memory address. Again the directory is searched for the associated directory and system must maintain a read pointer to the location in the file where next read is to take place.
4. Delete a file:- System will search for the directory for which file to be deleted. If entry is found it releases all free space. That free space can be reused by another file.
- Truncating the file:- User may want to erase the contents of the file but keep its attributes. Rather than forcing the user to delete a file and then recreate it, truncation allows all attributes to remain unchanged except for file length.
- Repositioning within a file:- The directory is searched for appropriate entry and the current file position is set to a given value. Repositioning within a file does not need to involve actual i/o. The file operation is also known as file seeks.

In addition to this basis 6 operations the other two operations include appending new information to the end of the file and renaming the existing file. These primitives can be combined to perform other two operations.

Most of the file operation involves searching the entire directory for the entry associated with the file. To avoid this OS keeps a small table containing information about an open file (the open table). When a file operation is requested, the file is specified via index in to this table. So searching is not required.

Several piece of information are associated with an open file:-

- File pointer:- on systems that does not include offset an a part of the read and write system calls, the system must track the last read-write location as current file position pointer. This pointer is unique to each process operating on a file.
- File open count:- As the files are closed, the OS must reuse its open file table entries, or it could run out of space in the table. Because multiple processes may open a file, the system must wait for the last file to close before removing the open file table entry. The counter tracks the number of copies of open and closes and reaches zero to last close.

- **Disk location of the file:-** The information needed to locate the file on the disk is kept in memory to avoid having to read it from the disk for each operation.
- **Access rights:-** Each process opens a file in an access mode. This information is stored on per-process table the OS can allow OS deny subsequent i/o request.

Access Methods:-

The information in the file can be accessed in several ways.

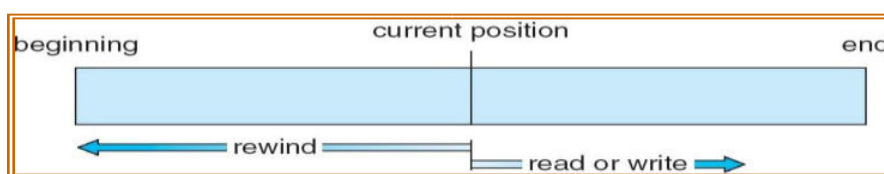
Different file access methods are:-

Sequential Access:-

Sequential access is the simplest access method. Information in the file is processed in order, one record after another. Editors and compilers access the files in this fashion. Normally read and write operations are done on the files. A read operation reads the next portion of the file and automatically advances a file pointer, which track next i/I track.

Write operation appends to the end of the file and such a file can be next to the beginning.

Sequential access depends on a tape model of a file.



Direct Access:-

Direct access allows random access to any file block. This method is based on disk model of a file. A file is made up of fixed length logical records. It allows the program to read and write records rapidly in any order. A direct access file allows arbitrary blocks to be read or written.

Eg:- User may need block 13, then read block 99 then write block 12.

For searching the records in large amount of information with immediate result, the direct access method is suitable. Not all OS support sequential and direct access. Few OS use sequential access and some OS uses direct access. It is easy to simulate sequential access on a direct access but the reverse is extremely inefficient.

Indexing Method:-

- The index is like an index at the end of a book which contains pointers to various blocks.
- To find a record in a file, we search the index and then use the pointer to access the file directly and to find the desired record.
- With large files index file itself can be very large to be kept in memory. One solution to create an index to the index files itself. The primary index file would contain pointer to secondary index files which would point to the actual data items.

Two types of indexes can be used:-

- a. Exhaustive index:- Contain one entry for each of record in the main file.

An index itself is organized as a sequential file.

- b. Partial index:- Contains entries to records where the field of interest exists with records of variable length, some record will not contain all fields. When a new record is added to the main file, all index files must be updated.

Directory Structure:-

The file systems can be very large. Some systems store millions of files on the disk. To manage all this data we need to organize them. This organization is done in two parts:-

1. Disks are split into one or more partitions also known as minidisks.
2. Each partition contains information about files within it. This information is kept in entries in a device directory or volume table of contents.

The device directory or simple directory records information as name, location, size, type for all files on the partition.

The directory can be viewed as a symbol table that translates the file names into the directory entries. The directory itself can be organized in many ways.

When considering a particular directory structure, we need to keep in mind the operations that are to be performed on a directory.

- Search for a file:- Directory structure is searched for finding particular file in the directory. Files have symbolic names and similar names may indicate a relationship between files, we may want to be able to find all the files whose names match a particular pattern.
- Create a file:- New files can be created and added to the directory.
- Delete a file:- When a file is no longer needed, we can remove it from the directory.
- List a directory:- We need to be able to list the files in the directory and the contents of the directory entry for each file in the list.
- Rename a file:- Name of the file must be changeable when the contents or use of the file is changed. Renaming allows the position within the directory structure to be changed.
- Traverse the file:- It is always good to keep the backup copy of the file so that it can be used when the system fails or when the file system is not in use.

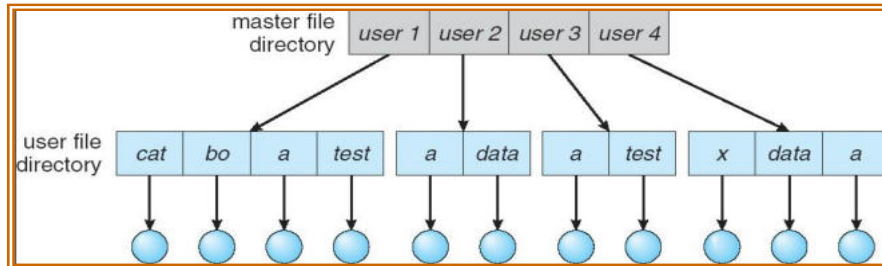
1. Single-level directory:-

This is the simplest directory structure. All the files are contained in the same directory which is easy to support and understand.

Disadvantage:-

- Not suitable for a large number of files and more than one user.
- Because of single directory files, files require unique file names.
- Difficult to remember names of all the files as the number of files increases.

MS-DOS OS allows only 11 character file name where as UNIX allows 255 character



2. Two-level directory:-

- A single level directory often leads to the confusion of file names between different users. The solution here is to create separate directory for each user.
- In two level directories each user has its own directory. It is called User File Directory (UFD). Each UFD has a similar structure, but lists only the files of a single user.
- When a user job starts or users logs in, the systems Master File Directory (MFD) is searched. The MFD is indexed by the user name or account number and each entry points to the UFD for that user.
- When a user refers to a particular file, only his own UFD is searched. Thus different users may have files with the same name.
- To create a file for a user, OS searches only those users UFD to ascertain whether another file of that name exists.

To delete a file checks in the local UFD so that accidentally delete another user's file with the same name. Although two-level directories solve the name collision problem but it still has some disadvantage.

This structure isolates one user from another. This isolation is an advantage. When the users are independent but disadvantage, when some users want to co-operate on some table and to access one another file.

3. Tree-structured directories:-

MS-DOS use Tree structure directory. It allows users to create their own subdirectory and to organize their files accordingly. A subdirectory contains a set of files or subdirectories. A directory is simply another file, but it is treated in a special way. The entire directory will have the same internal format. One bit in each entry defines the entry as a file (0) and as a subdirectory (1). Special system calls are used to create and delete directories.

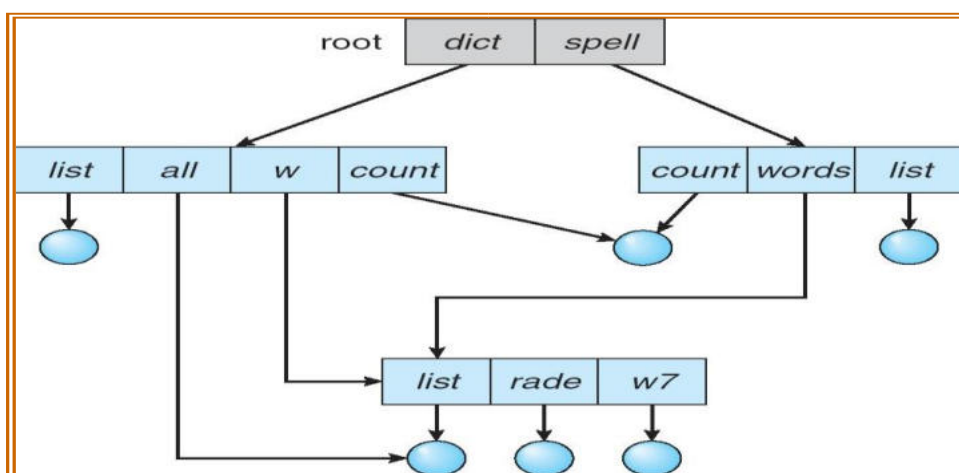
In normal use each user has a current directory. Current directory should contain most of the files that are of the current interest of users. When a reference to a file is needed the current directory is searched. If file is needed i.e., not in the current directory to be the directory currently holding that file.

Path name can be of two types:-

- a. **Absolute path name**:- Begins at the root and follows a path down to the specified file, giving the directory names on the path.
- b. **Relative path name**:- Defines a path from the current directory.
- c. One important policy in this structure is how to handle the deletion of a directory.
 - a. If a directory is empty, its entry can simply be deleted.
 - b. If a directory is not empty, one of the two approaches can be used.
 - In MS-DOS, the directory is not deleted until it becomes empty.
 - In UNIX, RM command is used with some options for deleting directory.

4. **Acyclic graph directories**:-

- It allows directories to have shared subdirectories and files.
- Same file or directory may be in two different directories.
- A graph with no cycles is a generalization of the tree structure subdirectories scheme.
- Shared files and subdirectories can be implemented by using links.
- A link is a pointer to another file or a subdirectory.
- A link is implemented as absolute or relative path.
- An acyclic graph directory structure is more flexible than is a simple tree structure but some times it is more complex.



File System Mounting:-

The file system must be mounted before it can be available to processes on the system. The procedure for mounting the file is:

- o The OS is given the name of the device and the location within the file structure at which to attach the file system (mount point). A mount point will be an empty directory at which the mounted file system will be attached.
- *Eg:-* On UNIX a file system containing users home directory might be mounted as /home then to access the directory structure within that file system. We must precede the directory names as /home/jane.
- o Then OS verifies that the device contains this valid file system. OS uses device drivers for this verification.
- o Finally the OS mounts the file system at the specified mount point.

File System Structure:-

Disks provide bulk of secondary storage on which the file system is maintained. Disks have two characteristics:-

- They can be rewritten in place i.e., it is possible to read a block from the disk to modify the block and to write back in to same place.
- They can access any given block of information on the disk. Thus it is simple to access any file either sequentially or randomly and switching from one file to another.

The lowest level is the i/o control consisting of device drivers and interrupt handlers to transfer the information between memory and the disk system. The device driver is like a translator. Its input is a high level command and the o/p consists of low level hardware specific instructions, which are used by the hardware controllers which interface I/O device to the rest of the system.

The basic file system needs only to issue generic commands to the appropriate device drivers to read and write physical blocks on the disk.

The file organization module knows about files and their logical blocks as well as physical blocks. By knowing the type of file allocation used and the location of the file, the file organization module can translate logical block address to the physical block address. Each logical block is numbered 0 to N. Since the physical blocks containing the data usually do not match the logical numbers, So a translation is needed to locate each block. The file allocation modules also include free space manager which tracks the unallocated blocks and provides these blocks when requested.

The logical file system uses the directory structure to provide the file organization module with the information, given a symbolic file name. The logical file system also responsible for protection and security.

Logical file system manages metadata information. Metadata includes all the file system structures excluding the actual data.

The file structure is maintained via file control block (FCB). FCB contains information about the file including the ownership permission and location of the file contents.

File System Implementation:-

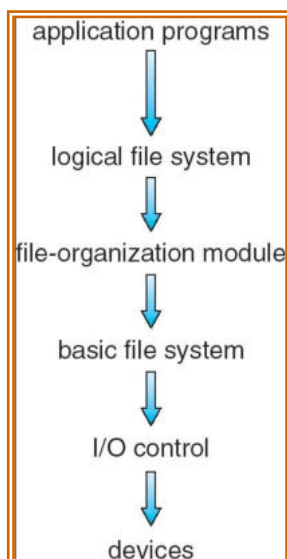
- File system is implemented on the disk and the memory.
- The implementation of the file system varies according to the OS and the file system, but there are some general principles.

If the file system is implemented on the disk it contains the following information:-

- Boot Control Block:-** can contain information needed by the system to boot an OS from that partition. If the disk has no OS, this block is empty. It is the first block of the partition. In UFS boot block, In NTFS partition boot sector.
- Partition control Block:-** contains partition details such as the number of blocks in partition, size of the blocks, number of free blocks, free block pointer, free FCB count and FCB pointers. In NTFS master file tables, In UFS super block.
- Directory structure is used to organize the files.
- An FCB contains many of the files details, including file permissions, ownership, size, location of the data blocks. In UFS inode, In NTFS this information is actually stored within master file table.
- Structure of the file system management in memory is as follows:-
 - An in-memory partition table containing information about each mounted information.
 - An in-memory directory structure that holds the directory information of recently accessed directories.
 - The system wide open file table contains a copy of the FCB of each open file as well as other information.
 - The per-process open file table contains a pointer to the appropriate entry in the system wide open file table as well as other information.

File System Organization:-

To
the
A



provide efficient and convenient access to the disks, the OS provides file system to allow the data to be stored, located and retrieved.

file system has two design problems:-

- How the file system should look to the user.
- Selecting algorithms and data structures that must be created to map logical file system on to the physical secondary storage devices.

The file system itself is composed of different levels. Each level uses the feature of the lower levels to create new features for use by higher levels. The following structures shows an example of layered design

The lowest level is the i/o control consisting of device drivers and interrupt handlers to transfer the information between memory and the disk system. The device driver is like a translator. Its input is a high level command and the o/p consists of low level hardware specific instructions, which are used by the hardware controllers which interface I/O device to the rest of the system. The basic file system needs only to issue generic commands to the appropriate device drivers to read and write physical blocks on the disk.

The file organization module knows about files and their logical blocks as well as physical blocks. By knowing the type of file allocation used and the location of the file, the file organization module can translate logical block address to the physical block address. Each logical block is numbered 0 to N. Since the physical blocks containing the data usually do not match the logical numbers, So a translation is needed to locate each block. The file allocation modules also include free space manager which tracks the unallocated blocks and provides these blocks when requested.

The logical file system uses the directory structure to provide the file organization module with the information, given a symbolic file name. The logical file system also responsible for protection and security.

Logical file system manages metadata information. Metadata includes all the file system structures excluding the actual data.

The file structure is maintained via file control block (FCB). FCB contains information about the file including the ownership permission and location of the file contents.

File System Implementation:-

- File system is implemented on the disk and the memory.
- The implementation of the file system varies according to the OS and the file system, but there are some general principles.

If the file system is implemented on the disk it contains the following information:-

- **Boot Control Block:-** can contain information needed by the system to boot an OS from that partition. If the disk has no OS, this block is empty. It is the first block of the partition. In UFS it's boot block, In NTFS it's partition boot sector.
- **Partition control Block:-** contains partition details such as the number of blocks in partition, size of the blocks, number of free blocks, free block pointer, free FCB count and FCB pointers. In NTFS it's master file tables, In UFS it's super block.
- Directory structure is used to organize the files.
- An FCB contains many of the files details, including file permissions, ownership, size, location of the data blocks. In UFS it's inode, In NTFS this information is actually stored within master file table.

Structure of the file system management in memory is as follows:-

- An in-memory partition table containing information about each mounted information.
- An in-memory directory structure that holds the directory information of recently accessed directories.
- The system wide open file table contains a copy of the FCB of each open file as well as other information.
- The per-process open file table contains a pointer to the appropriate entry in the system wide open file table as well as other information.

A typical file control blocks is shown below

File permission
File dates (create, access, write)
File owner, group, Acc
File size
File data blocks

Partition and Mounting:-

- A disk can be divided into multiple partitions. Each partition can be either raw i.e., containing no file system and cooked i.e., containing a file system.

- Raw disk is used where no file system is appropriate. UNIX swap space can use a raw partition and do not use file system.
- Some db uses raw disk and format the data to suit their needs. Raw disks can hold information need by disk RAID (Redundant Array of Independent Disks) system.
- Boot information can be stored in a separate partition. Boot information will have their own format. At the booting time system does not load any device driver for the file system. Boot information is a sequential series of blocks, loaded as an image in to memory.
- Dual booting is also possible on some PCs, more than one OS are loaded on a system. A boot loader understands multiple file system and multiple OS can occupy the boot space once loaded it can boot one of the OS available on the disk. The disks can have multiple portions each containing different types of file system and different types of OS. Root partition contains the OS kernel and is mounted at a boot time. Microsoft window based systems mount each partition in a separate name space denoted by a letter and a colon. On UNIS file system can be mounted at any directory.

Directory Implementation:-

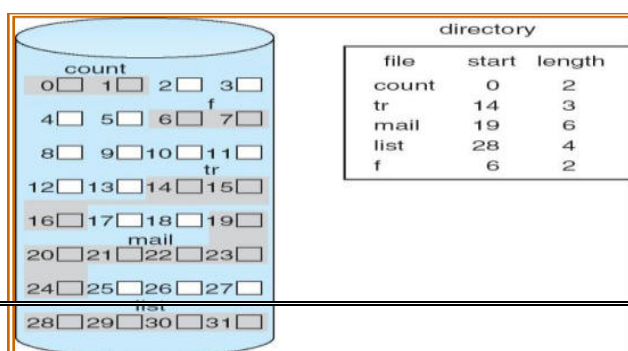
Directory is implemented in two ways:-

1. Linear list:-

- Linear list is a simplest method.
- It uses a linear list of file names with pointers to the data blocks.
- Linear list uses a linear search to find a particular entry.
- Simple for programming but time consuming to execute.
- For creating a new file, it searches the directory for the name whether same name already exists.
- Linear search is the main disadvantage.
- Directory implementation is used frequently and uses would notice a slow implementation of access to it.

2. Hash table:-

- Hash table decreases the directory search time.
- Insertion and deletion are fairly straight forward.
- Hash table takes the value computed from that file name.
- Then it returns a pointer to the file name in the linear list.
- Hash table uses fixed size.



Allocation Methods:-

The space allocation strategy is closely related to the efficiency of the file accessing and of logical to physical mapping of disk addresses.

A good space allocation strategy must take in to consideration several factors such as:- 1. Processing speed of sequential access to files, random access to files and allocation and de-allocation of blocks.

2. Disk space utilization.

3. Ability to make multi-user and multi-track transfers.

4. Main memory requirement of a given algorithm.

Three major methods of allocating disk space is used.

1. Contiguous Allocation:-

A single set of blocks is allocated to a file at the time of file creation. This is a preallocation strategy that uses portion of variable size. The file allocation table needs just a single entry for each file, showing the starting block and the length of the file.

The figure shows the contiguous allocation method.

If the file is n blocks long and starts at location b , then it occupies blocks $b, b+1, b+2, \dots, b+n-1$. The file allocation table entry for each file indicates the address of starting block and the length of the area allocated for this file.

Contiguous allocation is the best from the point of view of individual sequential file. It is easy to retrieve a single block. Multiple blocks can be brought in one at a time to improve I/O performance for sequential processing. Sequential and direct access can be supported by contiguous allocation.

Contiguous allocation algorithm suffers from external fragmentation. Depending on the amount of disk storage the external fragmentation can be a major or minor problem.

Compaction is used to solve the problem of external fragmentation.

The following figure shows the contiguous allocation of space after compaction. The original disk was then freed completely creating one large contiguous space.

If the file is n blocks long and starts at location b , then it occupies blocks $b, b+1, b+2, \dots, b+n-1$. The file allocation table entry for each file indicates the address of starting block and the length of the area allocated for this file. Contiguous allocation is the best from the point of view of individual sequential file. It is easy to retrieve a single block. Multiple blocks can be brought in one at a time to improve I/O performance for sequential processing. Sequential and direct access can be supported by contiguous allocation. Contiguous allocation algorithm suffers

from external fragmentation. Depending on the amount of disk storage the external fragmentation can be a major or minor problem. Compaction is used to solve the problem of external fragmentation. The following figure shows the contiguous allocation of space after compaction. The original disk was then freed completely creating one large contiguous space. Another problem with contiguous allocation algorithm is pre-allocation, i.e., it is necessary to declare the size of the file at the time of creation.

Characteristics:-

- Supports variable size portion.
- Pre-allocation is required.
- Requires only single entry for a file.
- Allocation frequency is only once.

Advantages:-

- Supports variable size problem.
- Easy to retrieve single block.
- Accessing a file is easy.
- Provides good performance

Disadvantage:-

- Pre-allocation is required.
- It suffers from external fragmentation.

2. **Linked Allocation:-**

- It solves the problem of contiguous allocation. This allocation is on the basis of an individual block. Each block contains a pointer to the next block in the chain.
- The disk block can be scattered any where on the disk.
- The directory contains a pointer to the first and the last blocks of the file.
- The following figure shows the linked allocation. To create a new file, simply create a new entry in the directory.

- There is no external fragmentation since only one block is needed at a time.

• The size of a file need not be declared when it is created. A file can continue to grow as long as free blocks are available.

Advantages:-

- No external fragmentation.
- Compaction is never required.
- Pre-allocation is not required.

Disadvantage:-

- Files are accessed sequentially.
- Space required for pointers.
- Reliability is not good.
- Cannot support direct access.

3. Indexed Allocation:-

The file allocation table contains a separate one level index for each file. The index has one entry for each portion allocated to the file. The i th entry in the index block points to the i th block of the file.

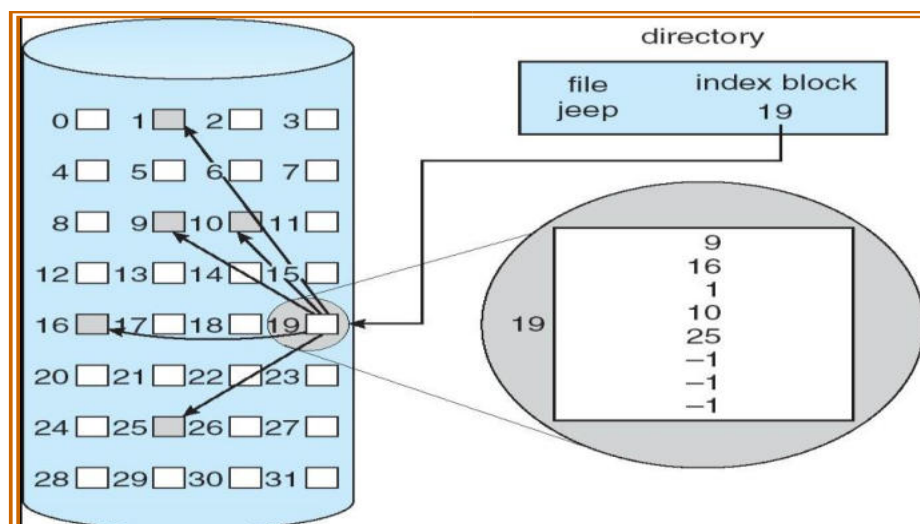
The indexes are not stored as a part of file allocation table rather than the index is kept as a separate block and the entry in the file allocation table points to that block.

Allocation can be made on either fixed size blocks or variable size blocks. When the file is created all pointers in the index block are set to nil. When an entry is made a block is obtained from free space manager.

Allocation by fixed size blocks eliminates external fragmentation where as allocation by variable size blocks improves locality.

Indexed allocation supports both direct access and sequential access to the file.

The following figure shows indexed allocation.



Advantages:-

- Supports both sequential and direct access.
- No external fragmentation.
- Faster than other two methods.
- Supports fixed size and variable sized blocks.

Disadvantage:-

- Suffers from wasted space.
- Pointer overhead is generally greater.

Mass Storage Structure Disk Structure:-

- Disks provide a bulk of secondary storage. Disks come in various sizes, speed and information can be stored optically or magnetically.
- Magnetic tapes were used early as secondary storage but the access time is less than disk.
- Modern disks are organized as single one-dimensional array of logical blocks.
- The actual details of disk i/o open depends on the computer system, OS, nature of i/o channels and disk controller hardware.
- The basic unit of information storage is a sector. The sectors are stored on flat, circular, media disk. This disk media spins against one or more read-write heads. The head can move from the inner portion of the disk to the outer portion.
- When the disk drive is operating the disks is rotating at a constant speed.
- To read or write the head must be positioned at the desired track and at the beginning of the desired sector on that track.
- Track selection involves moving the head in a movable head system or electronically selecting one head on a fixed head system.
- These characteristics are common to floppy disks, hard disks, CD-ROM and DVD.

Disk Performance Parameters:-

1. **Seek Time:-** Seek time is the time required to move the disk arm to the required track.
Seek time can be given by $T_s = m * n + s$. Where T_s = seek time
 n = number of track traversed.
 m = constant that depends on the disk drive s = startup time.
2. **Rotational Latency:-** Rotational latency is the additional addition time for waiting for the disk to rotate to the desired sector to the disk head.
3. **Rotational Delay:-** Disks other than the floppy disk rotate at 3600 rpm which is one revolution per 16.7ms.
4. **Disk Bandwidth:-** Disk bandwidth is the total number of bytes transferred divided by total time between the first request for service and the completion of last transfer. Transfer time = $T = b / rN$

Where b = number of bytes transferred.

T = transfer time.

r = rotational speed in RpS.

N = number of bytes on the track.

Average

access time = $T_a = T_s + 1/2r + b/rN$

Where T_s = seek time.

5. Total capacity of the disk:- It is calculated by using following formula.

Number of cylinders * number of heads * number of sector/track * number of bytes/sector.

Disk Scheduling:-

The amount of head movement needed to satisfy a series of i/o request can affect the performance. If the desired drive and the controller are available the request can be serviced immediately. If the device or controller is busy any new requests for service will be placed on the queue of pending requests for that drive when one request is complete the OS chooses which pending request to service next. Different types of scheduling algorithms are as follows:-

1. FCFS scheduling algorithm:-

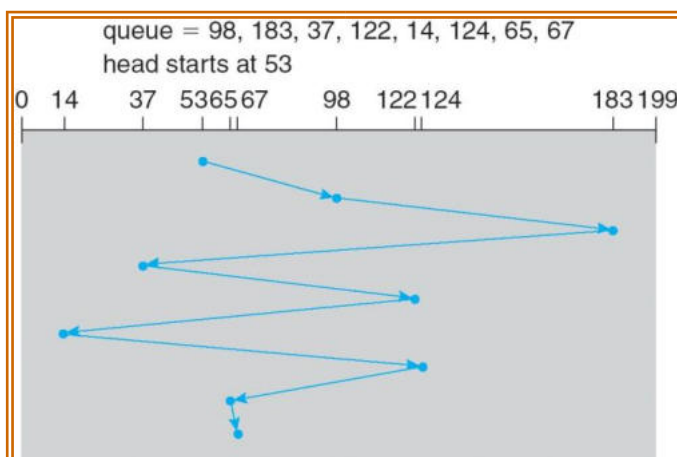
This is the simplest form of disk scheduling algorithm. This services the request in the order they are received. This algorithm is fair but do not provide fastest service.

It takes no special time to minimize the overall seek time.

Eg:- consider a disk queue with request for i/o to blocks on cylinders. 98, 183, 37, 122, 14, 124, 65, 67

If the disk head is initially at 53, it will first move from 53 to 98 then to 183 and then to 37, 122, 14, 124, 65, 67 for a total head movement of 640 cylinders.

The wild swing from 122 to 14 and then back to 124 illustrates the problem with this schedule.



If the requests for cylinders 37 and 14 could be serviced together before or after 122 and 124 the total head movement could be decreased substantially and performance could be improved.

2. **SSTF** (Shortest seek time first) **algorithm**:-

This selects the request with minimum seek time from the current head position. Since seek time increases with the number of cylinders traversed by head, SSTF chooses the pending request closest to the current head position.

Eg:- :- consider a disk queue with request for i/o to blocks on cylinders. 98, 183, 37, 122, 14, 124, 65, 67

If the disk head is initially at 53, the closest is at cylinder 65, then 67, then 37 is closer then 98 to 67. So it services 37, continuing we service 14, 98, 122, 124 and finally 183.

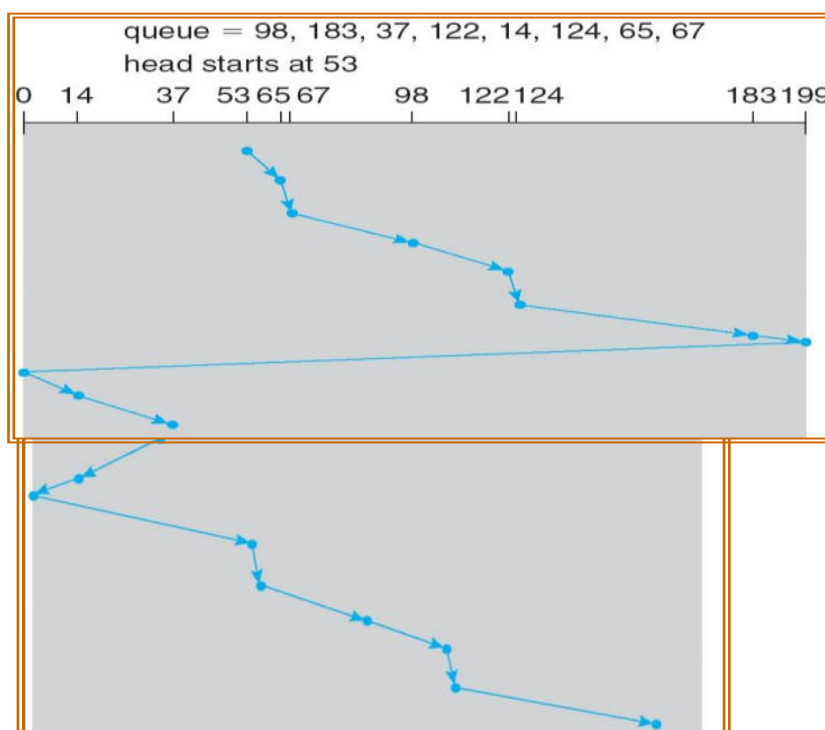
The total head movement is only 236 cylinders. SSTF is essentially a form of SJF and it may cause starvation of some requests. SSTF is a substantial improvement over FCFS, it is not optimal.

3. **SCAN** **algorithm**:-

In this the disk arm starts at one end of the disk and moves towards the other end, servicing the request as it reaches each cylinder until it gets to the other end of the disk. At the other end, the direction of the head movement is reversed and servicing continues.

Eg:- :- consider a disk queue with request for i/o to blocks on cylinders. 98, 183, 37, 122, 14, 124, 65, 67

If the disk head is initially at 53 and if the head is moving towards 0, it services 37 and then 14. At cylinder 0 the arm will reverse and will move towards the other end of the disk servicing 65, 67, 98, 122, 124 and 183.



If a request arrives just in from of head, it will be serviced immediately and the request just behind the head will have to wait until the arms reach other end and reverses direction. The SCAN is also called as elevator algorithm.

4. **C-SCAN** (Circular scan) **algorithm**:-

C-SCAN is a variant of SCAN designed to provide a more uniform wait time. Like SCAN, CSCAN moves the head from end of the disk to the other servicing the request along the way. When the head reaches the other end, it immediately returns to the beginning of the disk, without servicing any request on the return.

The C-SCAN treats the cylinders as circular list that wraps around from the final cylinder to the first one.

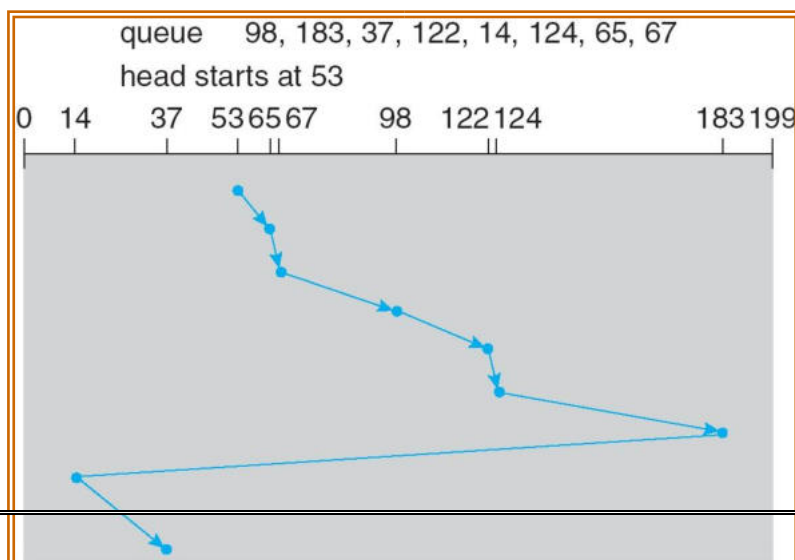
Eg:-

4.Look Scheduling algorithm:-

Both SCAN and C-SCAN move the disk arm across the full width of the disk. In practice neither of the algorithms is implemented in this way.

The arm goes only as far as the final request in each direction. Then it reverses, without going all the way to the end of the disk. These versions of SCAN and C-SCAN are called Look and C-Look

Eg:-



scheduling because they look for a request before continuing to move in a given direction.

Selection of Disk Scheduling Algorithm:-

SSTF is common and it increases performance over FCFS.

SCAN and C-SCAN algorithm is better for a heavy load on disk.

SCAN and C-SCAN have less starvation problem.

SSTF or Look is a reasonable choice for a default algorithm.

UNIT V CASE STUDY

Basic Concepts

- Linux looks and feels much like any other UNIX system; indeed, UNIX compatibility has been a major design goal of the Linux project. However, Linux is much younger than most UNIX systems. Its development began in 1991, when a Finnish university student, Linus Torvalds, began developing a small but self-contained kernel for the 80386 processor, the first true 32-bit processor in Intel's range of PC-compatible CPUs. of arbitrary files (but only read-only memory mapping was implemented in 1.0).
- A range of extra hardware support was included in this release. Although still restricted to the Intel PC platform, hardware support had grown to include floppy-disk and CD-ROM devices, as

well as sound cards, a range of mice, and international keyboards. Floating-point emulation was provided in the kernel for 80386 users who had no 80387 math coprocessor. System V UNIX-style interprocess communication (IPC), including shared memory, semaphores, and message queues, was implemented.

- At this point, development started on the 1.1 kernel stream, but numerous bug-fix patches were released subsequently for 1.0. A pattern was adopted as the standard numbering convention for Linux kernels. Kernels with an odd minor-version number, such as 1.1 or 2.5, are development kernels; even numbered minor-version numbers are stable production kernels. Updates for the stable kernels are intended only as remedial versions, whereas the development kernels may include newer and relatively untested functionality.

- As we will see, this pattern remained in effect until version 3.0 was given a major version-number increment because of two major new capabilities: support for multiple architectures, including a 64-bit native Alpha port, and symmetric multiprocessing (SMP) support. Additionally, the memory management code was substantially improved to provide a unified cache for file-system data independent of the caching of block devices.

- As a result of this change, the kernel offered greatly increased file-system and virtual memory performance. For the first time, file-system caching was extended to networked file systems, and writable memory-mapped regions were also supported. Other major improvements included the addition of internal kernel threads, a mechanism exposing dependencies between loadable modules, support for the automatic loading of modules on demand, file-system quotas, and POSIX-compatible real-time process-scheduling classes.

- Improvements continued with the release of Linux 2.2 in 1999. A port to Ultra SPARC systems was added. Networking was enhanced with more flexible firewalling, improved routing and traffic management, and support for TCP large window and selective acknowledgement. Acorn, Apple, and NT disks could now be read, and NFS was enhanced with a new kernel-mode NFS daemon. Signal handling, interrupts, and some I/O were locked at a finer level than before to improve symmetric multiprocessor (SMP) performance.

The Linux System

- Linux kernel is composed entirely of code written from scratch specifically for the Linux project, much of the supporting software that makes up the Linux system is not exclusive to Linux but is common to a number of UNIX-like operating systems. In particular, Linux uses many tools developed as part of Berkeley's BSD operating system, MIT's X Window System, and the Free Software Foundation's GNU project.

- This sharing of tools has worked in both directions. The main system libraries of Linux were originated by the GNU project, but the Linux community greatly improved the libraries by addressing omissions, inefficiencies, and bugs. Other components, such as the GNU C compiler (gcc), were already of sufficiently high quality to be used directly in Linux. The network administration tools under Linux were derived from code first developed for 4.3 BSD, but more recent BSD derivatives, such as

FreeBSD, have borrowed code from Linux in return. Examples of this sharing include the Intel floating-point-emulation math library and the PC sound-hardware device drivers.

- The Linux system as a whole is maintained by a loose network of developers collaborating over the Internet, with small groups or individuals having responsibility for maintaining the integrity of specific components.

- A small number of public Internet file-transfer-protocol (FTP) archive sites act as de facto standard repositories for these components. The File System Hierarchy Standard document is also maintained by the Linux community as a means of ensuring compatibility across the various system components.

- This standard specifies the overall layout of a standard Linux file system; it determines under which directory names configuration files, libraries, system binaries, and run-time data files should be stored.

Linux Distributions

- In theory, anybody can install a Linux system by fetching the latest revisions of the necessary system components from the FTP sites and compiling them. In Linux's early days, this is precisely what a Linux user had to do. As Linux has matured, however, various individuals and groups have attempted to make this job less painful by providing standard, precompiled sets of packages for easy installation.

- These collections, or distributions, include much more than just the basic Linux system. They typically include extra system-installation and management utilities, as well as precompiled and ready-to-install packages of many of the common UNIX tools, such as news servers, web browsers, text-processing and editing tools, and even games.

- The first distributions managed these packages by simply providing a means of unpacking all the files into the appropriate places. One of the important contributions of modern distributions, however, is advanced package management. Today's Linux distributions include a package-tracking database that allows packages to be installed, upgraded, or removed painlessly.

Linux Licensing

- The Linux kernel is distributed under version 2.0 of the GNU General Public License (GPL), the terms of which are set out by the Free Software Foundation. Linux is not public-domain software. Public domain implies that the authors have waived copyright rights in the software, but copyright rights in Linux code are still held by the code's various authors. Linux is free software, however, in the sense that people can copy it, modify it, use it in any manner they want, and give away (or sell) their own copies.

- The main implication of Linux's licensing terms is that nobody using Linux, or creating a derivative of Linux (a legitimate exercise), can distribute the derivative without including the source code. Software released under the GPL cannot be redistributed as a binary-only product.

- If you release software that includes any components covered by the GPL, then, under the GPL, you must make source code available alongside any binary distributions. (This restriction does not prohibit making—or even selling—binary software distributions, as long as anybody who

receives binaries is also given the opportunity to get the originating source code for a reasonable distribution charge.)

SYSTEM ADMINISTRATION

- In its overall design, Linux resembles other traditional, nonmicrokernel UNIX implementations. It is a multiuser, preemptively multitasking system with a full set of UNIX-compatible tools. Linux's file system adheres to traditional UNIX semantics, and the standard UNIX networking model is fully implemented. The internal details of Linux's design have been influenced heavily by the history of this operating system's development.

- Although Linux runs on a wide variety of platforms, it was originally developed exclusively on PC architecture. A great deal of that early development was carried out by individual enthusiasts rather than by wellfunded development or research facilities, so from the start Linux attempted to squeeze as much functionality as possible from limited resources. Today, Linux can run happily on a multiprocessor machine with many gigabytes of main memory and many terabytes of disk space, but it is still capable of operating usefully in under 16 MB of RAM.

1. Components of a Linux System

The Linux system is composed of three main bodies of code, in line with most traditional UNIX implementations:

Kernel. The kernel is responsible for maintaining all the important abstractions of the operating system, including such things as virtual memory and processes.

System libraries. The system libraries define a standard set of functions through which applications can interact with the kernel. These functions implement much of the operating-system functionality that does not need the full privileges of kernel code. The most important system library is the C library, known as `libc`. In addition to providing the standard C library, `libc` implements the user mode side of the Linux system call interface, as well as other critical system-level interfaces.

System utilities. The system utilities are programs that perform individual, specialized management tasks. Some system utilities are invoked just once to initialize and configure some aspect of the system. Others —known as daemons in UNIX terminology—run permanently, handling such tasks as responding to incoming network connections, accepting logon requests from terminals, and updating log files.

2. Kernel Modules

The Linux kernel has the ability to load and unload arbitrary sections of kernel code on demand. These loadable kernel modules run in privileged kernel mode and as a consequence have full access to all the hardware capabilities of the machine on which they run. In theory, there is no restriction on what a kernel module is allowed to do. Among other things, a kernel module can implement a device driver, a file system, or a networking protocol.

Kernel modules are convenient for several reasons. Linux's source code is free, so anybody wanting to write kernel code is able to compile a modified kernel and to reboot into that new functionality. However, recompiling, relinking, and reloading the entire kernel is a cumbersome cycle to undertake when you are developing a new driver. If you use kernel modules, you do not have to make a new

kernel to test a new driver—the driver can be compiled on its own and loaded into the already running kernel. Of course, once a new driver is written, it can be distributed as a module so that other users can benefit from it without having to rebuild their kernels.

The module support under Linux has four components:

1. The module-management system allows modules to be loaded into memory and to communicate with the rest of the kernel.
2. The module loader and unloader, which are user-mode utilities, work with the module-management system to load a module into memory.
3. The driver-registration system allows modules to tell the rest of the kernel that a new driver has become available.
4. A conflict-resolution mechanism allows different device drivers to reserve hardware resources and to protect those resources from accidental use by another driver.

1. Module Management

Loading a module requires more than just loading its binary contents into kernel memory. The system must also make sure that any references the correct locations in the kernel's address space. Linux deals with this reference updating by splitting the job of module loading into two separate sections: the management of sections of module code in kernel memory and the handling of symbols that modules are allowed to reference.

Linux maintains an internal symbol table in the kernel. This symbol table does not contain the full set of symbols defined in the kernel during the latter's compilation; rather, a symbol must be explicitly exported. The set of exported symbols constitutes a well-defined interface by which a module can interact with the kernel.

2. Driver Registration

Once a module is loaded, it remains no more than an isolated region of memory until it lets the rest of the kernel know what new functionality it provides. The kernel maintains dynamic tables of all known drivers and provides a set of routines to allow drivers to be added to or removed from these tables at any time. The kernel makes sure that it calls a module's startup routine when that module is loaded and calls the module's cleanup routine before that module is unloaded. These routines are responsible for registering the module's functionality.

A module may register many types of functionality; it is not limited to only one type. For example, a device driver might want to register two separate mechanisms for accessing the device. Registration tables include, among others, the following items:

- Device drivers. These drivers include character devices (such as printers, terminals, and mice), block devices (including all disk drives), and network interface devices.

File systems. The file system may be anything that implements

Linux's virtual file system calling routines. It might implement a format for storing files on a disk, but it might equally well be a network file system, such as NFS, or a virtual file system whose contents are generated on demand, such as Linux's /proc file system.

Network protocols. A module may implement an entire networking protocol, such as TCP or simply a new set of packet-filtering rules for a network firewall.

Binary format. This format specifies a way of recognizing, loading, and executing a new type of executable file.

(3) Conflict Resolution

Commercial UNIX implementations are usually sold to run on a vendor's own hardware. One advantage of a single-supplier solution is that the software vendor has a good idea about what hardware configurations are possible. PC hardware, however, comes in a vast number of configurations, with large numbers of possible drivers for devices such as network cards and video display adapters. The problem of managing the hardware configuration becomes more severe when modular device drivers are supported, since the currently active set of devices becomes dynamically variable.

Linux provides a central conflict-resolution mechanism to help arbitrate access to certain hardware resources. Its aims are as follows:

To prevent modules from clashing over access to hardware resources

To prevent autoprobes—device-driver probes that auto-detect device configuration— from interfering with existing device drivers

To resolve conflicts among multiple drivers trying to access the same hardware—as, for example, when both the parallel printer driver and the parallel line IP (PLIP) network driver try to talk to the parallel port.

REQUIREMENTS FOR LINUX SYSTEM ADMINISTRATOR

Hardware-Abstraction Layer

The HAL is the layer of software that hides hardware chipset differences from upper levels of the operating system. The HAL exports a virtual hardware drivers. Only a single version of each device driver is required for each CPU architecture, no matter what support chips might be present. Device drivers map devices and access them directly, but the chipset-specific details of mapping memory, configuring I/O buses, setting up DMA, and coping with motherboard-specific facilities are all provided by the HAL interfaces.

Kernel

The kernel layer of Windows has four main responsibilities: thread scheduling, lowlevel processor synchronization, interrupt and exception handling, and switching between user mode and kernel mode. The kernel is implemented in the C language, using assembly language only where absolutely necessary to interface with the lowest level of the hardware architecture.

Kernel Dispatcher

The kernel dispatcher provides the foundation for the executive and the subsystems. Most of the dispatcher is never paged out of memory, and its execution is never preempted. Its main responsibilities are thread scheduling and context switching, implementation of synchronization primitives, timer management, software interrupts (asynchronous and deferred procedure calls), and exception dispatching.

Threads and Scheduling

Like many other modern operating systems, Windows uses processes and threads for executable code. Each process has one or more threads, and each thread has its own scheduling state, including actual priority, processor affinity, and CPU usage information.

There are six possible thread states: ready, standby, running, waiting, transition, and terminated. Ready indicates that the thread is waiting to run. The highestpriority ready thread is moved to the standby state, which means it is the next thread to run. In a multiprocessor system, each processor keeps one thread in a standby state. A thread is running when it is executing on a processor. It runs until it is preempted by a higher-priority thread, until it terminates, until its allotted execution time (quantum) ends, or until it waits on a dispatcher object, such as an event signaling I/O completion. A thread is in the waiting state when it is waiting for a dispatcher object to be signaled. A thread is in the transition state while it waits for resources necessary for execution; for example, it may be waiting for its kernel stack to be swapped in from disk. A thread enters the terminated state when it finishes execution.

Implementation of Synchronization Primitives

Key operating-system data structures are managed as objects using common facilities for allocation, reference counting, and security. Dispatcher objects control dispatching and synchronization in the system. Examples of these objects include the following:

The event object is used to record an event occurrence and to synchronize this occurrence with some action. Notification events signal all waiting threads, and synchronization events signal a single waiting thread.

The mutant provides kernel-mode or user-mode mutual exclusion associated with the notion of ownership.

The mutex, available only in kernel mode, provides deadlock-free mutual exclusion.

The semaphore object acts as a counter or gate to control the number of threads that access a resource.

The thread object is the entity that is scheduled by the kernel dispatcher. It is associated with a process object, which encapsulates a virtual address space. The thread object is signaled when the thread exits, and the process object, when the process exits.

The timer object is used to keep track of time and to signal timeouts when operations take too long and need to be interrupted or when a periodic activity needs to be scheduled.

SETTING UP A LINUX MULTIFUNCTION SERVER

Follow the steps below to avoid any complications during the hardware installation:

1. Confirm that the printer you will use to connect to the DPR-1020 is operating correctly.
2. When you have confirmed that the printer is operating correctly, switch its power OFF.
3. Confirm that your network is operating normally.
4. Using a CAT 5 Ethernet cable, connect the DPR-1020 Ethernet Port (labelled LAN) to the network.
5. While the printer is turned OFF, connect the USB printer cable to the printer and then to the USB port on the Print Server.
6. Switch on the printer.
7. Insert the power adapter's output plug into the DC 5V power socket on the rear panel of the Print Server.
8. Connect the other end of the power adapter into a power outlet. This will supply power to the Print Server. The blue LED on the Print Server's front panel should turn on and the Print Server's self-test will proceed.

Power ON Self-Test

When the DPR-1020 is powered ON, it automatically performs a Self-Test on each of its major components. The final result of the Self-Test is signaled by the state of the USB LED indicator following the Self-Test. Preliminary to the actual component tests, the three LED indicators are tested to confirm their operation.

Immediately after power-up, all three of the blue LEDs should illuminate steadily for several seconds. Then the USB LED should light OFF simultaneously. Irregularity of any of the three LEDs during these LED tests may mean there is a problem with the LEDs themselves.

The actual component tests immediately follow the LED tests. A normal (no fault) result is signaled by simultaneous flashing of the LEDs three times, followed by a quiescent state with all three LEDs dark.

If the Self-Test routine traps any component error, then following the LED tests the Self-Test will halt and the LEDs will continuously signal the error according to the following table. In the event of any such error signal, contact your dealer for correction of the faulty unit.

Getting Started

Below is a sample network using the DPR-1020. The DPR-1020 has a built-in web configurator that allows users to easily configure the Print Server and manage multiple print queues through TCP/IP.



Auto-Run Installation

Insert the included installation CD into your computer's CD-ROM drive to initiate the auto-run program. If auto-run does not start, click My Computer > [CD ROM Drive Letter].

The content of the installation CD-ROM includes:

- Install PS Software – click this to install the PS Software, which contains PS-Link and PS-Wizard that can configure more settings for the MFP Server, such as:
 - o Change the IP address
 - p Support the multi-functions (Print/Scan/Copy/Fax) of a MFP printer, GDI printing, and other software from any MFP/GDI printer.
 - `- Easily add a printer to your computer.

View Quick Installation Guide – click this to preview the Quick Installation Guide in PDF format for step-by-step instructions of the MFP Server Installation.

View Manual – click this to preview the User Manual in PDF format for detailed information regarding the MFP Server.

Install Acrobat Reader – click this to install Acrobat Reader for the viewing and printing of PDF files found in this Installation CD-ROM.

Exit – click to close the Auto-Run program.

DOMAIN NAME SYSTEM

The domain name, or network name, is a unique name followed by a standard Internet suffixes such as .com, .org, .mil, .net, etc. You can pretty much name your LAN anything if it has a simple dial-up connection and your LAN is not a server providing some type of service to other hosts directly. In addition, our sample network is considered private since it uses IP addresses in the range of 192.168.1.x. Most importantly, the domain name of choice should not be accessible from the Internet if the above constraints are strictly enforced. Lastly, to obtain an "official" domain name you could register through InterNIC, Network Solutions or Register.com. See the Resources section later in this article for the Web sites with detailed instructions for obtaining official domain names.

Hostnames

- Another important step in setting up a LAN is assigning a unique hostname to each computer in the LAN. A hostname is simply a unique name that can be made up and is used to identify a unique computer in the LAN. Also, the name should not contain any blank spaces or punctuation. For example, the following are valid hostnames that could be assigned to each computer in a LAN consisting of 5 hosts: hostname 1 - Morpheus; hostname 2 - Trinity; hostname 3 - Tank; hostname 4 - Oracle; and hostname 5 - Dozer. Each of these hostnames conforms to the requirement that no blank spaces or punctuation marks are present. Use short hostnames to eliminate excessive typing, and choose a name that is easy to remember.

- Every host in the LAN will have the same network address, broadcast address, subnet mask, and domain name because those addresses identify the network in its entirety. Each computer in the LAN will have a hostname and IP address that uniquely identifies that particular host. The network address is 192.168.1.0, and the broadcast address is 192.168.1.128. Therefore, each host in the LAN must have an IP address between 192.168.1.1 to 192.168.1.127.

IP address	Example	Same/unique
Network address	192.168.1.0	Same for all hosts
Domain name	www.yourcompanyname.com	Same for all hosts
Broadcast address	192.168.1.128	Same for all hosts
Subnet mask	255.255.255.0	Same for all hosts
Hostname	Any valid name	Unique to each host
Host addresses	192.168.1.x	x must be unique to each host

SETTING UP LOCAL NETWORK SERVICES

Linux is increasingly popular in the computer networking/telecommunications industry. Acquiring the Linux operating system is a relatively simple and inexpensive task since virtually all of the source code can be downloaded from several different FTP or HTTP sites on the Internet. In addition, the most recent version of Red Hat Linux can be purchased from computer retail stores for between \$25 and \$50, depending on whether you purchase the standard or full version. The retail brand is indeed a worthwhile investment (vs. the free FTP or HTTP versions) since valuable technical support is included directly from the Red Hat Linux engineers for at least a year. This can be very helpful if, for instance, you can not resolve an installation/configuration problem after consulting the Red Hat Linux manuals.

This article describes how to put together a Local Area Network (LAN) consisting of two or more computers using the Red Hat Linux 6.2 operating system. A LAN is a communications network that interconnects a variety of devices and provides a means for exchanging information among those devices. The size and scope of a LAN is usually small, covering a single building or group of buildings. In a LAN, modems and phone lines are not required, and the computers should be close enough to run a network cable between them.

For each computer that will participate in the LAN, you'll need a network interface card (NIC) to which the network cable will be attached. You will also need to assign a unique hostname and IP address to each computer in the LAN (described later in this article), but this requires a basic understanding of TCP/IP (Transmission Control Protocol/Internet Protocol).

Introduction to TCP/IP

TCP/IP is the suite of protocols used by the Internet and most LANs throughout the world. In TCP/IP, every host (computer or other communications device) that is connected to the network has a unique IP address. An IP address is composed of four octets (numbers in the range of 0 to 255) separated by decimal points. The IP address is used to uniquely identify a host or computer on the LAN. For example, a computer with the hostname Morpheus could have an IP address of 192.168.7.127. You should avoid giving two or more computers the same IP address by using the range of IP addresses that are reserved for private, local area networks; this range of IP addresses usually begins with the octets 192.168.

LAN network address The first three octets of an IP address should be the same for all computers in the LAN. For example, if a total of 128 hosts exist in a single LAN, the IP addresses could be assigned starting with 192.168.1.x, where x represents a number in the range of 1 to 128. You could create consecutive LANs within the same company in a similar manner consisting of up to another 128 computers. Of course, you are not limited to 128 computers, as there are other ranges of IP addresses that allow you to build even larger networks.

There are different classes of networks that determine the size and total possible unique IP addresses of any given LAN. For example, a class A LAN can have over 16 million unique IP addresses. A class B LAN can have over 65,000 unique IP addresses. The size of your LAN depends on which reserved address range you use and the subnet mask (explained later in the article) associated with that range (see Table 1.).

Address range	Subnet mask	Provides	Addresses per LAN
10.0.0.0 - 10.255.255.255	255.0.0.0	1 class A LAN	16,777,216
172.16.0.0 - 172.31.255.255	255.255.0.0	16 class B LANs	65,536
192.168.0.0 - 192.168.255.255	25.255.255.0	256 class C LANs	256

Address ranges and LAN sizes

Network and broadcast addresses

Another important aspect of building a LAN is that the addresses at the two extreme ends of the address range are reserved for use as the LAN's network address and broadcast address. The network address is used by an application to represent the overall network. The broadcast address is used by an application to send the same message to all other hosts in the network simultaneously.

- For example, if you use addresses in the range of 192.168.1.0 to 192.168.1.128, the first address (192.168.1.0) is reserved as the network address, and the last address (192.168.1.128) is reserved as the broadcast address. Therefore, you only assign individual computers on the LAN IP addresses in the range of 192.168.1.1 to 192.168.1.127:

Network address: 192.168.1.0

Individual hosts: 192.168.1.1 to 192.168.1.127

Broadcast address: 192.168.1.128

Subnet masks

Each host in a LAN has a subnet mask. The subnet mask is an octet that uses the number 255 to represent the network address portion of the IP address and a zero to identify the host portion of the address. For example, the subnet mask 255.255.255.0 is used by each host to determine which LAN or class it belongs to. The zero at the end of the subnet mask represents a unique host within that network.

Assigning IP addresses in a LAN

There are two ways to assign IP addresses in a LAN. You can manually assign a static IP address to each computer in the LAN, or you can use a special type of server that automatically assigns a dynamic IP address to each computer as it logs into the network.

Static IP addressing

Static IP addressing means manually assigning a unique IP address to each computer in the LAN. The first three octets must be the same for each host, and the last digit must be a unique number for each host. In addition, a unique hostname will need to be assigned to each computer. Each host in the LAN will have the same network address (192.168.1.0), broadcast address (192.168.1.128), subnet mask (255.255.255.0), and domain name (yourcompanyname.com). It's a good idea to start by visiting each computer in the LAN and jotting down the hostname and IP address for future reference.

Dynamic IP addressing

Dynamic IP addressing is accomplished via a server or host called DHCP (Dynamic Host Configuration Program) that automatically assigns a unique IP address to each computer as it connects to the LAN. A similar service called BootP can also automatically assign unique IP addresses to each host in the network. The DHCP/ BootP service is a program or device that will act as a host with a unique IP address. An example of a DHCP device is a router that acts as an Ethernet hub (a communications device that allows multiple host to be connected via an Ethernet jack and a specific port) on one end and allows a connection to the Internet on the opposite end. Furthermore, the DHCP server will also assign the network and broadcast addresses. You will not be required to manually assign hostnames and domain names in a dynamic IP addressing scheme.

SETTING UP Xen , Vmware ON LINUX HOST AND ADDING GUEST OS

What Is VMware Player?

VMware Player is a free desktop application that lets you run virtual machines on a Windows or Linux PC.

VMware Player is the only product on the market that lets you run virtual machines without investing in virtualization software, making it easier than ever to take advantage of the security, flexibility, and portability of virtual machines. VMware Player lets you use host machine devices, such as CD and DVD drives, from the virtual machine.

VMware Player provides an intuitive user interface for running preconfigured virtual machines created with VMware Workstation, ESX Server, VMware Server, and GSX Server. On Windows host machines, VMware Player also opens and runs Microsoft Virtual PC and Virtual Server virtual machines and Symantec Backup Exec System Recovery (formerly LiveState Recovery) system images. VMware Player makes VMware virtual machines accessible to colleagues, partners, customers, and

clients, whether or not they have purchased VMware products. Anyone who downloads VMware Player can open and run compatible virtual machines.

What You Can Do with VMware Player

With VMware Player, you can:

Use and evaluate prebuilt applications-Download and safely run prebuilt application environments in virtual machines that are available from the Virtual Appliance Marketplace at <http://vam.vmware.com>. The Virtual Appliance Marketplace includes virtual machines from leading software vendors, including Oracle, Red Hat, Novell, BEA, SpikeSource, IBM, and MySQL, as well as virtual machines that are preconfigured with popular open source software.

Transform software distribution-Simplify software distribution by shipping preconfigured software in virtual machines. End users can experience the benefits of your products immediately, without setup hassles. VMware Player is ideal for shipping evaluation copies or beta software. You can package complex, sophisticated applications, complete with a full working environment, in a virtual machine that can be used by anyone who downloads VMware Player.

Collaborate with colleagues-VMware Player makes it easy for support, development, and QA to share customer scenarios in virtual machines.

Features in VMware Player

VMware Player is a free desktop application for running virtual machines. VMware Player does not include features found in other VMware products, such as the ability to create virtual machines.

VMware Player provides the following features:

You can connect, disconnect, and use configured host devices, including USB devices, in the virtual machine.

You can set preferences, such as how devices are displayed in VMware Player.

You can change the amount of memory allocated to the virtual machine.

You can drag and drop files between a Linux or Windows host and a Linux, Windows, or Solaris guest. (Linux hosts and Linux and Solaris guests must be running X Windows.) You can use this feature if the person who created the virtual machine you are running also installed VMware Tools in it.

You can copy and paste text between a Windows or Linux host and a Windows, Linux, or Solaris guest.

You can use this feature if the person who created the virtual machine you are running also installed VMware Tools in it.

You can copy and paste files between a Windows or Linux host and a Windows, Linux, or Solaris guest.

You can use this feature if the person who created the virtual machine you are running also installed VMware Tools in it.

To install VMware Player on a Linux host

- 1 Log on to your Linux host with the user name you plan to use when running VMware Player.
- 2 In a terminal window, become root so you can perform the initial installation steps:
- 3 Mount the VMware Player CD ROM.
- 4 Change to the Linux directory on the CD.
- 5 To use the RPM installer, skip to Step 6. To use the tar installer, follow these steps:

a. If you have a previous tar installation, delete the VMware Player distribution directory before installing from a tar file again. The default location of this directory is:

`/tmp/vmware-player-distrib`

- b Copy the tar archive to a temporary directory on your hard drive, for example,

`/tmp:`

`cp VMware-<xxxx>.tar.gz /tmp`

VMware-<xxxx>.tar.gz is the installation file. (In the filename, <xxxx-xxxx> is a series of numbers representing the version and build numbers.)

- c Change to the directory to which you copied the file: `cd /tmp`

d Unpack the archive:

`tar xzpf VMware-<xxxx>.tar.gz`

e Change to the installation directory:

`cd vmware-player-distrib`

f Run the installation program:

`./vmware-install.pl`

g Press return to accept the default values at the prompts.

h Press return (Yes) when prompted to run `vmware-config.pl`.

i Skip to Step 7.

Adding Guest OS

To install the OS from an ISO image in a virtual machine:

1. Save the ISO image file in any location accessible to your host. For example:

Windows: `C:\Temp` or `%TEMP%`

Linux: `/tmp` or `/usr/tmp`

Note: For best performance, place this image on the host computer's hard drive. However, to make the ISO image accessible to multiple users, you can also place the ISO image on a network share drive

(Windows) or exported filesystem (Linux). If your OS install spans multiple discs, you need to use an ISO image of each disc and place them all of them in a location accessible to the host.

2. Create a new virtual machine. Go to File > New > Virtual Machine.
3. Select Typical to accept Workstation's recommendations for various settings (such as processors, RAM, and disk controller type). Select Custom if you want to select these options yourself.
4. On the Guest Operating System Installation screen, when prompted where to install from, select Installer disc image file (iso).
5. Click Browse, and navigate to the location where you saved the ISO image file.
6. Click next, and proceed through the new virtual machine wizard.
7. Before you click Finish, to create the virtual machine, deselect Power on this virtual machine after creation.
8. Edit the virtual machine settings so that its virtual CD/DVD device is configured to use the ISO image rather than the physical CD/DVD drive:
 - a. Select the tab for the virtual machine you just created.
 - b. Click Edit virtual machine settings.
 - c. On the Hardware tab, select the CD/DVD drive.
 - d. On the right side:
 - i. Select Connect at power on.
 - ii. Use ISO image file.
 - iii. Click Browse and navigate to where you saved the ISO image file.
 - e. Click OK.
9. Power on the virtual machine.

When you are finished installing the guest OS, you can edit the virtual machine settings so that it is once more set to use the host computer's physical drive. You do not need to leave the drive set to connect at power on.