

➤ INTRODUCTION

This chapter provides an introduction to Computer networks and covers fundamental topics like data, information to the definition of communication and computer networks. The main objective of data communication and networking is to enable seamless exchange of data between any two points in the world. This exchange of data takes place over a computer network.

DATA & INFORMATION

Data refers to the raw facts that are collected while information refers to processed data that enables us to take decisions. Ex. When result of a particular test is declared it contains data of all students, when you find the marks you have scored you have the information that lets you know whether you have passed or failed. The word data refers to any information which is presented in a form that is agreed and accepted upon by its creators and users.

DATA COMMUNICATION

Data Communication is a process of exchanging data or information. In case of computer networks this exchange is done between two devices over a transmission medium. This process involves a communication system which is made up of hardware and software. The hardware part involves the sender and receiver devices and the intermediate devices through which the data passes. The software part involves certain rules which specify what is to be communicated, how it is to be communicated and when. It is also called as a Protocol.

Characteristics of Data Communication

The effectiveness of any data communications system depends upon the following four fundamental characteristics:

1. **Delivery:** The data should be delivered to the correct destination and correct user.
2. **Accuracy:** The communication system should deliver the data accurately, without introducing any errors. The data may get corrupted during transmission affecting the accuracy of the delivered data.
3. **Timeliness:** Audio and Video data has to be delivered in a timely manner without any delay; such a data delivery is called real time transmission of data.
4. **Jitter:** It is the variation in the packet arrival time. Uneven Jitter may affect the timeliness of data being transmitted.

Components of Data Communication

A Data Communication system has five components as shown in the diagram below: Fig. Components of a Data Communication System

1. **Message** Message is the information to be communicated by the sender to the receiver.
2. **Sender** The sender is any device that is capable of sending the data (message).
3. **Receiver** The receiver is a device that the sender wants to communicate the data (message).
4. **Transmission Medium** It is the path by which the message travels from sender to receiver. It can be wired or wireless and many subtypes in both.
5. **Protocol** It is an agreed upon set or rules used by the sender and receiver to communicate data. A protocol is a set of rules that governs data communication. A Protocol is a necessity in data communications without which the communicating entities are like two persons trying to talk to each other in a different language without knowing the other language.



DATA REPRESENTATION

Data is collection of raw facts which is processed to deduce information. There may be different forms in which data may be represented. Some of the forms of data used in communications are as follows:

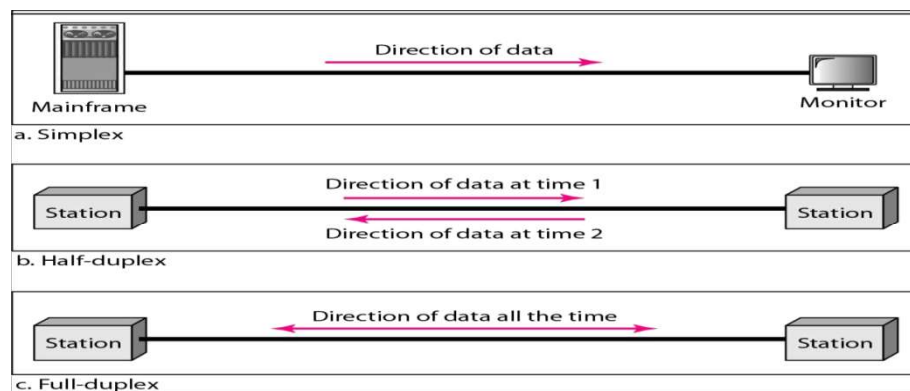
1. Text: Text includes combination of alphabets in small case as well as upper case. It is stored as a pattern of bits. Prevalent encoding system: ASCII, Unicode
2. Numbers: Numbers include combination of digits from 0 to 9. It is stored as a pattern of bits. Prevalent encoding system: ASCII, Unicode
3. Images —An image is worth a thousand words. In computers images are digitally stored. A Pixel is the smallest element of an image. To put it in simple terms, a picture or image is a matrix of pixel elements. The pixels are represented in the form of bits. Depending upon the type of image (black n white or color) each pixel would require different number of bits to represent the value of a pixel. The size of an image depends upon the number of pixels (also called resolution) and the bit pattern used to indicate the value of each pixel. Example: if an image is purely black and white (two color) each pixel can be represented by a value either 0 or 1, so an image made up of 10 x 10 pixel elements would require only 100 bits in memory to be stored. On the other hand an image that includes gray may require 2 bits to represent every pixel value (00 - black, 01 – dark gray, 10 – light gray, 11 –white). So the same 10 x 10 pixel image would now require 200 bits of memory to be stored. Commonly used Image formats: jpg, png, bmp, etc
4. Audio: Data can also be in the form of sound which can be recorded and broadcasted. Example: What we hear on the radio is a source of data or information. Audio data is continuous, not discrete.
5. Video: Video refers to broadcasting of data in form of picture or movie

Data Flow

In **simplex mode(SX)**, the communication is unidirectional, as on a one-way street. Only one of the two devices on a link can transmit; the other can only receive. Commercial radio broadcasting is an example. Simplex lines are also called receive-only, transmit-only or one-way-only lines.

In **half-duplex(HDX)** mode, each station can both transmit and receive, but not at the same time. When one device is sending, the other can only receive, and vice versa. The half-duplex mode is used in case where there is no need for communication in both directions at the same time; the entire capacity of the channel can be utilized for each direction. Citizens band (CB) radio is an example where push to talk (PTT) is to be pressed or depressed while sending and transmitting.

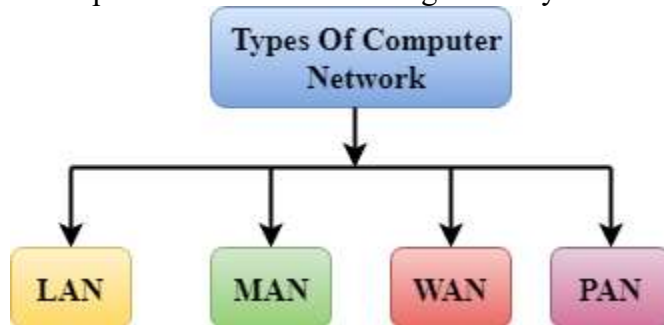
In **full-duplex mode(FDX)** (called duplex), both stations can transmit and receive simultaneously. One common example of full-duplex communication is the telephone network. The full-duplex mode is used when communication in both directions is required all the time. The capacity of the channel must be divided between the two directions.



Categories of Networks

A computer network is a group of computers linked to each other that enables the computer to communicate with another computer and share their resources, data, and applications.

A computer network can be categorized by their size. A **computer network** is mainly of **four types**:



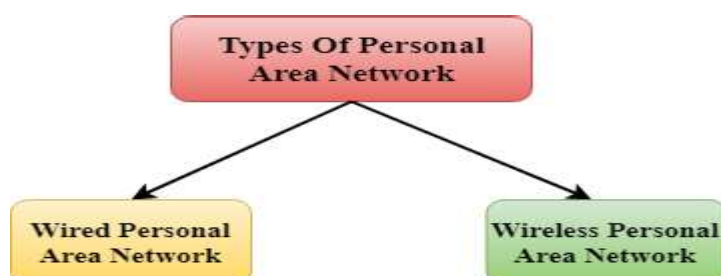
- LAN(Local Area Network)
- PAN(Personal Area Network)
- MAN(Metropolitan Area Network)
- WAN(Wide Area Network)

LAN(Local Area Network)

- Local Area Network is a group of computers connected to each other in a small area such as building, office.
- LAN is used for connecting two or more personal computers through a communication medium such as twisted pair, coaxial cable, etc.
- It is less costly as it is built with inexpensive hardware such as hubs, network adapters, and ethernet cables.
- The data is transferred at an extremely faster rate in Local Area Network.
- Local Area Network provides higher security.

PAN(Personal Area Network)

- Personal Area Network is a network arranged within an individual person, typically within a range of 10 meters.
- Personal Area Network is used for connecting the computer devices of personal use is known as Personal Area Network.
- **Thomas Zimmerman** was the first research scientist to bring the idea of the Personal Area Network.
- Personal Area Network covers an area of **30 feet**.
- Personal computer devices that are used to develop the personal area network are the laptop, mobile phones, media player and play stations.



There are two types of Personal Area Network:

- Wired Personal Area Network
- Wireless Personal Area Network

Wireless Personal Area Network: Wireless Personal Area Network is developed by simply using wireless technologies such as WiFi, Bluetooth. It is a low range network.

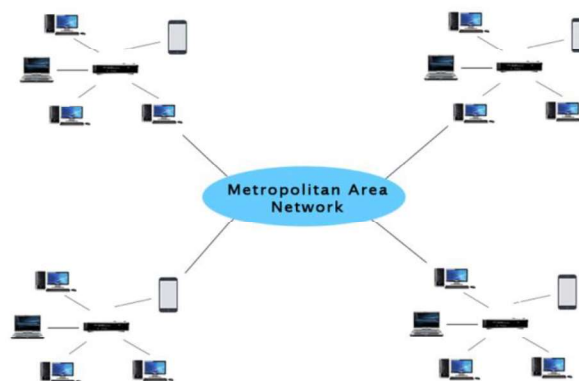
Wired Personal Area Network: Wired Personal Area Network is created by using the USB.

Examples Of Personal Area Network:

- **Body Area Network:** Body Area Network is a network that moves with a person. **For example**, a mobile network moves with a person. Suppose a person establishes a network connection and then creates a connection with another device to share the information.
- **Offline Network:** An offline network can be created inside the home, so it is also known as a **home network**. A home network is designed to integrate the devices such as printers, computer, television but they are not connected to the internet.
- **Small Home Office:** It is used to connect a variety of devices to the internet and to a corporate network using a VPN

MAN(Metropolitan Area Network)

- A metropolitan area network is a network that covers a larger geographic area by interconnecting a different LAN to form a larger network.
- Government agencies use MAN to connect to the citizens and private industries.
- In MAN, various LANs are connected to each other through a telephone exchange line.
- The most widely used protocols in MAN are RS-232, Frame Relay, ATM, ISDN, OC-3, ADSL, etc.



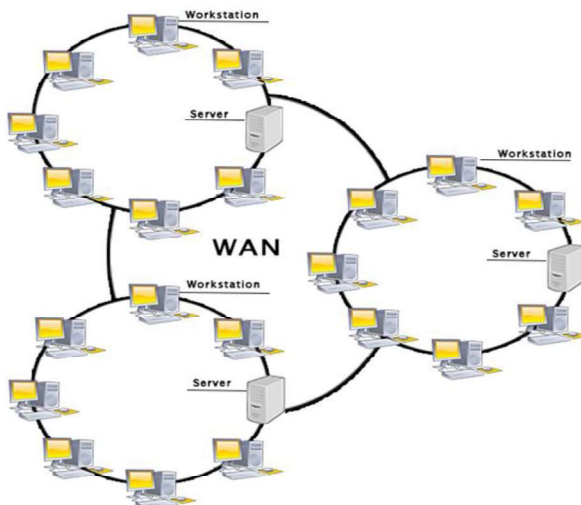
- It has a higher range than Local Area Network(LAN).

Uses Of Metropolitan Area Network:

- MAN is used in communication between the banks in a city.
- It can be used in an Airline Reservation.
- It can be used in a college within a city.
- It can also be used for communication in the military.

WAN(Wide Area Network)

- A Wide Area Network is a network that extends over a large geographical area such as states or countries.
- A Wide Area Network is quite bigger network than the LAN.
- A Wide Area Network is not limited to a single location, but it spans over a large geographical area through a telephone line, fibre optic cable or satellite links.
- The internet is one of the biggest WAN in the world.



- A Wide Area Network is widely used in the field of Business, government, and education.

Examples Of Wide Area Network:

- **Mobile Broadband:** A 4G network is widely used across a region or country.
- **Last mile:** A telecom company is used to provide the internet services to the customers in

hundreds of cities by connecting their home with fiber.

- **Private network:** A bank provides a private network that connects the 44 offices. This network is made by using the telephone leased line provided by the telecom company.

Advantages Of Wide Area Network:

Following are the advantages of the Wide Area Network:

- **Geographical area:** A Wide Area Network provides a large geographical area. Suppose if the branch of our office is in a different city then we can connect with them through WAN. The internet provides a leased line through which we can connect with another branch.
- **Centralized data:** In case of WAN network, data is centralized. Therefore, we do not need to buy the emails, files or back up servers.
- **Get updated files:** Software companies work on the live server. Therefore, the programmers get the updated files within seconds.
- **Exchange messages:** In a WAN network, messages are transmitted fast. The web application like Facebook, Whatsapp, Skype allows you to communicate with friends.
- **Sharing of software and resources:** In WAN network, we can share the software and other resources like a hard drive, RAM.
- **Global business:** We can do the business over the internet globally.
- **High bandwidth:** If we use the leased lines for our company then this gives the high bandwidth. The high bandwidth increases the data transfer rate which in turn increases the productivity of our company.

Disadvantages of Wide Area Network:

The following are the disadvantages of the Wide Area Network:

- **Security issue:** A WAN network has more security issues as compared to LAN and MAN network as all the technologies are combined together that creates the security problem.
- **Needs Firewall & antivirus software:** The data is transferred on the internet which can be changed or hacked by the hackers, so the firewall needs to be used. Some people can inject the virus in our system so antivirus is needed to protect from such a virus.
- **High Setup cost:** An installation cost of the WAN network is high as it involves the purchasing of routers, switches.
- **Troubleshooting problems:** It covers a large area so fixing the problem is difficult.

Internetwork

- An internetwork is defined as two or more computer network LANs or WAN or computer network segments are connected using devices, and they are configured by a local addressing scheme. This process is known as **internetworking**.
- An interconnection between public, private, commercial, industrial, or government computer networks can also be defined as **internetworking**.
- An internetworking uses the **internet protocol**.
- The reference model used for internetworking is **Open System Interconnection (OSI)**.

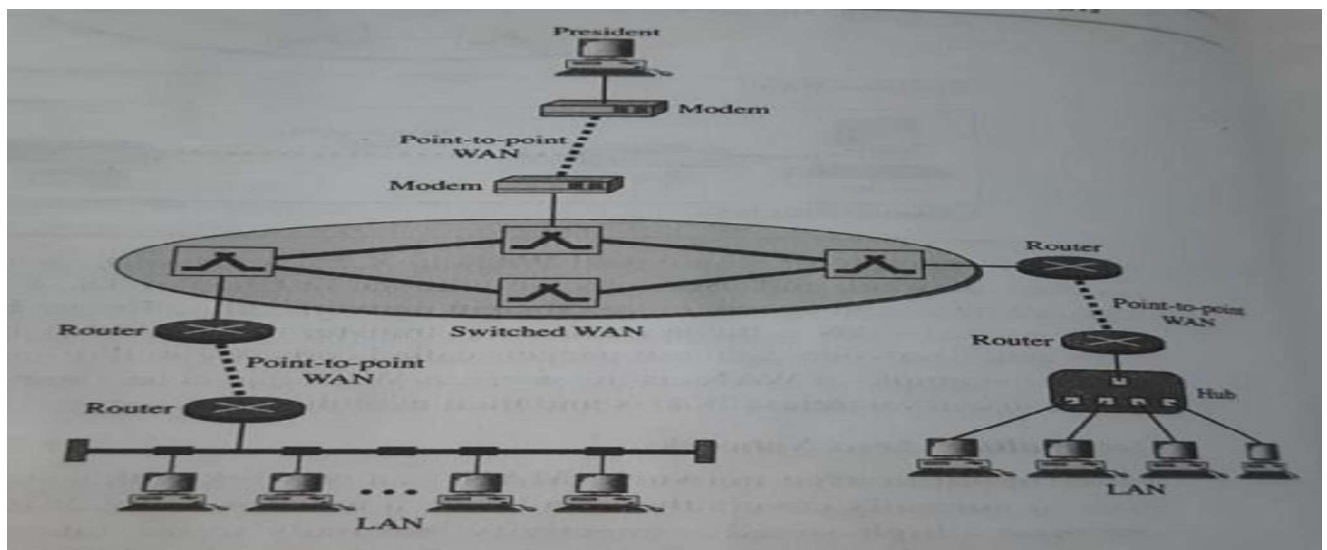
Types Of Internetwork:

1. Extranet: An extranet is a communication network based on the internet protocol such as **Transmission Control protocol** and **internet protocol**. It is used for information sharing. The access to the extranet is restricted to only those users who have login credentials. An extranet is the lowest level of internetworking. It can be categorized as **MAN**, **WAN** or other computer networks. An extranet cannot have a single **LAN**, atleast it must have one connection to the external network.

2. **Intranet:** An intranet is a private network based on the internet protocol such as **Transmission Control protocol** and **internet protocol**. An intranet belongs to an organization which is only accessible by the **organization's employee** or members. The main aim of the intranet is to share the information and resources among the organization employees. An intranet provides the facility to work in groups and for teleconferences.

Intranet advantages:

- **Communication:** It provides a cheap and easy communication. An employee of the organization can communicate with another employee through email, chat.
- **Time-saving:** Information on the intranet is shared in real time, so it is time-saving.
- **Collaboration:** Collaboration is one of the most important advantage of the intranet. The information is distributed among the employees of the organization and can only be accessed by the authorized user.
- **Platform independency:** It is a neutral architecture as the computer can be connected to another device with different architecture.
- **Cost effective:** People can see the data and documents by using the browser and distributes the duplicate copies over the intranet. This leads to a reduction in the cost



➤ Network Models

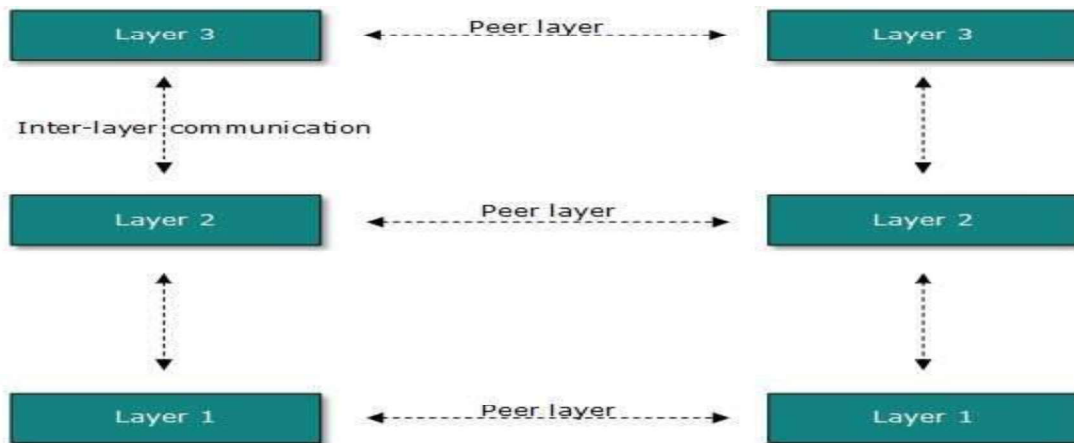
A Network is a complex piece of Hardware and software. Early attempts for implementing the software for such subsystems were based on a single, complex, unstructured program with many interacting components. The resultant software was very difficult to test and modify. To overcome such problem, the ISO has developed a layered approach. In a layered approach, networking concept is divided into several layers, and each layer is assigned a particular task. Therefore, we can say that networking tasks depend upon the layers.

Layered Tasks

In layered architecture of Network Model, one whole network process is divided into small tasks. Each small task is then assigned to a particular layer which works dedicatedly to process the task only. Every layer does only specific work.

In layered communication system, one layer of a host deals with the task done by or to be done by its peer layer at the same level on the remote host. The task is either initiated by layer at the lowest level or at the top most level. If the task is initiated by the-top most layer, it is passed on to the layer below it for further

processing. The lower layer does the same thing, it processes the task and passes on to lower layer. If the task is initiated by lower most layer, then the reverse path is taken.



Every layer clubs together all procedures, protocols, and methods which it requires to execute its piece of task. All layers identify their counterparts by means of encapsulation header and tail

➤ Internet Model

Internet uses TCP/IP protocol suite, also known as Internet suite. This defines Internet Model which contains four layered architecture. OSI Model is general communication model but Internet Model is what the internet uses for all its communication. The internet is independent of its underlying network architecture so is its Model. This model has the following layers:

Application Layer: This layer defines the protocol which enables user to interact with the network. For example, FTP, HTTP etc.

- **Transport Layer:** This layer defines how data should flow between hosts. Major protocol at this layer is Transmission Control Protocol (TCP). This layer ensures data delivered between hosts is in-order and is responsible for end-to-end delivery.
- **Internet Layer:** Internet Protocol (IP) works on this layer. This layer facilitates host addressing and recognition. This layer defines routing.
- **Link Layer:** This layer provides mechanism of sending and receiving actual data. Unlike its OSI Model counterpart, this layer is independent of underlying network architecture and hardware



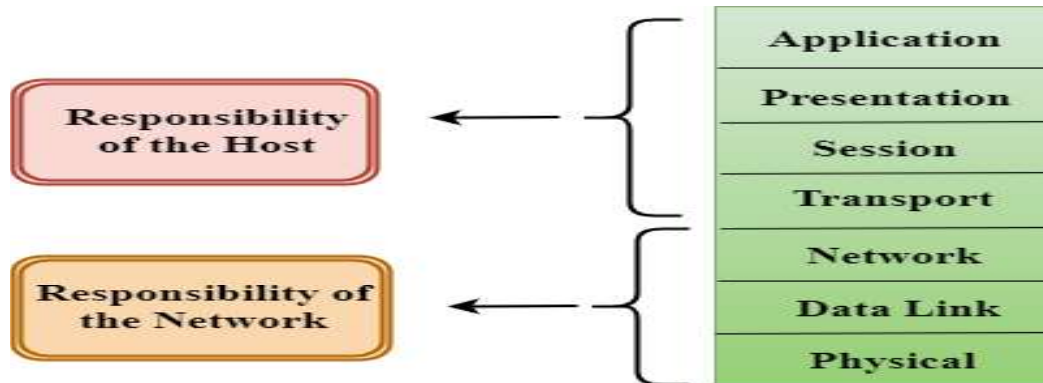
•

➤ OSI Model

- OSI stands for **Open System Interconnection** is a reference model that describes how information from a software application in one computer moves through a physical medium to the software application in another computer.

- OSI consists of seven layers, and each layer performs a particular network function.
- OSI model was developed by the International Organization for Standardization (ISO) in 1984, and it is now considered as an architectural model for the inter-computer communications.
- OSI model divides the whole task into seven smaller and manageable tasks. Each layer is assigned a particular task.
- Each layer is self-contained, so that task assigned to each layer can be performed independently.

Characteristics of OSI Model:

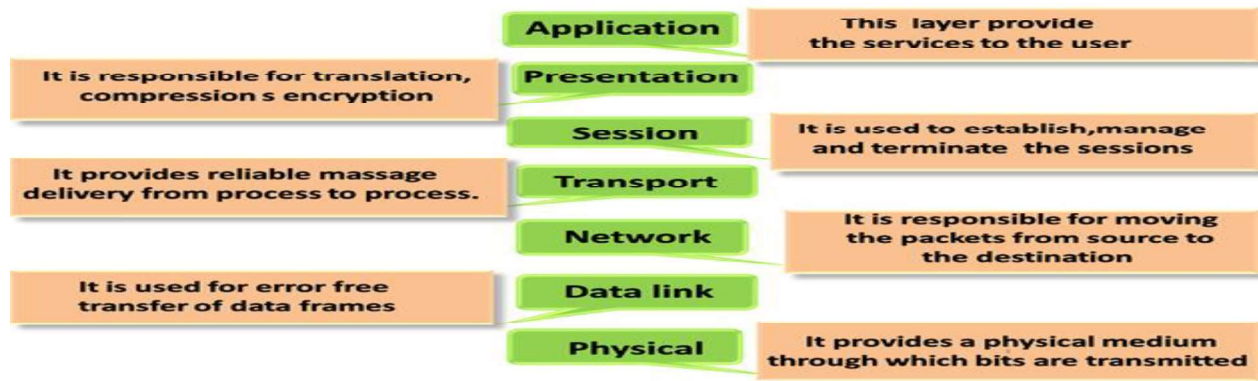


- The OSI model is divided into two layers: upper layers and lower layers.
- The upper layer of the OSI model mainly deals with the application related issues, and they are implemented only in the software. The application layer is closest to the end user. Both the end user and the application layer interact with the software applications. An upper layer refers to the layer just above another layer.
- The lower layer of the OSI model deals with the data transport issues. The data link layer and the physical layer are implemented in hardware and software. The physical layer is the lowest layer of the OSI model and is closest to the physical medium. The physical layer is mainly responsible for placing the information on the physical medium.

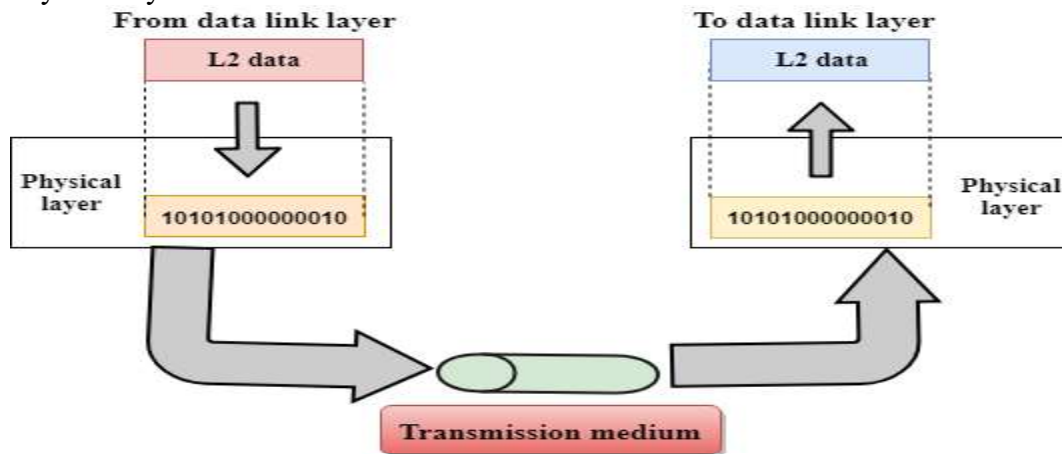
Functions of the OSI Layers

There are the seven OSI layers. Each layer has different functions. A list of seven layers are given below:

1. Physical Layer
2. Data-Link Layer
3. Network Layer
4. Transport Layer
5. Session Layer
6. Presentation Layer
7. Application Layer



Physical layer



- The main functionality of the physical layer is to transmit the individual bits from one node to another node.
- It is the lowest layer of the OSI model.
- It establishes, maintains and deactivates the physical connection.
- It specifies the mechanical, electrical and procedural network interface specifications.

➤ **Functions of a Physical layer:**

- **Line Configuration:** It defines the way how two or more devices can be connected physically.
- **Data Transmission:** It defines the transmission mode whether it is simplex, half-duplex or full-duplex mode between the two devices on the network.
- **Topology:** It defines the way how network devices are arranged.
- **Signals:** It determines the type of the signal used for transmitting the information.

Physical layer in the OSI model plays the role of interacting with actual hardware and signaling mechanism. Physical layer is the only layer of OSI network model which actually deals with the physical connectivity of two different stations. This layer defines the hardware equipment, cabling, wiring, frequencies, pulses used to represent binary signals etc.

Physical layer provides its services to Data-link layer. Data-link layer hands over frames to physical layer. Physical layer converts them to electrical pulses, which represent binary data. The binary data is then sent over the wired or wireless media.

➤ **Signals**

When data is sent over physical medium, it needs to be first converted into electromagnetic signals. Data itself can be analog such as human voice, or digital such as file on the disk. Both analog and digital data can be represented in digital or analog signals.

- **Digital Signals**

Digital signals are discrete in nature and represent sequence of voltage pulses. Digital signals are used within the circuitry of a computer system.

- **Analog Signals**

Analog signals are in continuous wave form in nature and represented by continuous electromagnetic waves.

Transmission Impairment

When signals travel through the medium they tend to deteriorate. This may have many reasons as given:

- **Attenuation**

For the receiver to interpret the data accurately, the signal must be sufficiently strong. When the signal passes through the medium, it tends to get weaker. As it covers distance, it loses strength.

- **Dispersion**

As signal travels through the media, it tends to spread and overlaps. The amount of dispersion depends upon the frequency used.

- **Delay distortion**

Signals are sent over media with pre-defined speed and frequency. If the signal speed and frequency do not match, there are possibilities that signal reaches destination in arbitrary fashion. In digital media, this is very critical that some bits reach earlier than the previously sent ones.

- **Noise**

Random disturbance or fluctuation in analog or digital signal is said to be Noise in signal, which may distort the actual information being carried. Noise can be characterized in one of the following class:

- **Thermal Noise**

Heat agitates the electronic conductors of a medium which may introduce noise in the media. Up to a certain level, thermal noise is unavoidable.

- **Intermodulation**

When multiple frequencies share a medium, their interference can cause noise in the medium. Intermodulation noise occurs if two different frequencies are sharing a medium and one of them has excessive strength or the component itself is not functioning properly, then the resultant frequency may not be delivered as expected.

- **Crosstalk**

This sort of noise happens when a foreign signal enters into the media. This is because signal in one medium affects the signal of second medium.

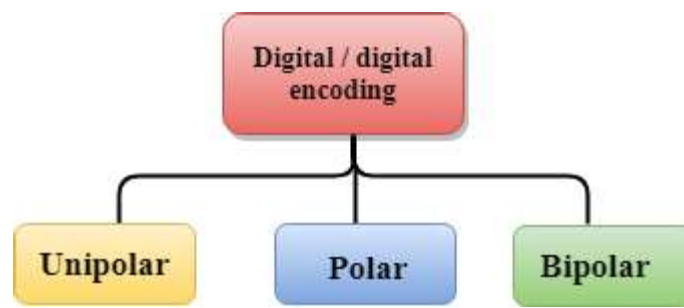
- **Impulse**

This noise is introduced because of irregular disturbances such as lightening, electricity, short-circuit, or faulty components. Digital data is mostly affected by this sort of noise.

➤ **Digital Transmission**

Data can be represented either in analog or digital form. The computers used the digital form to store the information. Therefore, the data needs to be converted in digital form so that it can be used by a computer.

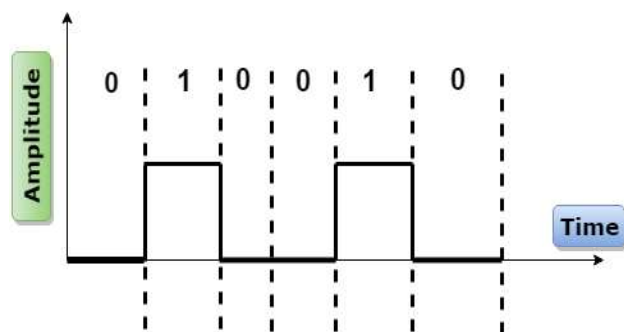
Digital-To-Digital Conversion



- Unipolar Encoding
- Polar Encoding
- Bipolar Encoding

Digital-to-digital encoding is the representation of digital information by a digital signal. When binary 1s and 0s generated by the computer are translated into a sequence of voltage pulses that can be propagated over a wire, this process is known as digital-to-digital encoding.

Digital-to-digital encoding is divided into three categories:

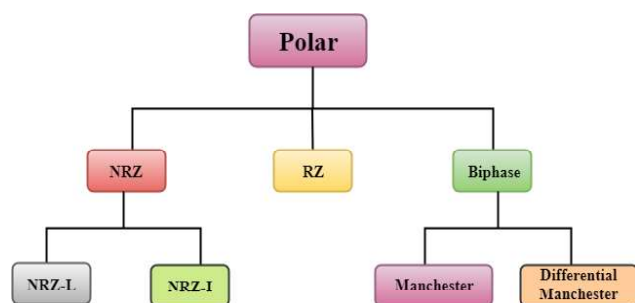


Unipolar

- Digital transmission system sends the voltage pulses over the medium link such as wire or cable.
- In most types of encoding, one voltage level represents 0, and another voltage level represents 1.
- The polarity of each pulse determines whether it is positive or negative.
- This type of encoding is known as Unipolar encoding as it uses only one polarity.

- In Unipolar encoding, the polarity is assigned to the 1 binary state.
- In this, 1s are represented as a positive value and 0s are represented as a zero value.
- In Unipolar Encoding, '1' is considered as a high voltage and '0' is considered as a zero voltage.
- Unipolar encoding is simpler and inexpensive to implement.

Unipolar encoding has two problems that make this scheme less desirable:



Polar

- Polar encoding is an encoding scheme that uses two voltage levels: one is positive, and another is negative.
- By using two voltage levels, an average voltage level is reduced, and the DC component problem of unipolar encoding scheme is alleviated.

NRZ

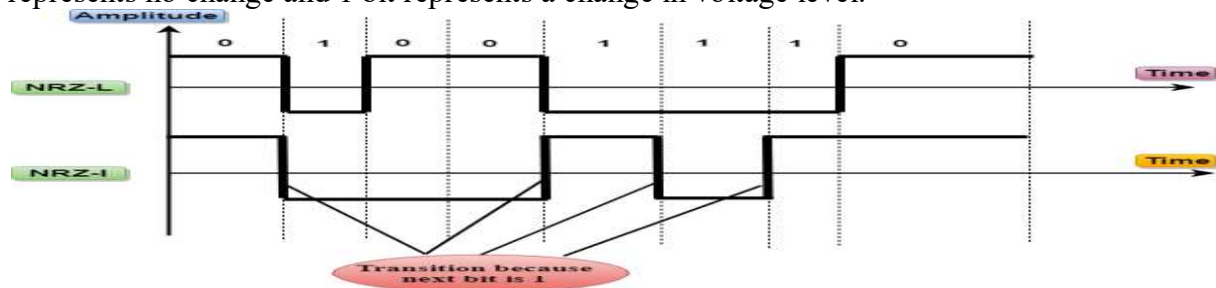
- NRZ stands for Non-return zero.
- In NRZ encoding, the level of the signal can be represented either positive or negative.

The two most common methods used in NRZ are:

NRZ-L: In NRZ-L encoding, the level of the signal depends on the type of the bit that it represents. If a bit is 0 or 1, then their voltages will be positive and negative respectively. Therefore, we can say that the level of the signal is dependent on the state of the bit.

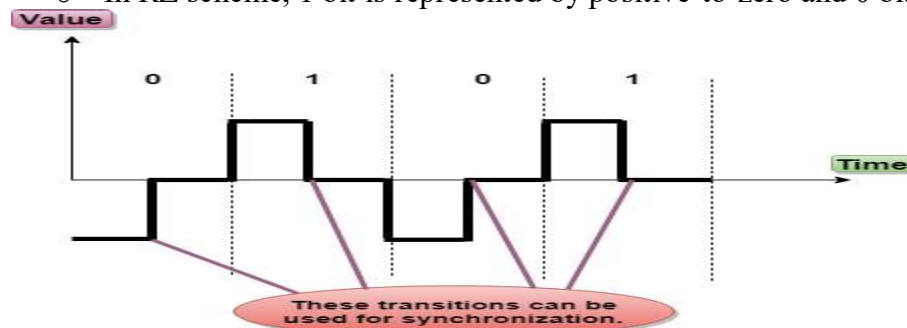
NRZ-I: NRZ-I is an inversion of the voltage level that represents 1 bit. In the NRZ-I encoding scheme, a transition occurs between the positive and negative voltage that represents 1 bit. In this scheme, 0 bit

represents no change and 1 bit represents a change in voltage level.



RZ

- RZ stands for Return to zero.
- There must be a signal change for each bit to achieve synchronization. However, to change with every bit, we need to have three values: positive, negative and zero.
- RZ is an encoding scheme that provides three values, positive voltage represents 1, the negative voltage represents 0, and zero voltage represents none.
- In the RZ scheme, halfway through each interval, the signal returns to zero.
- In RZ scheme, 1 bit is represented by positive-to-zero and 0 bit is represented by negative-to-zero.



Disadvantage of RZ:

It performs two signal changes to encode one bit that acquires more bandwidth.

Biphase

- Biphase is an encoding scheme in which signal changes at the middle of the bit interval but does not return to zero.

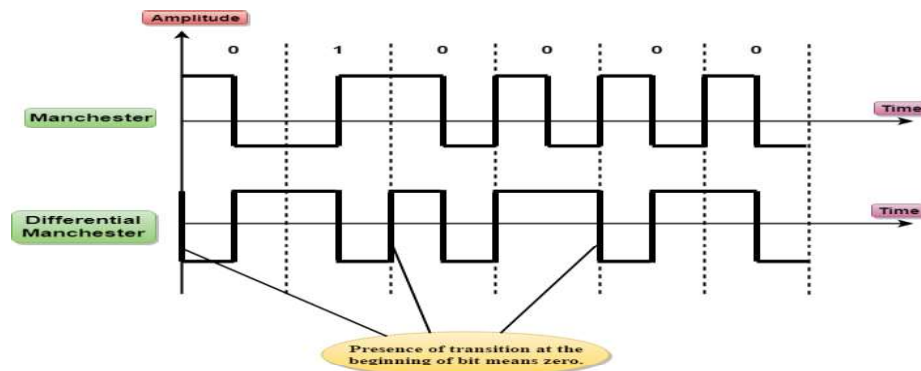
Biphase encoding is implemented in two different ways:

Manchester

- It changes the signal at the middle of the bit interval but does not return to zero for synchronization.
- In Manchester encoding, a negative-to-positive transition represents binary 1, and positive-to-negative transition represents 0.
- Manchester has the same level of synchronization as RZ scheme except that it has two levels of amplitude.

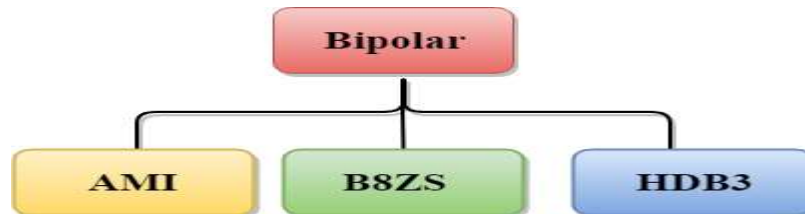
Differential Manchester

- It changes the signal at the middle of the bit interval for synchronization, but the presence or absence of the transition at the beginning of the interval determines the bit. A transition means binary 0 and no transition means binary 1.
- In Manchester Encoding scheme, two signal changes represent 0 and one signal change represent 1.



Bipolar

- Bipolar encoding scheme represents three voltage levels: positive, negative, and zero.
- In Bipolar encoding scheme, zero level represents binary 0, and binary 1 is represented by alternating positive and negative voltages.
- If the first 1 bit is represented by positive amplitude, then the second 1 bit is represented by negative voltage, third 1 bit is represented by the positive amplitude and so on. This alternation can also occur even when the 1 bits are not consecutive.



Bipolar can be classified as:

AMI

- AMI stands for **alternate mark inversion** where mark work comes from telegraphy which means 1. So, it can be redefined as **alternate 1 inversion**.

- In Bipolar AMI encoding scheme, 0 bit is represented by zero level and 1 bit is represented by alternating positive and negative voltages.

Advantage:

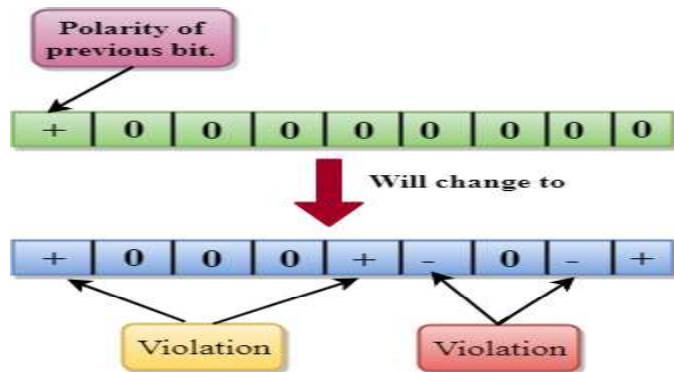
- DC component is zero.
- Sequence of 1s bits are synchronized.

Disadvantage:

- This encoding scheme does not ensure the synchronization of a long string of 0s bits.

B8ZS

- B8ZS stands for **Bipolar 8-Zero Substitution**.
- This technique is adopted in North America to provide synchronization of a long sequence of 0s bits.
- In most of the cases, the functionality of B8ZS is similar to the bipolar AMI, but the only difference is that it provides the synchronization when a long sequence of 0s bits occur.
- B8ZS ensures synchronization of a long string of 0s by providing force artificial signal changes called violations, within 0 string pattern.
- When eight 0 occurs, then B8ZS implements some changes in 0s string pattern based on the polarity of the previous 1 bit.
- If the polarity of the previous 1 bit is positive, the eight 0s will be encoded as zero, zero, zero, positive, negative, zero, negative, positive.

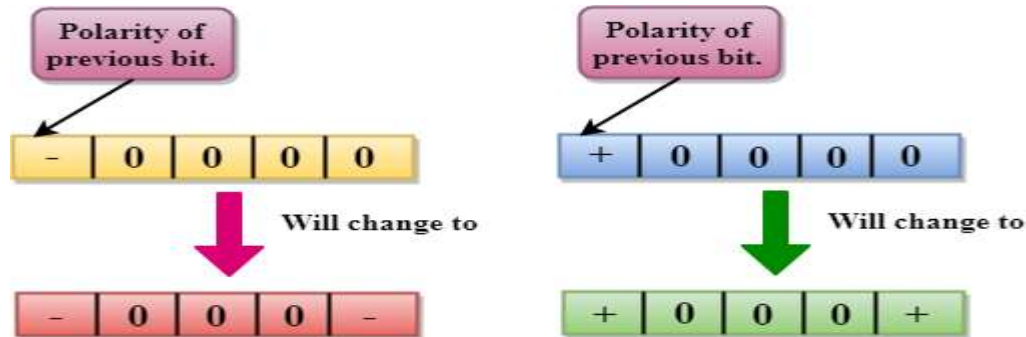


- If the polarity of previous 1 bit is negative, then the eight 0s will be encoded as zero, zero, zero, zero, positive, negative, zero, positive, negative.

HDB3

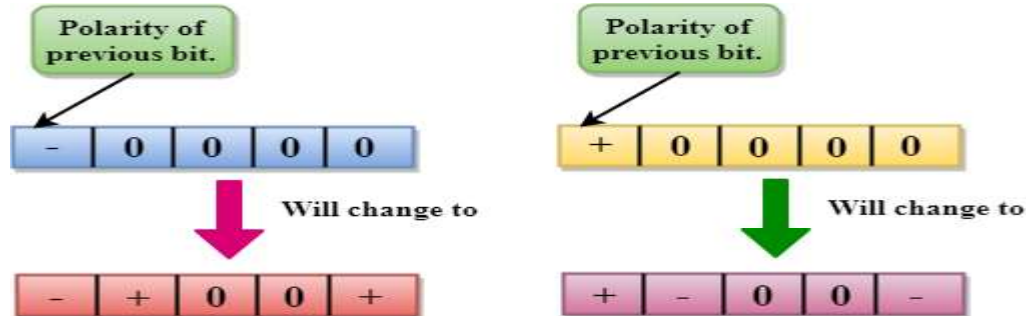
- HDB3 stands for **High-Density Bipolar 3**.
- HDB3 technique was first adopted in Europe and Japan.
- HDB3 technique is designed to provide the synchronization of a long sequence of 0s bits.
- In the HDB3 technique, the pattern of violation is based on the polarity of the previous bit.
- When four 0s occur, HDB3 looks at the number of 1s bits occurred since the last substitution.
- If the number of 1s bits is odd, then the violation is made on the fourth consecutive of 0. If the polarity of the previous bit is positive, then the violation is positive. If the polarity of the previous bit is negative, then the violation is negative.

If the number of 1s bits since the last substitution is odd.



If the number of 1s bits is even, then the violation is made on the place of the first and fourth consecutive 0s. If the polarity of the previous bit is positive, then violations are negative, and if the polarity of the previous bit is negative, then violations are positive.

If the number of 1s bits since the last substitution is even.



Analog-To-Digital Conversion

Microphones create analog voice and camera creates analog videos, which are treated as analog data. To

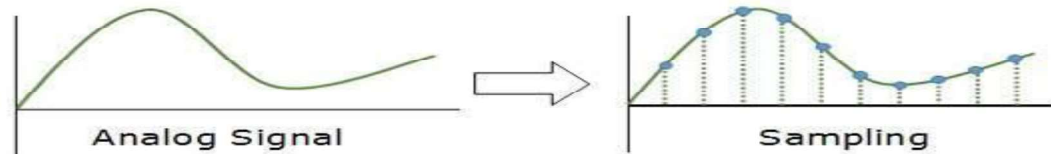
transmit this analog data over digital signals, we need analog to digital conversion.

Analog data is a continuous stream of data in the wave form whereas digital data is discrete. To convert analog wave into digital data, we use Pulse Code Modulation (PCM).

PCM is one of the most commonly used method to convert analog data into digital form. It involves three steps:

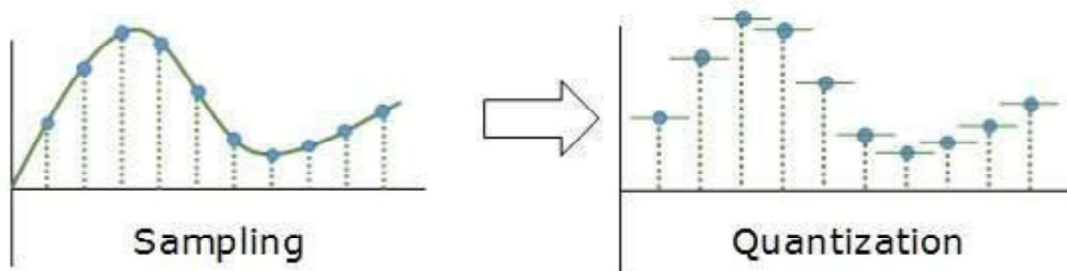
- Sampling
- Quantization
- Encoding.

Sampling



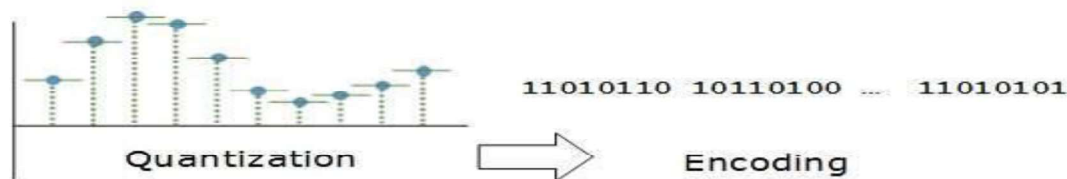
The analog signal is sampled every T interval. Most important factor in sampling is the rate at which analog signal is sampled. According to Nyquist Theorem, the sampling rate must be at least two times of the highest frequency of the signal.

Quantization



Sampling yields discrete form of continuous analog signal. Every discrete pattern shows the amplitude of the analog signal at that instance. The quantization is done between the maximum amplitude value and the minimum amplitude value. Quantization is approximation of the instantaneous analog value.

Encoding



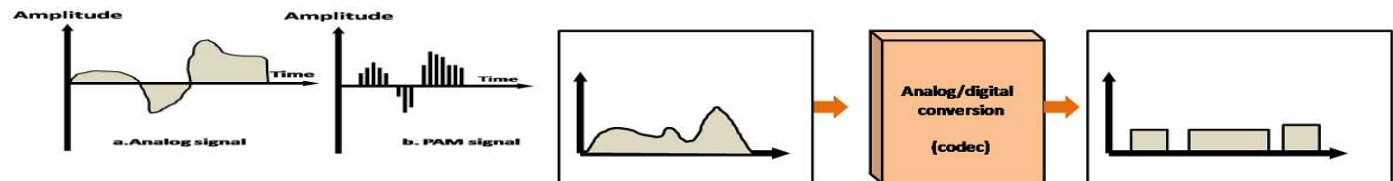
In encoding, each approximated value is then converted into binary format.

- When an analog signal is digitalized, this is called an analog-to-digital conversion.
- Suppose human sends a voice in the form of an analog signal, we need to digitalize the analog signal which is less prone to noise. It requires a reduction in the number of values in an analog message so that they can be represented in the digital stream.
- In analog-to-digital conversion, the information contained in a continuous wave form is converted in digital pulses.

Techniques for Analog-To-Digital Conversion

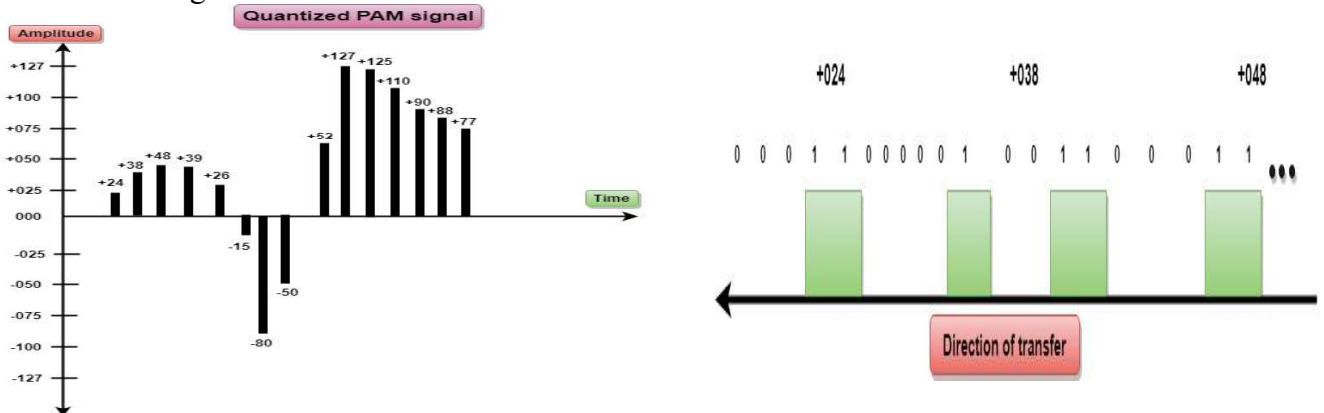
PAM

- PAM stands for **pulse amplitude modulation**.
- PAM is a technique used in analog-to-digital conversion.
- PAM technique takes an analog signal, samples it, and generates a series of digital pulses based on the result of sampling where sampling means measuring the amplitude of a signal at equal intervals.
- PAM technique is not useful in data communication as it translates the original wave form into pulses, but these pulses are not digital. To make them digital, PAM technique is modified to PCM technique.



PCM

- PCM stands for **Pulse Code Modulation**.
- PCM technique is used to modify the pulses created by PAM to form a digital signal. To achieve this, PCM quantizes PAM pulses. Quantization is a process of assigning integral values in a specific range to sampled instances.
- PCM is made of four separate processes: PAM, quantization, binary encoding, and digital-to-digital encoding.



PCM

➤ Analog Transmission

To send the digital data over an analog media, it needs to be converted into analog signal. There can be two cases according to data formatting.

Bandpass: The filters are used to filter and pass frequencies of interest. A bandpass is a band of frequencies which can pass the filter.

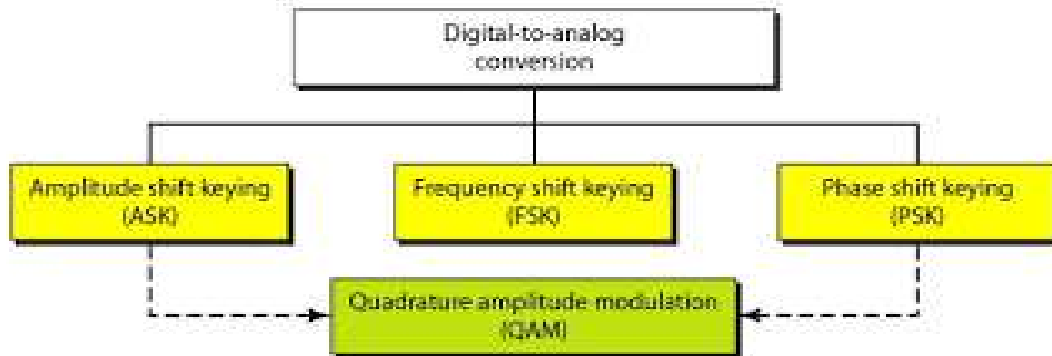
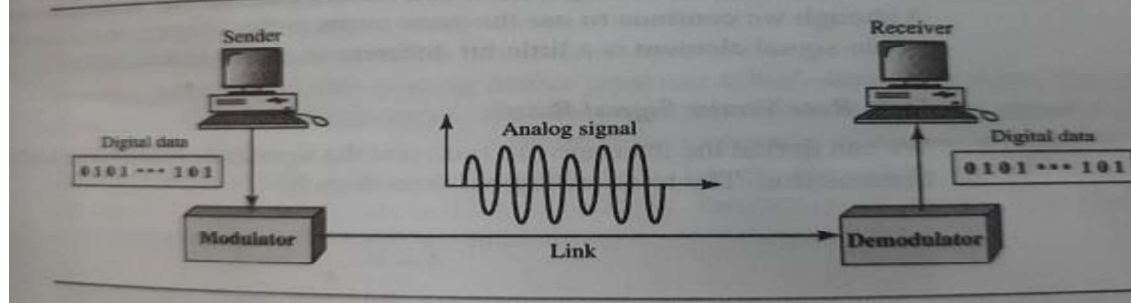
Low-pass: Low-pass is a filter that passes low frequencies signals.

When digital data is converted into a bandpass analog signal, it is called digital-to-analog conversion. When low-pass analog signal is converted into bandpass analog signal, it is called analog-to-analog conversion.

Digital-to-Analog Conversion

When data from one computer is sent to another via some analog carrier, it is first converted into analog signals. Analog signals are modified to reflect digital data.

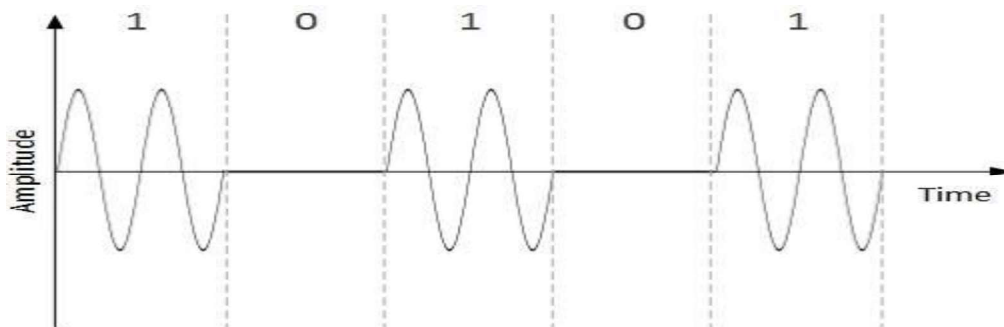
Figure 5.1 Digital-to-analog conversion



An analog signal is characterized by its amplitude, frequency, and phase. There are three kinds of digital-to-analog conversions:

- **Amplitude Shift Keying**

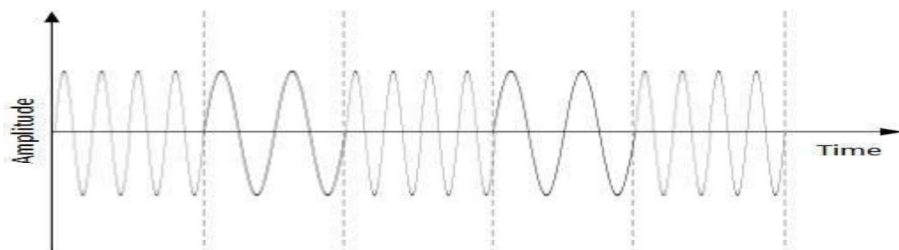
In this conversion technique, the amplitude of analog carrier signal is modified to reflect binary data.



When binary data represents digit 1, the amplitude is held; otherwise it is set to 0. Both frequency and phase remain same as in the original carrier signal.

- **Frequency Shift Keying**

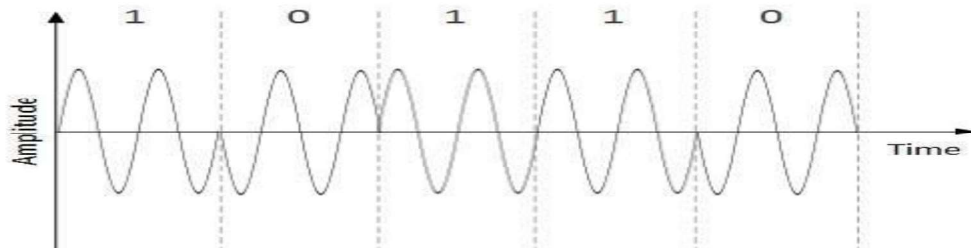
In this conversion technique, the frequency of the analog carrier signal is modified to reflect binary data.



This technique uses two frequencies, f_1 and f_2 . One of them, for example f_1 , is chosen to represent binary digit 1 and the other one is used to represent binary digit 0. Both amplitude and phase of the carrier wave are kept intact.

- **Phase Shift Keying**

In this conversion scheme, the phase of the original carrier signal is altered to reflect the binary data.



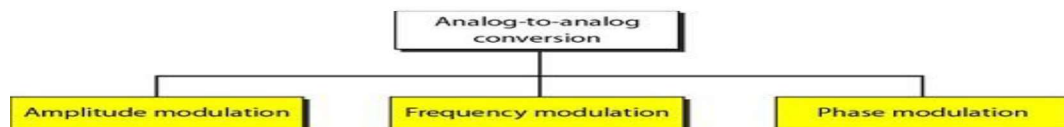
When a new binary symbol is encountered, the phase of the signal is altered. Amplitude and frequency of the original carrier signal is kept intact.

- **Quadrature Phase Shift Keying**

QPSK alters the phase to reflect two binary digits at once. This is done in two different phases. The main stream of binary data is divided equally into two sub-streams. The serial data is converted into parallel in both sub-streams and then each stream is converted to digital signal using NRZ technique. Later, both the digital signals are merged together.

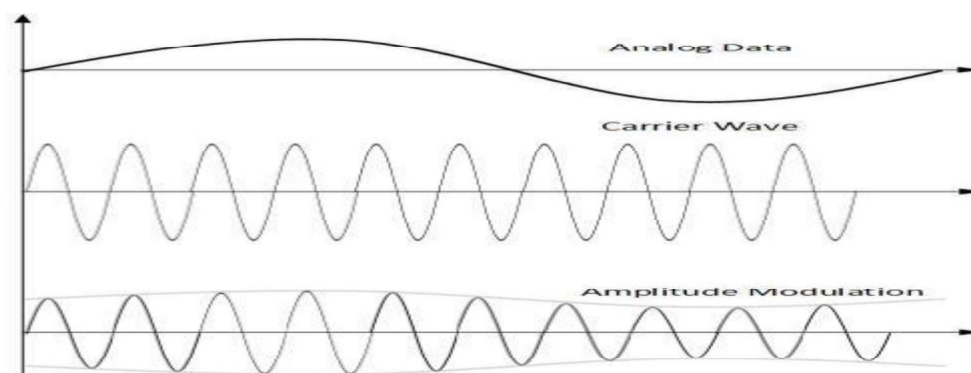
Analog-to-Analog Conversion

Analog signals are modified to represent analog data. This conversion is also known as Analog Modulation. Analog modulation is required when bandpass is used. Analog to analog conversion can be done in three ways:



- **Amplitude Modulation**

In this modulation, the amplitude of the carrier signal is modified to reflect the analog data.

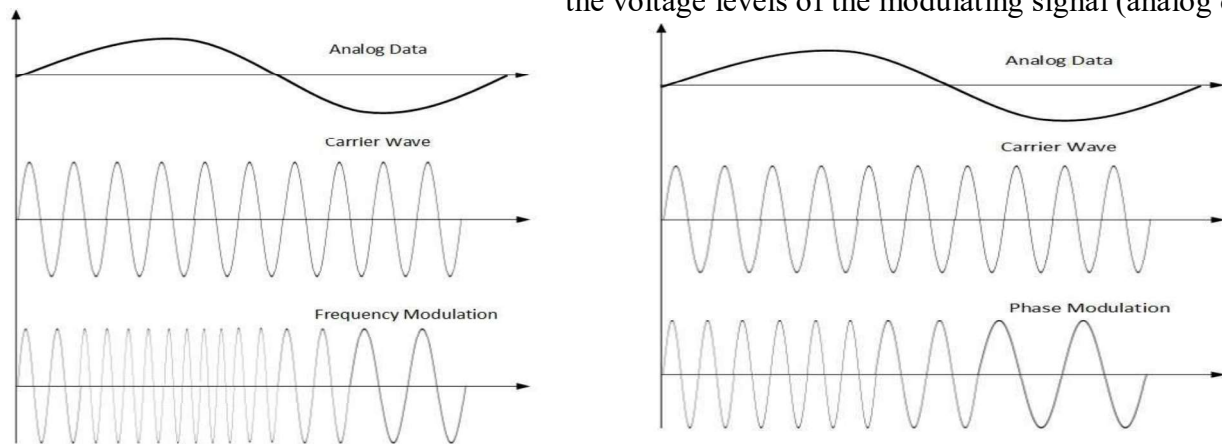


Amplitude modulation is implemented by means of a multiplier. The amplitude of modulating signal (analog data) is multiplied by the amplitude of carrier frequency, which then reflects analog data.

The frequency and phase of carrier signal remain unchanged.

- **Frequency Modulation**

In this modulation technique, the frequency of the carrier signal is modified to reflect the change in the voltage levels of the modulating signal (analog data).



The amplitude and phase of the carrier signal are not altered.

- **Phase Modulation**

In the modulation technique, the phase of carrier signal is modulated in order to reflect the change in voltage (amplitude) of analog data signal.

Phase modulation is practically similar to Frequency Modulation, but in Phase modulation frequency of the carrier signal is not increased. Frequency of carrier is signal is changed (made dense and sparse) to reflect voltage change in the amplitude of modulating signal

➤ **Multiplexing**

Multiplexing is a technique used to combine and send the multiple data streams over a single medium. The process of combining the data streams is known as multiplexing and hardware used for multiplexing is known as a multiplexer.

Multiplexing is achieved by using a device called Multiplexer (**MUX**) that combines n input lines to generate a single output line. Multiplexing follows many-to-one, i.e., n input lines and one output line.

Demultiplexing is achieved by using a device called Demultiplexer (**DEMUX**) available at the receiving end. DEMUX separates a signal into its component signals (one input and n outputs). Therefore, we can say that demultiplexing follows the one-to-many approach.

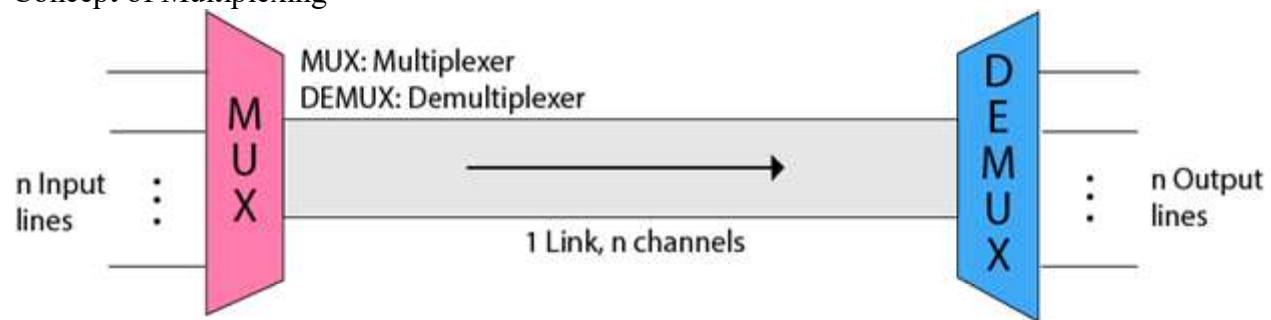
Why Multiplexing?

- The transmission medium is used to send the signal from sender to receiver. The medium can only have one signal at a time.
- If there are multiple signals to share one medium, then the medium must be divided in such a way that each signal is given some portion of the available bandwidth. For example: If there are 10 signals and bandwidth of medium is 100 units, then the 10 unit is shared by each signal.
- When multiple signals share the common medium, there is a possibility of collision. Multiplexing concept is used to avoid such collision.
- Transmission services are very expensive.

History of Multiplexing

- Multiplexing technique is widely used in telecommunications in which several telephone calls are carried through a single wire.
- Multiplexing originated in telegraphy in the early 1870s and is now widely used in communication.
- George Owen Squier developed the **telephone carrier multiplexing** in 1910.

Concept of Multiplexing



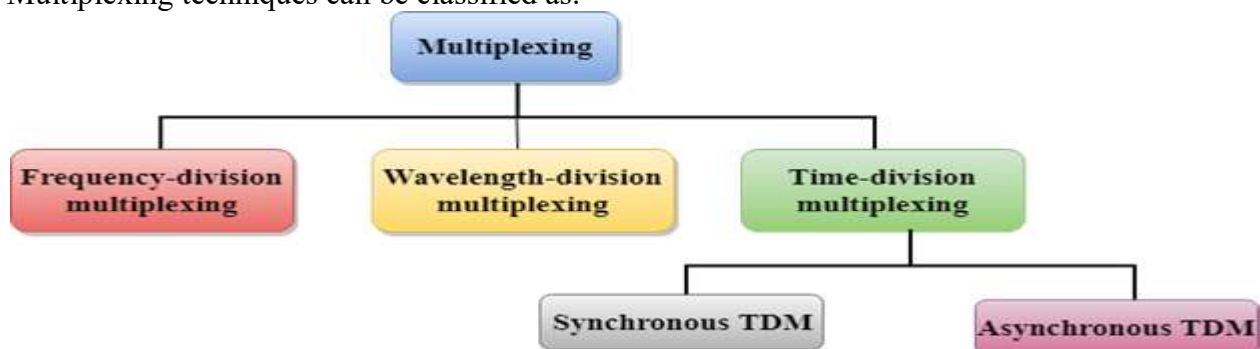
- The 'n' input lines are transmitted through a multiplexer and multiplexer combines the signals to form a composite signal.
- The composite signal is passed through a Demultiplexer and demultiplexer separates a signal to component signals and transfers them to their respective destinations.

Advantages of Multiplexing:

- More than one signal can be sent over a single medium.
- The bandwidth of a medium can be utilized effectively.

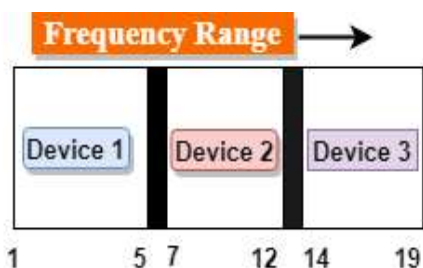
Multiplexing Techniques

Multiplexing techniques can be classified as:

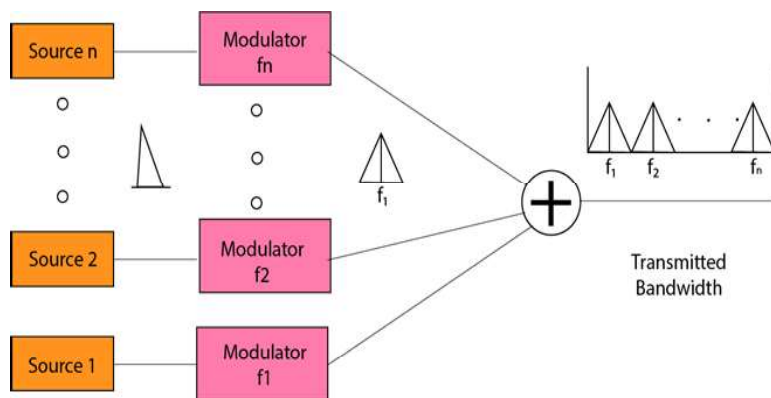


Frequency-division Multiplexing (FDM)

- It is an analog technique.
- **Frequency Division Multiplexing** is a technique in which the available bandwidth of a single transmission medium is subdivided into several channels.



- In the above diagram, a single transmission medium is subdivided into several frequency channels, and each frequency channel is given to different devices. Device 1 has a frequency channel of range from 1 to 5.
- The input signals are translated into frequency bands by using modulation techniques, and they are combined by a multiplexer to form a composite signal.
- The main aim of the FDM is to subdivide the available bandwidth into different frequency channels and allocate them to different devices.
- Using the modulation technique, the input signals are transmitted into frequency bands and then combined to form a composite signal.
- The carriers which are used for modulating the signals are known as **sub-carriers**. They are represented as $f_1, f_2, \dots, f_n$.



○ **FDM** is mainly used in radio broadcasts and TV networks.

Advantages Of FDM:

- FDM is used for analog signals.
- FDM process is very simple and easy modulation.
- A Large number of signals can be sent through an FDM simultaneously.
- It does not require any synchronization between sender and receiver.

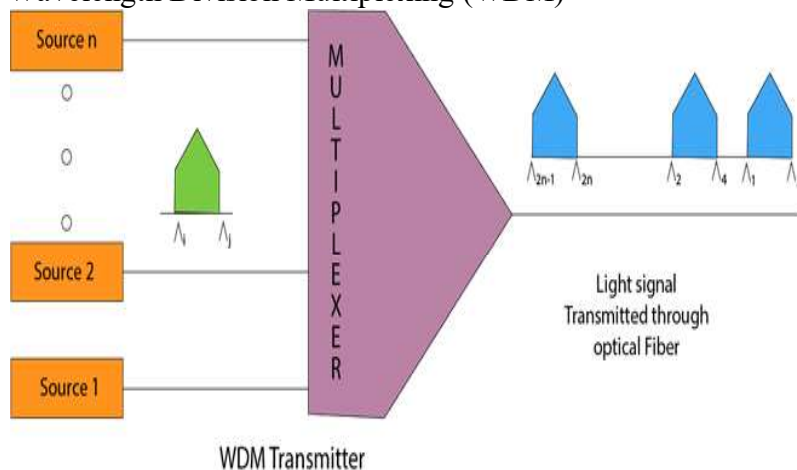
Disadvantages Of FDM:

- FDM technique is used only when low-speed channels are required.
- It suffers the problem of crosstalk.
- A Large number of modulators are required.
- It requires a high bandwidth channel.

Applications Of FDM:

- FDM is commonly used in TV networks.
- It is used in FM and AM broadcasting. Each FM radio station has different frequencies, and they are multiplexed to form a composite signal. The multiplexed signal is transmitted in the air.

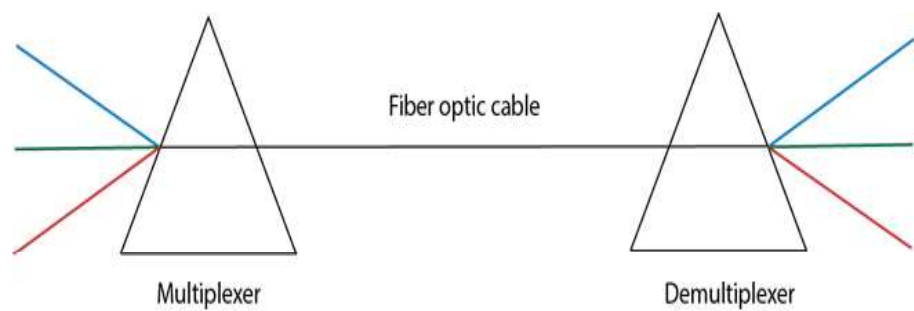
Wavelength Division Multiplexing (WDM)



- Wavelength Division Multiplexing is same as FDM except that the optical signals are transmitted through the fibre optic cable.
- WDM is used on fibre optics to increase the capacity of a single fibre.
- It is used to utilize the high data rate capability of fibre optic cable.
- It is an analog multiplexing technique.
- Optical signals from different source are combined to form a wider band of light with the help of multiplexer.

- At the receiving end, demultiplexer separates the signals to transmit them to their respective destinations.
- Multiplexing and Demultiplexing can be achieved by using a prism.
- Prism can perform a role of multiplexer by combining the various optical signals to form a composite signal, and the composite signal is transmitted through a fibre optical cable.
- Prism also performs a reverse operation, i.e., demultiplexing the signal.

Time Division Multiplexing



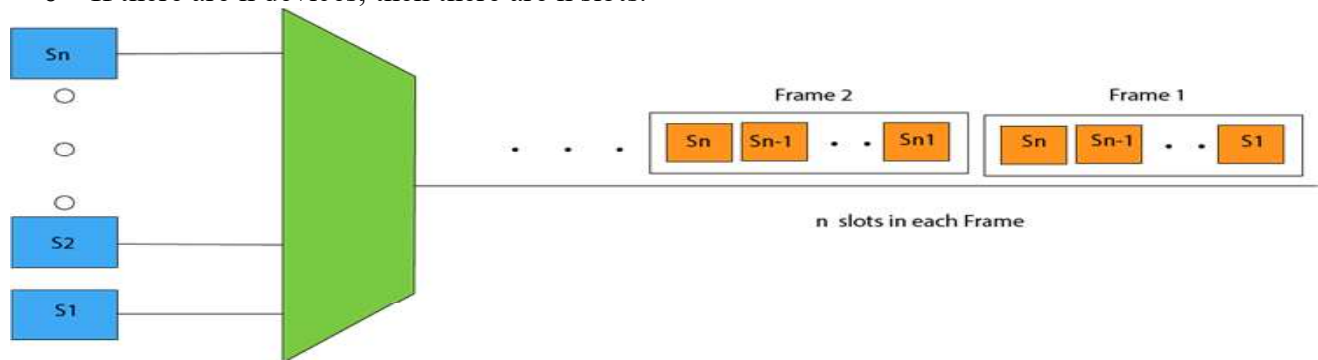
- It is a digital technique.
- In Frequency Division Multiplexing Technique, all signals operate at the same time with different frequency, but in case of Time Division Multiplexing technique, all signals operate at the same frequency with different time.
- In **Time Division Multiplexing technique**, the total time available in the channel is distributed among different users. Therefore, each user is allocated with different time interval known as a Time slot at which data is to be transmitted by the sender.
- A user takes control of the channel for a fixed amount of time.
- In Time Division Multiplexing technique, data is not transmitted simultaneously rather the data is transmitted one-by-one.
- In TDM, the signal is transmitted in the form of frames. Frames contain a cycle of time slots in which each frame contains one or more slots dedicated to each user.
- It can be used to multiplex both digital and analog signals but mainly used to multiplex digital signals.

There are two types of TDM:

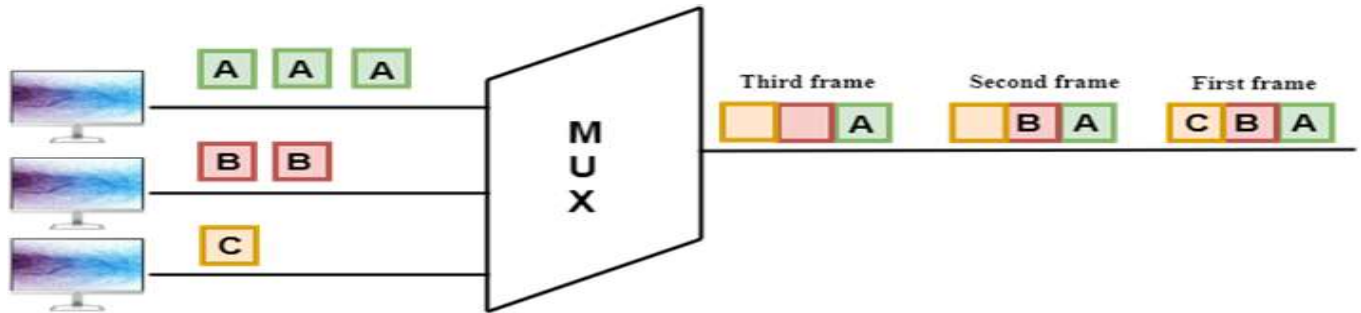
- Synchronous TDM
- Asynchronous TDM

Synchronous TDM

- A Synchronous TDM is a technique in which time slot is preassigned to every device.
- In Synchronous TDM, each device is given some time slot irrespective of the fact that the device contains the data or not.
- If the device does not have any data, then the slot will remain empty.
- In Synchronous TDM, signals are sent in the form of frames. Time slots are organized in the form of frames. If a device does not have data for a particular time slot, then the empty slot will be transmitted.
- The most popular Synchronous TDM are T-1 multiplexing, ISDN multiplexing, and SONET multiplexing.
- If there are n devices, then there are n slots.



Concept Of Synchronous TDM



In the above figure, the Synchronous TDM technique is implemented. Each device is allocated with some time slot. The time slots are transmitted irrespective of whether the sender has data to send or not.

Disadvantages Of Synchronous TDM:

- The capacity of the channel is not fully utilized as the empty slots are also transmitted which is having no data. In the above figure, the first frame is completely filled, but in the last two frames, some slots are empty. Therefore, we can say that the capacity of the channel is not utilized efficiently.
- The speed of the transmission medium should be greater than the total speed of the input lines. An alternative approach to the Synchronous TDM is Asynchronous Time Division Multiplexing.

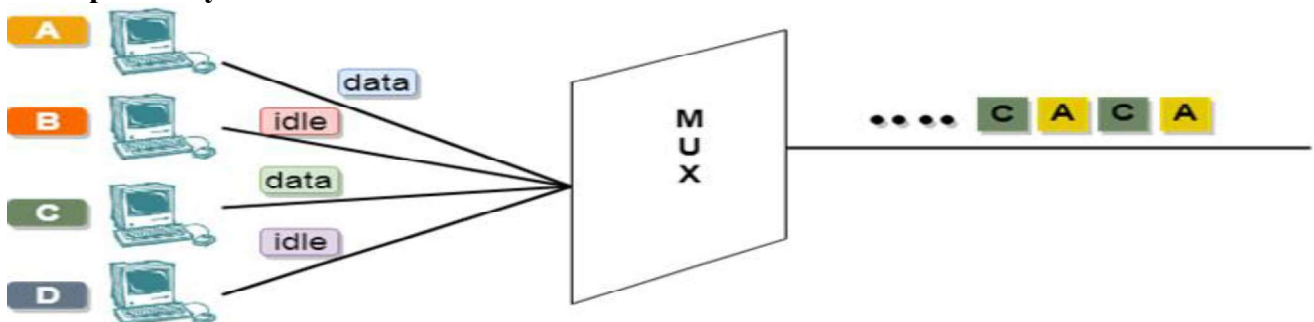
Asynchronous TDM

- An asynchronous TDM is also known as Statistical TDM.
- An asynchronous TDM is a technique in which time slots are not fixed as in the case of Synchronous TDM. Time slots are allocated to only those devices which have the data to send. Therefore, we can say that Asynchronous Time Division multiplexor transmits only the data from active workstations.
- An asynchronous TDM technique dynamically allocates the time slots to the devices.
- In Asynchronous TDM, total speed of the input lines can be greater than the capacity of the channel.
- Asynchronous Time Division multiplexor accepts the incoming data streams and creates a frame that contains only data with no empty slots.
- In Asynchronous TDM, each slot contains an address part that identifies the source of the data.



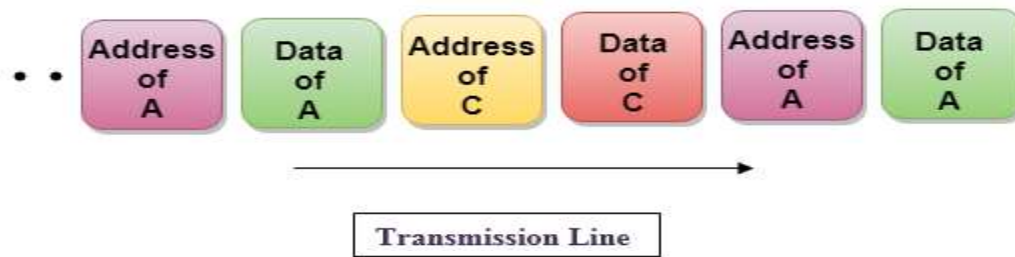
- The difference between Asynchronous TDM and Synchronous TDM is that many slots in Synchronous TDM are unutilized, but in Asynchronous TDM, slots are fully utilized. This leads to the smaller transmission time and efficient utilization of the capacity of the channel.
- In Synchronous TDM, if there are n sending devices, then there are n time slots. In Asynchronous TDM, if there are n sending devices, then there are m time slots where m is less than n ($m < n$).
- The number of slots in a frame depends on the statistical analysis of the number of input lines.

Concept Of Asynchronous TDM



In the above diagram, there are 4 devices, but only two devices are sending the data, i.e., A and C. Therefore, the data of A and C are only transmitted through the transmission line.

Frame of above diagram can be represented as:



The above figure shows that the data part contains the address to determine the source of the data.

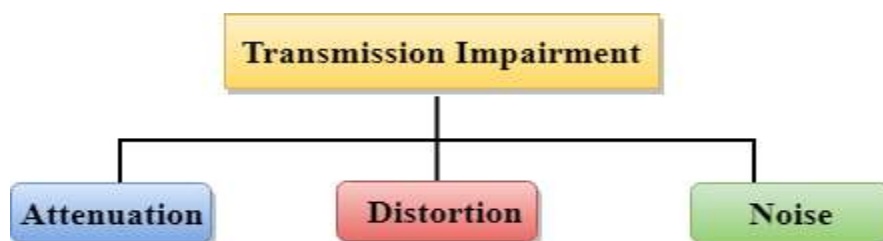
➤ Transmission Media

- Transmission media is a communication channel that carries the information from the sender to the receiver. Data is transmitted through the electromagnetic signals.
- The main functionality of the transmission media is to carry the information in the form of bits through LAN(Local Area Network).
- It is a physical path between transmitter and receiver in data communication.
- In a copper-based network, the bits in the form of electrical signals.
- In a fibre based network, the bits in the form of light pulses.
- In OSI(Open System Interconnection) phase, transmission media supports the Layer 1. Therefore, it is considered to be as a Layer 1 component.
- The electrical signals can be sent through the copper wire, fibre optics, atmosphere, water, and vacuum.
- The characteristics and quality of data transmission are determined by the characteristics of medium and signal.
- Transmission media is of two types are wired media and wireless media. In wired media, medium characteristics are more important whereas, in wireless media, signal characteristics are more important.
- Different transmission media have different properties such as bandwidth, delay, cost and ease of installation and maintenance.
- The transmission media is available in the lowest layer of the OSI reference model, i.e., **Physical layer**.

Some factors need to be considered for designing the transmission media:

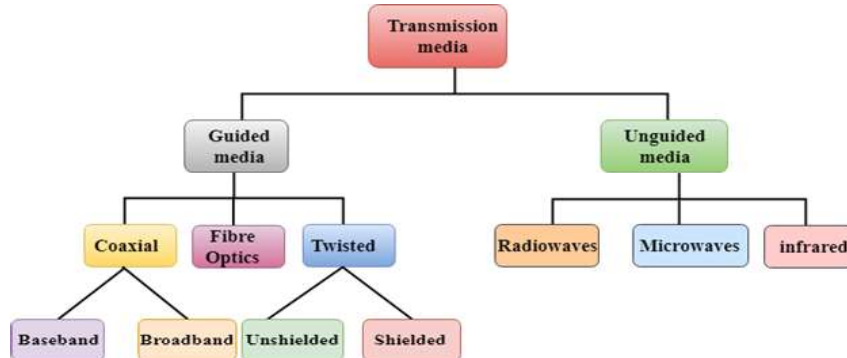
- **Bandwidth:** All the factors are remaining constant, the greater the bandwidth of a medium, the higher the data transmission rate of a signal.
- **Transmission impairment:** When the received signal is not identical to the transmitted one due to the transmission impairment. The quality of the signals will get destroyed due to transmission impairment.
- **Interference:** An interference is defined as the process of disrupting a signal when it travels over a communication medium on the addition of some unwanted signal.

Causes Of Transmission Impairment:



- **Attenuation:** Attenuation means the loss of energy, i.e., the strength of the signal decreases with increasing the distance which causes the loss of energy.
- **Distortion:** Distortion occurs when there is a change in the shape of the signal. This type of distortion is examined from different signals having different frequencies. Each frequency component has its own propagation speed, so they reach at a different time which leads to the delay distortion.
- **Noise:** When data is travelled over a transmission medium, some unwanted signal is added to it which creates the noise.

Classification of Transmission Media:



- Guided Transmission Media
- UnGuided Transmission Media

Guided Media

It is defined as the physical medium through which the signals are transmitted. It is also known as Bounded media.

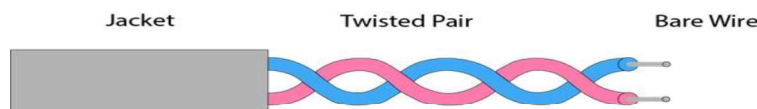
Types Of Guided media:

Twisted pair:

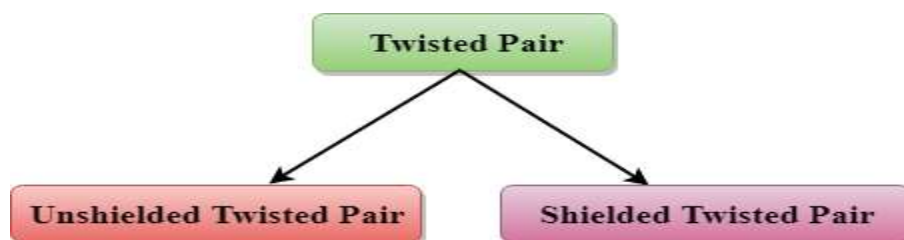
Twisted pair is a physical media made up of a pair of cables twisted with each other. A twisted pair cable is cheap as compared to other transmission media. Installation of the twisted pair cable is easy, and it is a lightweight cable. The frequency range for twisted pair cable is from 0 to 3.5KHz.

A twisted pair consists of two insulated copper wires arranged in a regular spiral pattern.

The degree of reduction in noise interference is determined by the number of turns per foot. Increasing the number of turns per foot decreases noise interference.



Types of Twisted pair:



Unshielded Twisted Pair:

An unshielded twisted pair is widely used in telecommunication. Following are the categories of the unshielded twisted pair cable:

- **Category 1:** Category 1 is used for telephone lines that have low-speed data.
- **Category 2:** It can support upto 4Mbps.
- **Category 3:** It can support upto 16Mbps.
- **Category 4:** It can support upto 20Mbps. Therefore, it can be used for long-distance communication.
- **Category 5:** It can support upto 200Mbps.

Advantages Of Unshielded Twisted Pair:

- It is cheap.
- Installation of the unshielded twisted pair is easy.
- It can be used for high-speed LAN.

Disadvantage:

- This cable can only be used for shorter distances because of attenuation.

Shielded Twisted Pair

A shielded twisted pair is a cable that contains the mesh surrounding the wire that allows the higher transmission rate.

Characteristics Of Shielded Twisted Pair:

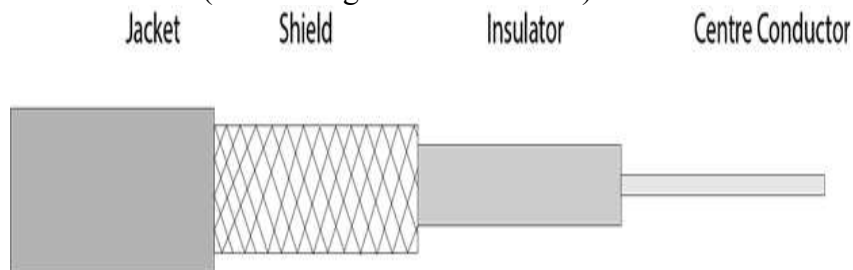
- The cost of the shielded twisted pair cable is not very high and not very low.
- An installation of STP is easy.
- It has higher capacity as compared to unshielded twisted pair cable.
- It has a higher attenuation.
- It is shielded that provides the higher data transmission rate.

Disadvantages

- It is more expensive as compared to UTP and coaxial cable.
- It has a higher attenuation rate.

Coaxial Cable

- Coaxial cable is very commonly used transmission media, for example, TV wire is usually a coaxial cable.
- The name of the cable is coaxial as it contains two conductors parallel to each other.
- It has a higher frequency as compared to Twisted pair cable.
- The inner conductor of the coaxial cable is made up of copper, and the outer conductor is made up of copper mesh. The middle core is made up of non-conductive cover that separates the inner conductor from the outer conductor.
- The middle core is responsible for the data transferring whereas the copper mesh prevents from the **EMI**(Electromagnetic interference).



Coaxial cable is of two types:

1. **Baseband transmission:** It is defined as the process of transmitting a single signal at high speed.
2. **Broadband transmission:** It is defined as the process of transmitting multiple signals simultaneously.

Advantages Of Coaxial cable:

- The data can be transmitted at high speed.
- It has better shielding as compared to twisted pair cable.
- It provides higher bandwidth.

Disadvantages Of Coaxial cable:

- It is more expensive as compared to twisted pair cable.
- If any fault occurs in the cable causes the failure in the entire network.

Fibre Optic

- Fibre optic cable is a cable that uses electrical signals for communication.
- Fibre optic is a cable that holds the optical fibres coated in plastic that are used to send the data by pulses of light.
- The plastic coating protects the optical fibres from heat, cold, electromagnetic interference from other types of wiring.
- Fibre optics provide faster data transmission than copper wires.

Diagrammatic representation of fibre optic cable:



Basic elements of Fibre optic cable:

- **Core:** The optical fibre consists of a narrow strand of glass or plastic known as a core. A core is a light transmission area of the fibre. The more the area of the core, the more light will be transmitted into the fibre.
- **Cladding:** The concentric layer of glass is known as cladding. The main functionality of the cladding is to provide the lower refractive index at the core interface as to cause the reflection within the core so that the light waves are transmitted through the fibre.
- **Jacket:** The protective coating consisting of plastic is known as a jacket. The main purpose of a jacket is to preserve the fibre strength, absorb shock and extra fibre protection.

Following are the advantages of fibre optic cable over copper:

- **Greater Bandwidth:** The fibre optic cable provides more bandwidth as compared copper. Therefore, the fibre optic carries more data as compared to copper cable.
- **Faster speed:** Fibre optic cable carries the data in the form of light. This allows the fibre optic cable to carry the signals at a higher speed.
- **Longer distances:** The fibre optic cable carries the data at a longer distance as compared to copper cable.
- **Better reliability:** The fibre optic cable is more reliable than the copper cable as it is immune to any temperature changes while it can cause obstruct in the connectivity of copper cable.
- **Thinner and Sturdier:** Fibre optic cable is thinner and lighter in weight so it can withstand more pull pressure than copper cable.

UnGuided Transmission

- An unguided transmission transmits the electromagnetic waves without using any

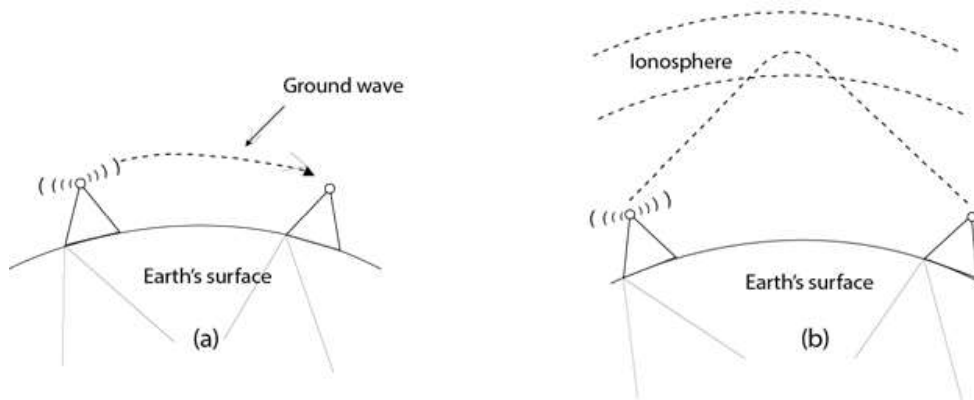
physical medium. Therefore it is also known as **wireless transmission**.

- In unguided media, air is the media through which the electromagnetic energy can flow easily.

Unguided transmission is broadly classified into three categories:

Radio waves

- Radio waves are the electromagnetic waves that are transmitted in all the directions of free space.
- Radio waves are unidirectional, i.e., the signals are propagated in all the directions.
- The range in frequencies of radio waves is from 3Khz to 1 khz.
- In the case of radio waves, the sending and receiving antenna are not aligned, i.e., the wave sent by the sending antenna can be received by any receiving antenna.
- An example of the radio wave is **FM radio**.

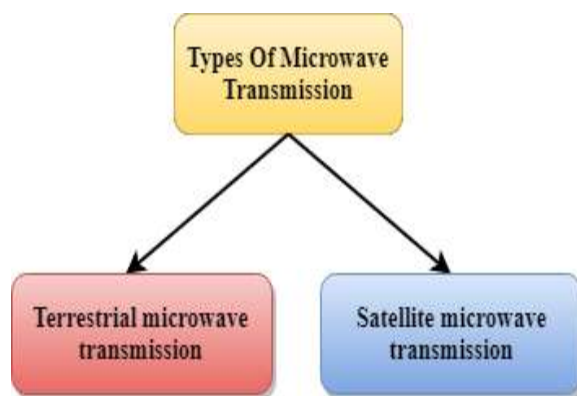


Applications Of Radio waves:

- A Radio wave is useful for multicasting when there is one sender and many receivers.
- An FM radio, television, cordless phones are examples of a radio wave.

Advantages Of Radio transmission:

- Radio transmission is mainly used for wide area networks and mobile cellular phones.
- Radio waves cover a large area, and they can penetrate the walls.
 - Radio transmission provides a higher transmission rate.



Microwaves are of two types:

- Terrestrial microwave
- Satellite microwave communication.

Terrestrial Microwave Transmission

○ Terrestrial Microwave transmission is a technology that transmits the focused beam of a radio signal from one ground-based microwave transmission antenna to another.

- Microwaves are the electromagnetic waves having the frequency in the range from 1GHz to 1000 GHz.
- Microwaves are unidirectional as the sending and receiving antenna is to be aligned, i.e., the waves sent by the sending antenna are narrowly focussed.
- In this case, antennas are mounted on the towers to send a beam to another antenna which

is km away.

- It works on the line of sight transmission, i.e., the antennas mounted on the towers are the direct sight of each other.

Characteristics of Microwave:

- **Frequency range:** The frequency range of terrestrial microwave is from 4-6 GHz to 21-23 GHz.
- **Bandwidth:** It supports the bandwidth from 1 to 10 Mbps.
- **Short distance:** It is inexpensive for short distance.
- **Long distance:** It is expensive as it requires a higher tower for a longer distance.
- **Attenuation:** Attenuation means loss of signal. It is affected by environmental conditions and antenna size.

Advantages Of Microwave:

- Microwave transmission is cheaper than using cables.
- It is free from land acquisition as it does not require any land for the installation of cables.
- Microwave transmission provides an easy communication in terrains as the installation of cable in terrain is quite a difficult task.
- Communication over oceans can be achieved by using microwave transmission.

Disadvantages of Microwave transmission:

- **Eavesdropping:** An eavesdropping creates insecure communication. Any malicious user can catch the signal in the air by using its own antenna.
- **Out of phase signal:** A signal can be moved out of phase by using microwave transmission.
- **Susceptible to weather condition:** A microwave transmission is susceptible to weather condition. This means that any environmental change such as rain, wind can distort the signal.
- **Bandwidth limited:** Allocation of bandwidth is limited in the case of microwave transmission.

Satellite Microwave Communication

- A satellite is a physical object that revolves around the earth at a known height.
- Satellite communication is more reliable nowadays as it offers more flexibility than cable and fibre optic systems.
- We can communicate with any point on the globe by using satellite communication.

How Does Satellite work?

The satellite accepts the signal that is transmitted from the earth station, and it amplifies the signal. The amplified signal is retransmitted to another earth station.

Advantages Of Satellite Microwave Communication:

- The coverage area of a satellite microwave is more than the terrestrial microwave.
- The transmission cost of the satellite is independent of the distance from the centre of the coverage area.
- Satellite communication is used in mobile and wireless communication applications.
- It is easy to install.
- It is used in a wide variety of applications such as weather forecasting, radio/TV signal broadcasting, mobile communication, etc.

Disadvantages Of Satellite Microwave Communication:

- Satellite designing and development requires more time and higher cost.
- The Satellite needs to be monitored and controlled on regular periods so that it remains in

orbit.

- The life of the satellite is about 12-15 years. Due to this reason, another launch of the satellite has to be planned before it becomes non-functional.

Infrared

- An infrared transmission is a wireless technology used for communication over short ranges.
- The frequency of the infrared is in the range from 300 GHz to 400 THz.
- It is used for short-range communication such as data transfer between two cell phones, TV remote operation, data transfer between a computer and cell phone resides in the same closed area.

Characteristics Of Infrared:

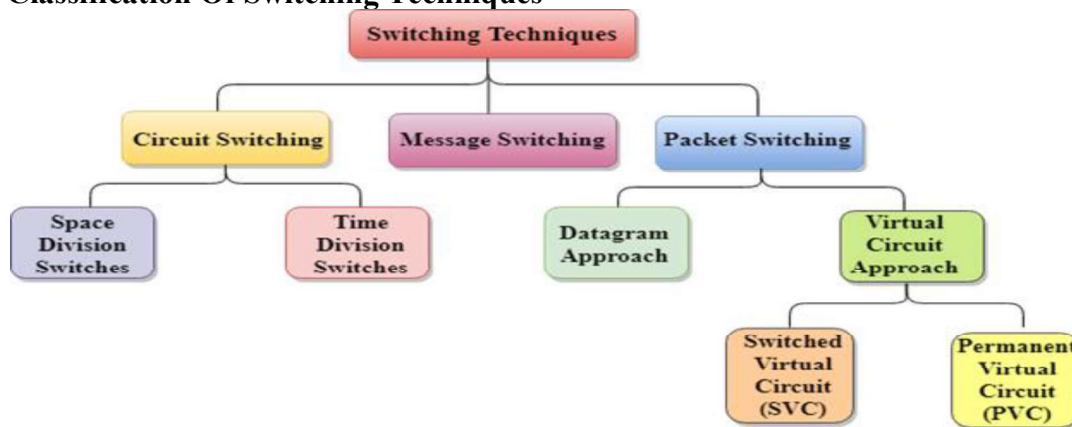
- It supports high bandwidth, and hence the data rate will be very high.
- Infrared waves cannot penetrate the walls. Therefore, the infrared communication in one room cannot be interrupted by the nearby rooms.
- An infrared communication provides better security with minimum interference.
- Infrared communication is unreliable outside the building because the sun rays will interfere with the infrared waves.

➤ Switching techniques

In large networks, there can be multiple paths from sender to receiver. The switching technique will decide the best route for data transmission.

Switching technique is used to connect the systems for making one-to-one communication.

Classification Of Switching Techniques



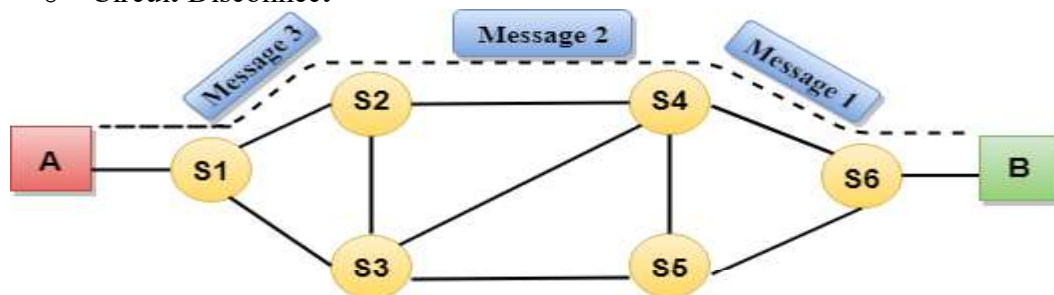
➤ Circuit Switching

- Circuit switching is a switching technique that establishes a dedicated path between sender and receiver.
- In the Circuit Switching Technique, once the connection is established then the dedicated path will remain to exist until the connection is terminated.
- Circuit switching in a network operates in a similar way as the telephone works.
- A complete end-to-end path must exist before the communication takes place.
- In case of circuit switching technique, when any user wants to send the data, voice, video, a request signal is sent to the receiver then the receiver sends back the acknowledgment to ensure the availability of the dedicated path. After receiving the acknowledgment, dedicated path transfers the data.

- Circuit switching is used in public telephone network. It is used for voice transmission.
- Fixed data can be transferred at a time in circuit switching technology.

Communication through circuit switching has 3 phases:

- Circuit establishment
- Data transfer
- Circuit Disconnect



Circuit Switching can use either of the two technologies:

Space Division Switches:

- Space Division Switching is a circuit switching technology in which a single transmission path is accomplished in a switch by using a physically separate set of crosspoints.
- Space Division Switching can be achieved by using crossbar switch. A crossbar switch is a metallic crosspoint or semiconductor gate that can be enabled or disabled by a control unit.
- The Crossbar switch is made by using the semiconductor. For example, Xilinx crossbar switch using FPGAs.
- Space Division Switching has high speed, high capacity, and nonblocking switches.

Space Division Switches can be categorized in two ways:

- **Crossbar Switch**
- **Multistage Switch**

Crossbar Switch

The Crossbar switch is a switch that has n input lines and n output lines. The crossbar switch has n^2 intersection points known as **crosspoints**.

Disadvantage of Crossbar switch:

The number of crosspoints increases as the number of stations is increased. Therefore, it becomes very expensive for a large switch. The solution to this is to use a multistage switch.

Multistage Switch

- Multistage Switch is made by splitting the crossbar switch into the smaller units and then interconnecting them.
- It reduces the number of crosspoints.
- If one path fails, then there will be an availability of another path.

Advantages Of Circuit Switching:

- In the case of Circuit Switching technique, the communication channel is dedicated.
- It has fixed bandwidth.

Disadvantages Of Circuit Switching:

- Once the dedicated path is established, the only delay occurs in the speed of data transmission.
- It takes a long time to establish a connection approx 10 seconds during which no data can be transmitted.

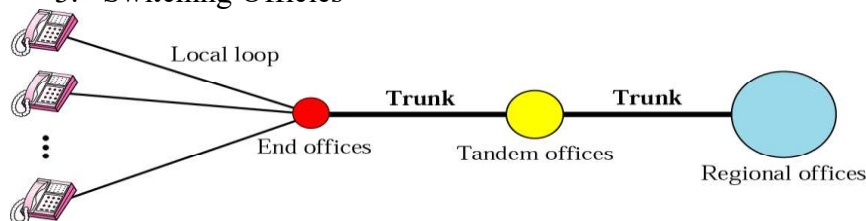
- It is more expensive than other switching techniques as a dedicated path is required for each connection.
- It is inefficient to use because once the path is established and no data is transferred, then the capacity of the path is wasted.
- In this case, the connection is dedicated therefore no other data can be transferred even if the channel is free.

➤ Telephone Network

Telephone networks use circuit switching, it had beginnings in the late 1980s. The entire network, which is referred to as the plain old telephone system (POTS), was originally an analog system using analog signals to transmit voice.

Major Components

1. Local Loops
2. Trunks
3. Switching Offices



Local Loops: Local Loop is a twisted-pair cable that connects the subscriber telephone to the nearest end office or local central office.

The local loop, when used for voice, has a bandwidth of 4000 Hz (4 KHz).

The first three digits of a local telephone number define the office, and the next four digits define the local loop number.

Trunks

Trunks are transmission media that handle the communication between offices. A trunk normally handles hundreds or thousands of connections through multiplexing. Transmission is usually through optical fibers or satellite links.

Switching Offices

To avoid having a permanent physical link between any two subscribers, the telephone company has switches located in a switching office. A switch connects several local loops or trunks and allows a connection between different subscribers.

LATAs

- After the divestiture of 1984 the United States was divided into more than 200 local-access transport areas (LATAs).
- A LATA can be a small or large metropolitan area.
- A small state may have one single LATA; a large state may have several LATAs.
- A LATA boundary may overlap the boundary of a state; part of a LATA can be in one state, part in another state.

IntraLATA Calls

IntraLATA refers to a telephone call or circuit which does not cross a LATA boundary. IntraLATA communications require the assistance of the Local Exchange Carrier, but not the IXC (InterExchange Carrier). Intra-LATA services are provided by local exchange carriers.

- ILEC (Incumbent Local Exchange Carriers)

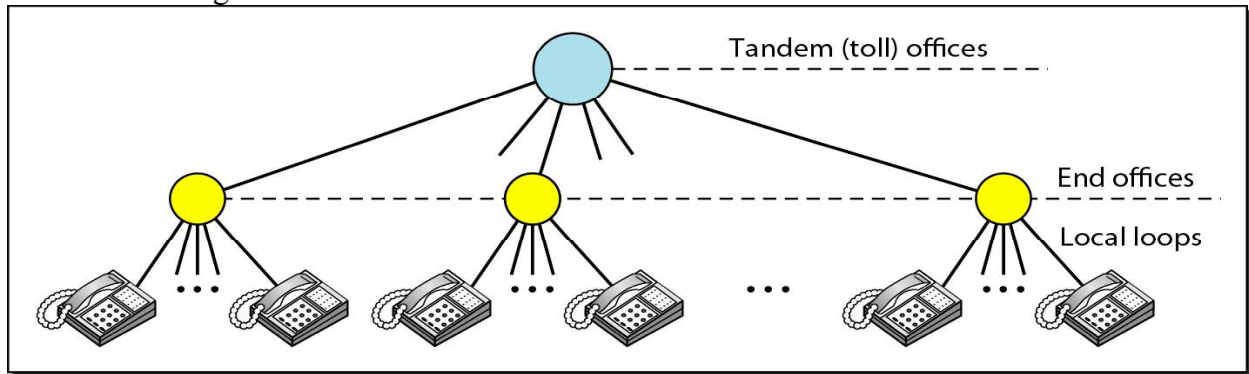
- CLEC (Competitive Local Exchange Carriers)

INTERLATA (INTERLATA AND INTERSTATE) CALLS

InterLATA calls refer to those that originate within one LATA and terminate in another. It is this type of call that is most referred to as a “long distance” call. Because InterLATA communications do cross LATA boundaries, they do require the help of an IXC (InterExchange Carrier) to be completed.

- Inter-LATA services are handled by IXCs (Inter-Exchange Carriers)
 - Called long-distance companies

Switching Offices in a LATA

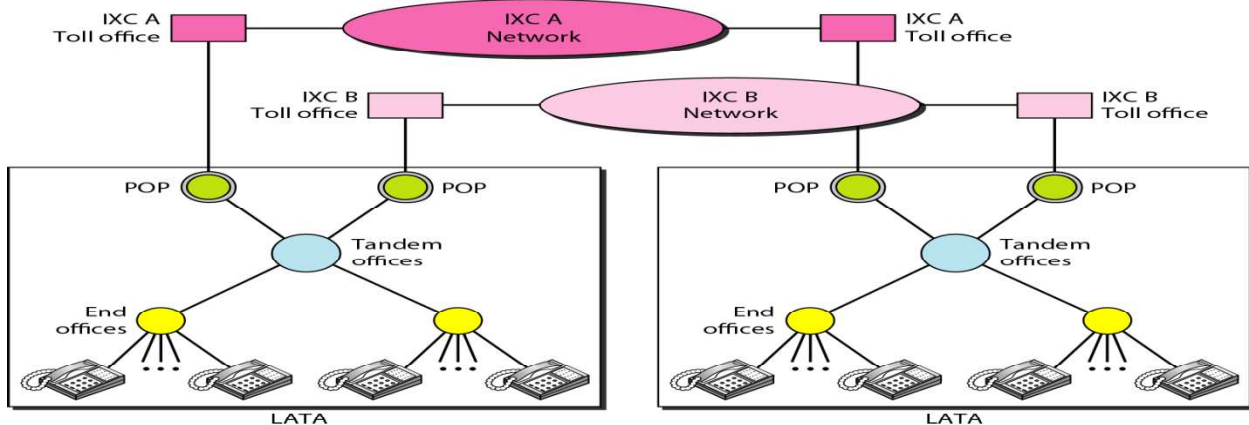


Points of Presence

Point of Presence is a switching office

Each IXC that wants to provide InterLATA services in a LATA must have a POP in that LATA.

The LECs that provide services inside the LATA must provide connections so that every subscriber can have access to all POPs.

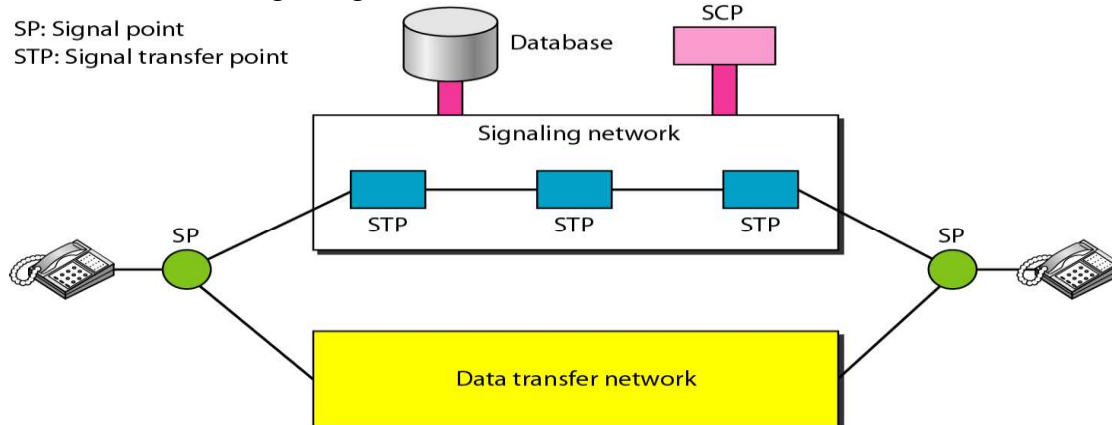


Signaling

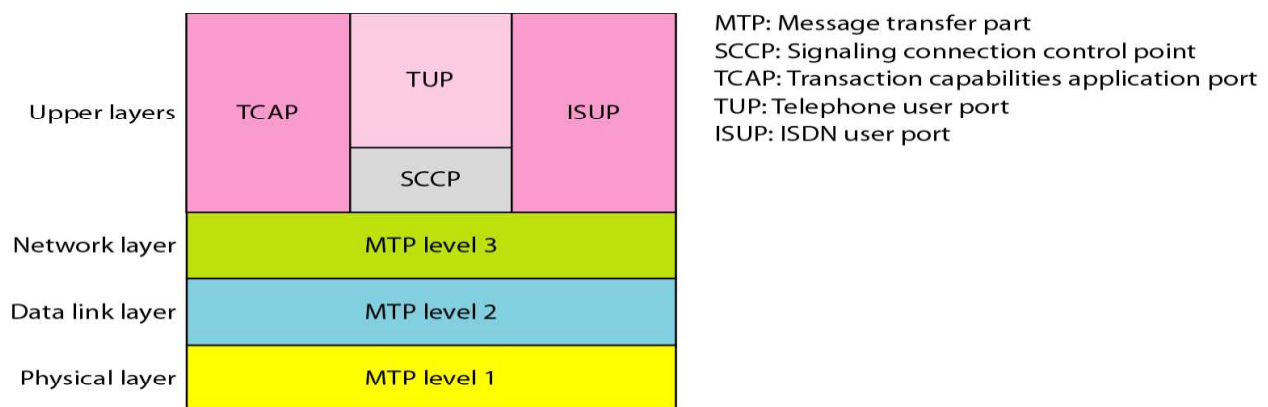
- In-band signaling : the same circuit used for both signaling and voice communication
- Out-of-band signaling
- The task of data transfer and signaling are separated in modern telephone network data transfer is done by one network, signaling by another
- Signaling system are required to
 - Providing dial tone, ring tone, and busy tone
 - Transferring telephone numbers between offices
 - Maintaining and monitoring the call
 - Keeping billing information
 - Maintaining and monitoring the status of the telephone network equipment

- Providing other functions such as caller ID, voice mail, and so on

Data Transfer and Signaling Networks



Layers in SS7 (Signaling System 7)



Physical Layer: MTPLevel1 The physical layerSS7 called message transport part (MTP) level1 uses several physical layer specifications such as T-1(1.544 Mbps) and DC0(64kbps).

DataLinkLayer: MTPLevel2The MTPLevel12 layer provides typical data link layer services such as packetizing, using source and destination address in the packet header, and CRC for error checking.

Network Layer: MTPLevel3The MTPLevel3layer provides end-to-end connectivity by using the datagram approach to switching .Routers and switches route the signal packets from the source to the destination.

Telephone Network Services

Transport Layer:SCCP The signaling connection control unit (SCCP) is used for special services such as 800-call processing.

Upper Layer: TUP, TCAP, and ISUP There are three protocols at the upper layers. Telephone user port (TUP) is responsible for setting up voice calls. It receives the dialed digits and route the calls. Transaction capabilities application port(TCAP) provides remote calls that let an application program on a computer invoke a procedure on another computer.ISDN user port(ISUP)can replace TUP to provide services similar to those of an ISDN network

- Analog services
 - Analog switched services
 - Dial-up service
 - 800 service, 900 service
 - WATS (wide-area telephone service)
 - Analog leased service: called a dedicated line

- Digital service
 - Switched/56 service
 - Digital data service: digital leased line

➤ DSL (DIGITAL SUBSCRIBER LINE)

DSL Provided higher-speed access to the Internet,DSL is one of the most promising for supporting high-speed digital communication over the existing local loops.DSL technology is a set of technologies, each differing in the first letter(ADSL,VDSL,HDSL and SDSL).The set is often referred to as xDSL,where x can be replaced by A,V,H or S.

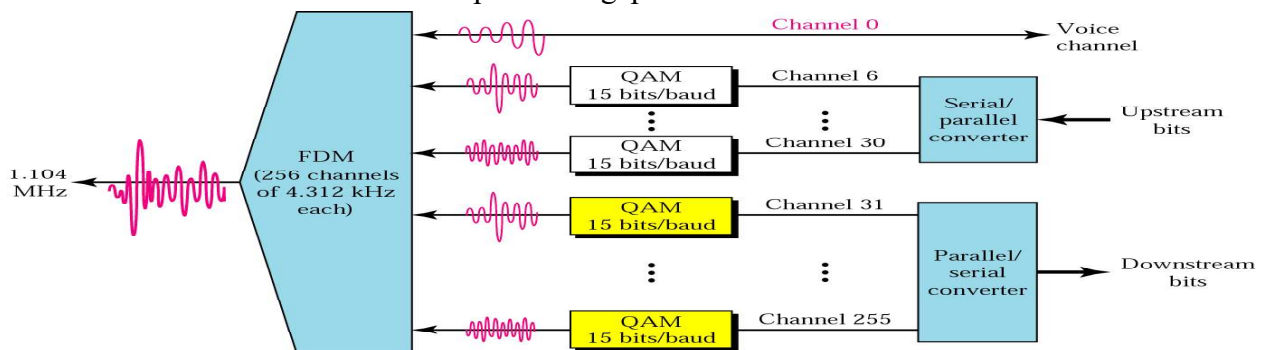
ADSL

- Asymmetrical Digital Subscriber Line
- ADSL is an asymmetric communication technology designed for residential users; it is not suitable for businesses
- The existing local loop can handle bandwidths up to 1.1 MHz
- ADSL is an adaptive technology. The system uses a data rate based on the condition of the local loop line.
- Discrete Multitone Technique:

The modulation technique that has become standard for ADSL is called the discrete multitone technique (DMT)which combines QAM and FDM.Typically an unavailable bandwidth of 1.04 MHz is divided into 256 channels. Each Channel use a bandwidth of 4.312KHZ.bandwidth can be divided into the following:

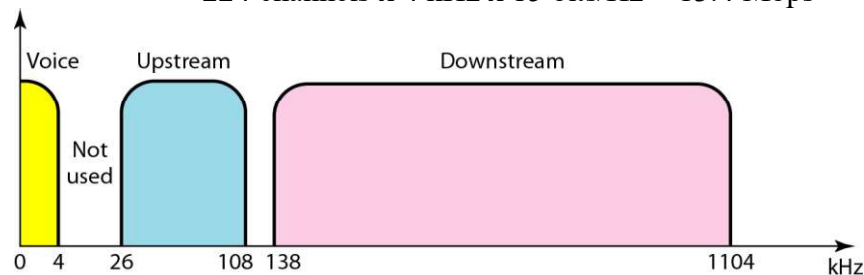
Voice Channel 0 is reserved for voice communication.

Idle Channel 1 to 5 are not used and provide a gap between voice and data communication



Bandwidth Division in ADSL

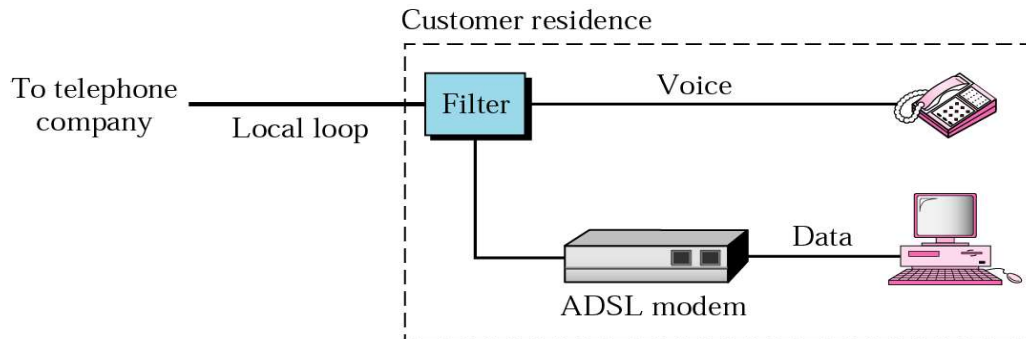
- There is no set way that the bandwidth is divided
- Upstream
 - $24 \text{ channels} \times 4 \text{ kHz} \times 15 \text{ bits/Hz} = 1.44 \text{ Mbps}$
- Downstream
 - $224 \text{ channels} \times 4 \text{ kHz} \times 15 \text{ bits/Hz} = 13.4 \text{ Mbps}$



DSL: Actual Bit Rate

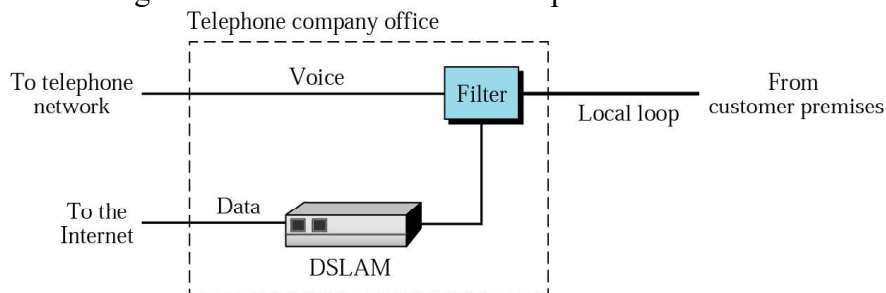
- Because of the high signal/noise ratio
- Upstream
 - Normally below 500 kbps
- Downstream
 - Normally below 8 Mbps

Customer Site: ADSL Modem



Telephone Company Site: DSLAM

- Digital subscriber line access multiplexer



Other DSL Technologies: xDSL

- HDSL (High-bit-rate DSL)
 - Designed as an alternative to the T-1 line (AMI encoding)
 - 2B1Q encoding used for less susceptible to attenuation
 - Up to a distance 12,000 ft without repeaters
 - Two twisted pairs for full-duplex transmission
- SDSL (Symmetric DSL)
 - One twisted pair version of HDSL
 - 768 kbps in each direction, symmetric communication
 - Send and receive data in large volumes in both directions
- VDSL (Very-high-bit-rate DSL): Alternative approach to ADSL
 - Coaxial, fiber-optic, or twisted-pair cable for short distances
 - Uses DMT with 25-55 Mbps (downstream) and 3.2 Mbps (upstream)

<i>Technology</i>	<i>Downstream Rate</i>	<i>Upstream Rate</i>	<i>Distance (ft)</i>	<i>Twisted Pairs</i>	<i>Line Code</i>
ADSL	1.5–6.1 Mbps	16–640 kbps	12,000	1	DMT
ADSL Lite	1.5 Mbps	500 kbps	18,000	1	DMT
HDSL	1.5–2.0 Mbps	1.5–2.0 Mbps	12,000	2	2B1Q
SDSL	768 kbps	768 kbps	12,000	1	2B1Q
VDSL	25–55 Mbps	3.2 Mbps	3000–10,000	1	DMT

➤ S
**ON
ET**
**Defi
niti
on**

Synchronous optical networking (SONET) is a standardized digital communication protocol that synchronously transfers multiple data streams over long distances through fiber optic cables. It is a physical layer specification that allows simultaneous transmission of voice, data, and video at speeds as high as 1Gbps through a single fiber. In telephone networks, it is used for transmission of a huge amount of telephone calls and data streams through fiber.

SONET was standardized by the American National Standards Institute (ANSI). It is equivalent to Synchronous Digital Hierarchy (SDH) standardized by the International Telecommunication Union (ITU).

Today SONET acts as a standard so that digital networks can interconnect and that existing conventional transmission systems can take advantage of optical media through tributary attachments. Backbone carrier networks will typically utilize SONET. Typically, telecommunication companies will share data over the line of a fiber optic cable instead of going through the more expensive process of digging trenches to bury new cables. Data is multiplexed by separating the cable into separate channels. The speed of data transmission is comparable to Gigabit Ethernet speeds.

The network elements defined in SONET include the STS multiplexer, STS demultiplexer, regenerator and the add/drop multiplexer. The STS multiplexer is the process that multiplexes signals and converts electrical signals to optical ones. STS demultiplexer condenses signals and converts optical signals back to electrical signals. Regenerators increase incoming optical signals, allowing them to travel farther. The add/drop multiplexer enables a signal to be added or removed from a source.

SONET connections are broken down between sections, lines and paths. A section is the part of a network which connects two devices. The line connects two multiplexers, and the path is the network end-to-end.

SONET also defines four different layers, the path, line, section and photonic layers. The path layer moves signals from its source to its destination. The line layer is where the signal moves across the cable. The section layer defines the movement of signals across cables. The photonic layer is the specification for optical fiber channels.

SONET standards

SONET provides standards for a number of line rates up to the maximum line rate of 10 gigabits per second (Gbps). Actual line rates approaching in the 30 gigabits per second range are possible.

Base units of SONET are defined as Optical Carrier level-1, or OC-1. OC-1 supports up to 51.84 megabytes per second (Mbps). The next level up, OC-3, supports up to triple the bandwidth. Each level increases by multiples of four. OC-3, OC-12, OC-24, OC-48 can be used as examples. The set of multiples of the base rate known as "Optical Carrier levels (OCx)."

SONET standards are specified in ANSI T1.105 and T1.117.

Benefits

SONET is seen to have multiple characteristics that are considered advantageous, such as:

- High data rates.
- Large transmit distances.
- Support of multiple data types (data, voice and video).

- Can carry high-level protocols such as IP.
- Defines interoperability standards for organizations.

The one main disadvantage of utilizing SONET, however, is that it is high in cost.

SONET vs. SDH

Synchronous digital hierarchy (SDH) is the international equivalent of SONET. SONET and SDH are very similar standards made for the same reason. However, the basic unit for SDH is called the synchronous transmission module level-1 (STM-1), as compared to SONET's Optical Carrier level.

SDH is an International Telecommunications Union (ITU) standard and can work with SONET line rates. However, SONET and SDH have differing structures when restructuring data. SDH frames are made out of 2,430 bytes and use Synchronous Transport Module (STM), while SONET frames are made up of 6,480 and use Synchronous Transport Signal (STS).

SONET Frames

The basic SONET frame comprises of a block of 810 bytes. The frames are transmitted at the rate of 8000 frames/bytes, which is the sampling rate of the telephone networks.

A SONET frame is represented as a rectangular block of bytes with 9 rows and 90 columns as shown in the following diagram.



The first three columns of the SONET frame contains the system information and is generally termed as system overhead, while the rest (marked in blue) contains the payload, i.e. the data to be transmitted. In this frame, the first three rows of the system overhead (marked in yellow) contain section overhead, and the next six rows (marked in orange) contain line overhead.

Modem

A modem is a small box that connects your devices to the Internet using cables. Unlike a router, a modem doesn't provide your home with Wi-Fi connectivity. A modem acts as a digital translator, taking an information signal from your cable, fiber or phone lines and making it accessible to your computer.

Types of Modems

There are three types of modems: cable, digital subscriber line (DSL) and dial-up. A cable modem uses coaxial cables that connect to the back of the modem and the bolt-like outlet in your wall or on your cable box. This type of modem delivers high speed internet to your device.

DSL and dial-up modems use a cable that connects to your phone line. DSL, however, still allows you to use your landline telephone while connected to the internet.

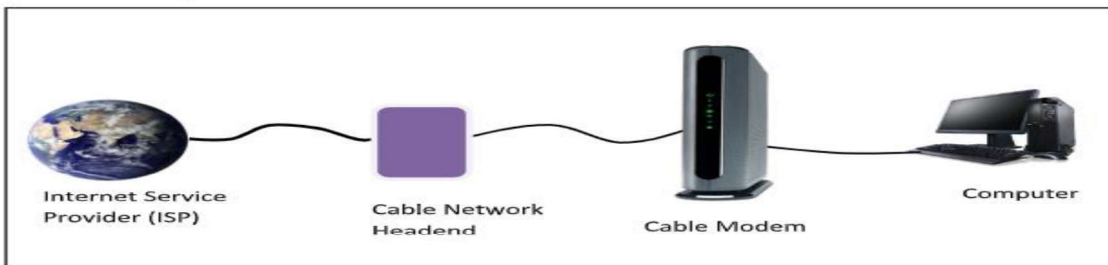
Fiber-optic technology doesn't require a modem for its Internet service.

➤ Cable Modem

Cable modem is a hardware device that is used to connect the computer with the Internet Service Provider (ISP) through the local cable TV line. It has two interfaces – one to the cable TV network outlet and the other to a computer or television or set-top box.

A **cable modem** is a type of network bridge that provides bi-directional data communication via radio frequency channels on a hybrid fiber-coaxial (HFC), radio frequency over glass (RFOG) and coaxial cable infrastructure. Cable modems are primarily used to deliver broadband Internet access in the form of cable Internet, taking advantage of the high bandwidth of a HFC and RFOG network

Configuration



Cable modems used to be proprietary in the initial days and had to be installed by the cable company. Nowadays, cable modems of open standards are available that can be personally installed by the user. The standard is called Data Over Cable Service Interface Spectrum (DOCSIS). The modem to computer interface is normally Ethernet or USB. The interface between the modem and the cable network outlet supports FDM, TDM, and CDMA so that the bandwidth of the cable can be shared among the subscribers.

Establishment of Connection

After a cable modem is plugged on to the cable TV network, it scans the downstream channels for a particular packet that is periodically sent over the network. On detecting it, the modem announces its presence over the network. If its authentication criteria are met, then it is assigned for both upstream and downstream communication.

Channels for Communication

For downstream data, 6MHz or 8MHz channels are used which are modulated using QAM-64. This gives the data rate of 36Mbps. For upstream data, there is more radio-frequency noise. Consequently, the data rate is around 9Mbps.

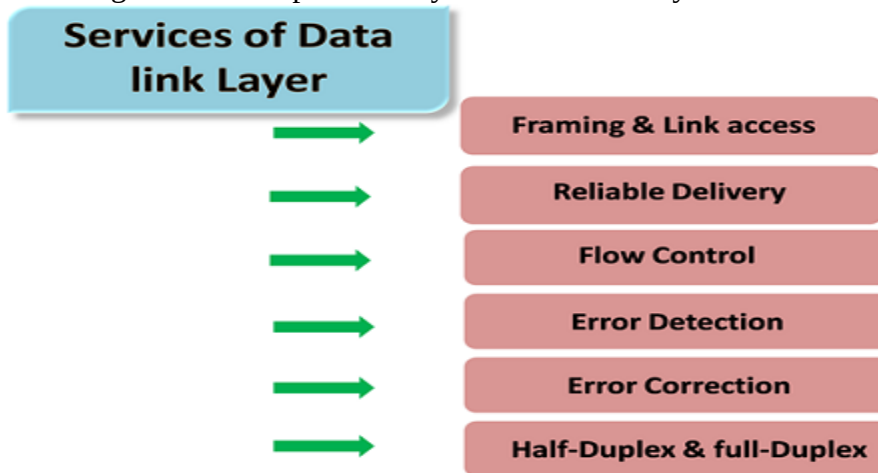
Communication Method

For sharing upstream data, time division multiplexing (TDM) is used. TDM divides the time in minislots, which are assigned to subscribers who want to send the data. When a computer has data to send, it sends data packets to the cable modem. The modem requests the number of minislots needed to send the data. If the request is granted, the modem receives an acknowledgment along with the allotted number of slots. The modem then transmits the data packets accordingly.

➤ Data Link Layer

- In the OSI model, the data link layer is a 4th layer from the top and 2nd layer from the bottom.
- The communication channel that connects the adjacent nodes is known as links, and in order to move the datagram from source to the destination, the datagram must be moved across an individual link.
- The main responsibility of the Data Link Layer is to transfer the datagram across an individual link.
- The Data link layer protocol defines the format of the packet exchanged across the nodes as well as the actions such as Error detection, retransmission, flow control, and random access.
- The Data Link Layer protocols are Ethernet, token ring, FDDI and PPP.
- An important characteristic of a Data Link Layer is that datagram can be handled by different link layer protocols on different links in a path. For example, the datagram is handled by Ethernet on the first link, PPP on the second link.

Following services are provided by the Data Link Layer:



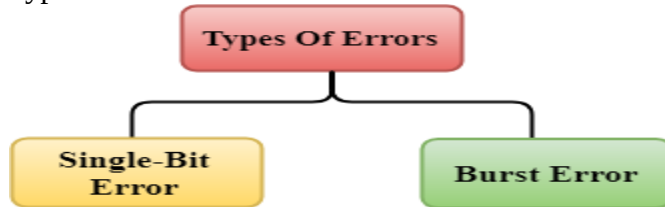
- **Framing & Link access:** Data Link Layer protocols encapsulate each network frame within a Link layer frame before the transmission across the link. A frame consists of a data field in which network layer datagram is inserted and a number of data fields. It specifies the structure of the frame as well as a channel access protocol by which frame is to be transmitted over the link.
- **Reliable delivery:** Data Link Layer provides a reliable delivery service, i.e., transmits the network layer datagram without any error. A reliable delivery service is accomplished with transmissions and acknowledgements. A data link layer mainly provides the reliable delivery service over the links as they have higher error rates and they can be corrected locally, link at which an error occurs rather than forcing to retransmit the data.
- **Flow control:** A receiving node can receive the frames at a faster rate than it can process the frame. Without flow control, the receiver's buffer can overflow, and frames can get lost. To overcome this problem, the data link layer uses the flow control to prevent the sending node on one side of the link from overwhelming the receiving node on another side of the link.
- **Error detection:** Errors can be introduced by signal attenuation and noise. Data Link Layer protocol provides a mechanism to detect one or more errors. This is achieved by adding error detection bits in the frame and then receiving node can perform an error check.
- **Error correction:** Error correction is similar to the Error detection, except that receiving node not only detect the errors but also determine where the errors have occurred in the frame.
- **Half-Duplex & Full-Duplex:** In a Full-Duplex mode, both the nodes can transmit the data at the same time. In a Half-Duplex mode, only one node can transmit the data at the same time.

➤ Error Detection and Correction:

Error Detection

When data is transmitted from one device to another device, the system does not guarantee whether the data received by the device is identical to the data transmitted by another device. An Error is a situation when the message received at the receiver end is not identical to the message transmitted.

Types Of Errors

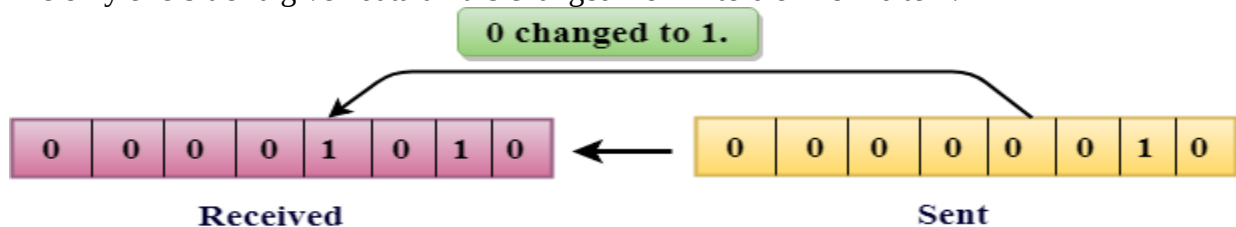


Errors can be classified into two categories:

- Single-Bit Error
- Burst Error

Single-Bit Error:

The only one bit of a given data unit is changed from 1 to 0 or from 0 to 1.



In the above figure, the message which is sent is corrupted as single-bit, i.e., 0 bit is changed to 1.

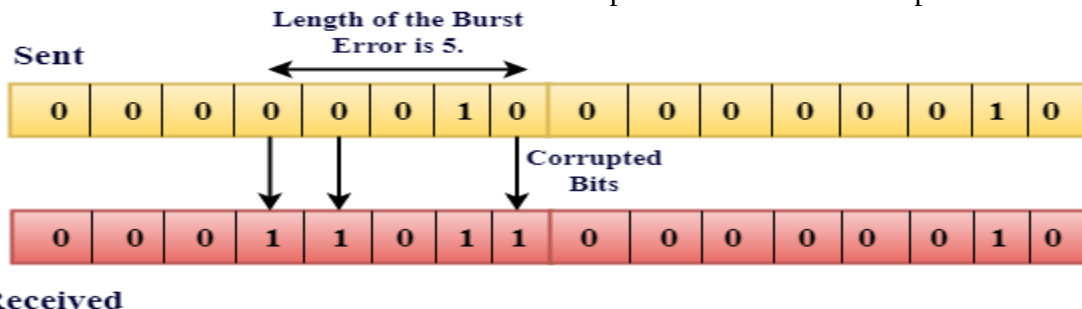
Single-Bit Error does not appear more likely in Serial Data Transmission. For example, Sender sends the data at 10 Mbps, this means that the bit lasts only for 1 ?s and for a single-bit error to occurred, a noise must be more than 1 ?s.

Single-Bit Error mainly occurs in Parallel Data Transmission. For example, if eight wires are used to send the eight bits of a byte, if one of the wire is noisy, then single-bit is corrupted per byte.

Burst Error:

The two or more bits are changed from 0 to 1 or from 1 to 0 is known as Burst Error.

The Burst Error is determined from the first corrupted bit to the last corrupted bit.



Received

The duration of noise in Burst Error is more than the duration of noise in Single-Bit.

Burst Errors are most likely to occur in Serial Data Transmission.

The number of affected bits depends on the duration of the noise and data rate.

Error Detecting Techniques:

The most popular Error Detecting Techniques are:

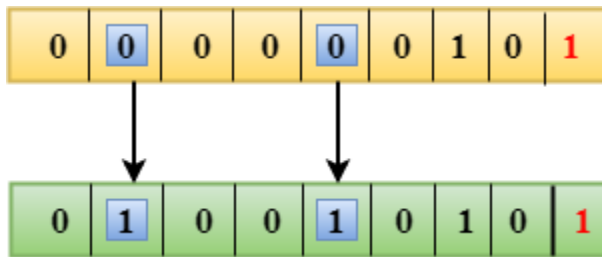
- **Single parity check**
- **Two-dimensional parity check**
- **Checksum**
- **Cyclic redundancy check**
- **Single Parity Check**
- Single Parity checking is the simple mechanism and inexpensive to detect the errors.

- In this technique, a redundant bit is also known as a parity bit which is appended at the end of the data unit so that the number of 1s becomes even. Therefore, the total number of transmitted bits would be 9 bits.
- If the number of 1s bits is odd, then parity bit 1 is appended and if the number of 1s bits is even, then parity bit 0 is appended at the end of the data unit.
- At the receiving end, the parity bit is calculated from the received data bits and compared with the received parity bit.
- This technique generates the total number of 1s even, so it is known as even-parity checking.



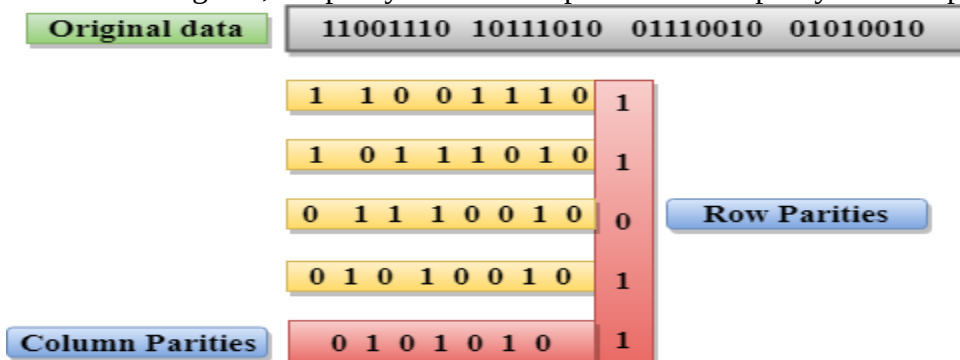
Drawbacks Of Single Parity Checking

- It can only detect single-bit errors which are very rare.
- If two bits are interchanged, then it cannot detect the errors.



Two-Dimensional Parity Check

- Performance can be improved by using **Two-Dimensional Parity Check** which organizes the data in the form of a table.
- Parity check bits are computed for each row, which is equivalent to the single-parity check.
- In Two-Dimensional Parity check, a block of bits is divided into rows, and the redundant row of bits is added to the whole block.
- At the receiving end, the parity bits are compared with the parity bits computed from the received data.



Drawbacks Of 2D Parity Check

- If two bits in one data unit are corrupted and two bits exactly the same position in another data unit are also corrupted, then 2D Parity checker will not be able to detect the error.
- This technique cannot be used to detect the 4-bit errors or more in some cases.

Checksum

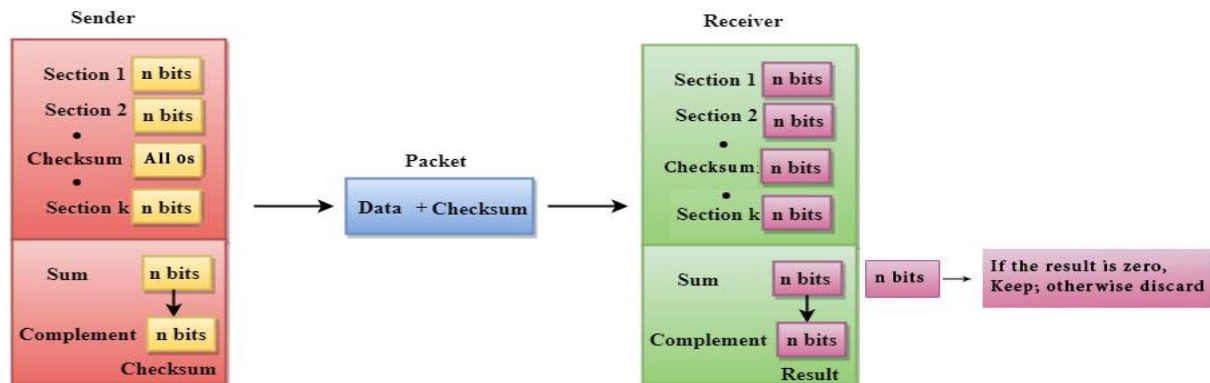
A Checksum is an error detection technique based on the concept of redundancy.

It is divided into two parts:

Checksum Generator

A Checksum is generated at the sending side. Checksum generator subdivides the data into equal segments of n bits each, and all these segments are added together by using one's complement arithmetic. The sum is complemented and appended to the original data, known as checksum field. The extended data is transmitted across the network.

Suppose L is the total sum of the data segments, then the checksum would be $?L$



1. The Sender follows the given steps:
2. The block unit is divided into k sections, and each of n bits.
3. All the k sections are added together by using one's complement to get the sum.
4. The sum is complemented and it becomes the checksum field.
5. The original data and checksum field are sent across the network.

Checksum Checker

A Checksum is verified at the receiving side. The receiver subdivides the incoming data into equal segments of n bits each, and all these segments are added together, and then this sum is complemented. If the complement of the sum is zero, then the data is accepted otherwise data is rejected.

1. The Receiver follows the given steps:
2. The block unit is divided into k sections and each of n bits.
3. All the k sections are added together by using one's complement algorithm to get the sum.
4. The sum is complemented.
5. If the result of the sum is zero, then the data is accepted otherwise the data is discarded.

Cyclic Redundancy Check (CRC)

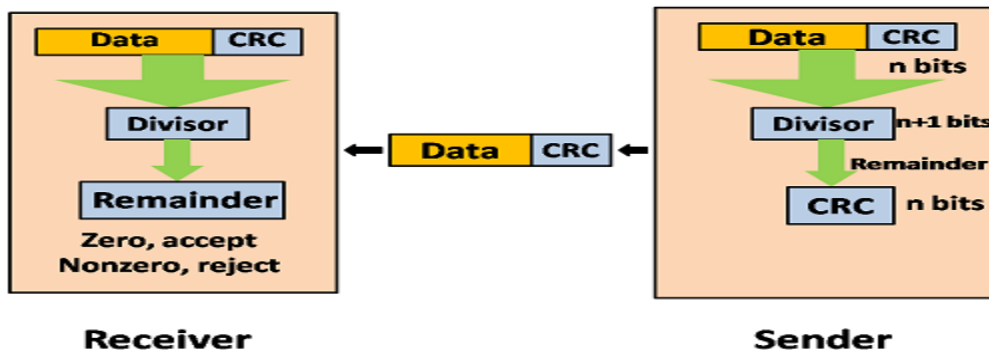
CRC is a redundancy error technique used to determine the error.

Following are the steps used in CRC for error detection:

- In CRC technique, a string of n 0s is appended to the data unit, and this n number is less than the number of bits in a predetermined number, known as division which is $n+1$ bits.
- Secondly, the newly extended data is divided by a divisor using a process known as binary division. The remainder generated from this division is known as CRC remainder.
- Thirdly, the CRC remainder replaces the appended 0s at the end of the original data. This newly generated unit is sent to the receiver.
- The receiver receives the data followed by the CRC remainder. The receiver will treat this whole unit as a single unit, and it is divided by the same divisor that was used to find the CRC remainder.

If the resultant of this division is zero which means that it has no error, and the data is accepted.

If the resultant of this division is not zero which means that the data consists of an error. Therefore, the data is discarded.

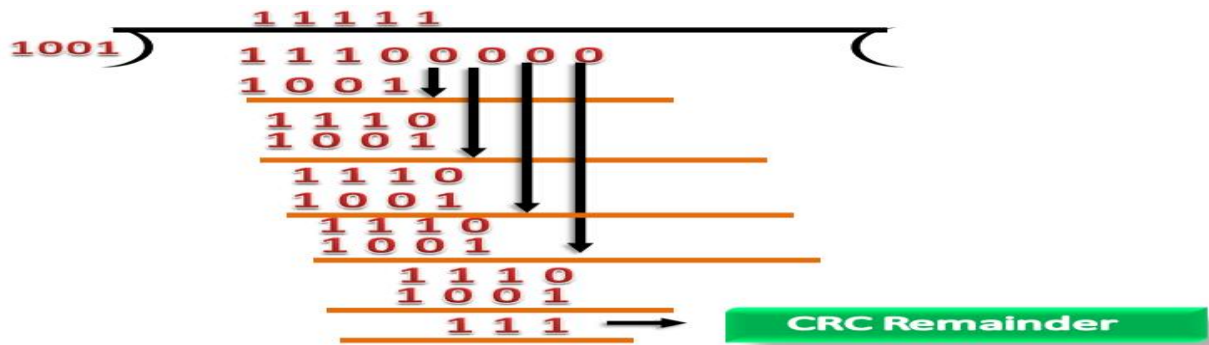


Let's understand this concept through an example:

Suppose the original data is 11100 and divisor is 1001.

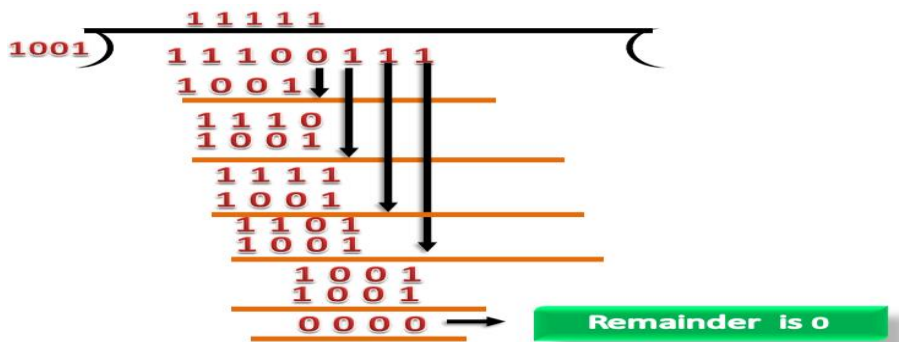
CRC Generator

- A CRC generator uses a modulo-2 division. Firstly, three zeroes are appended at the end of the data as the length of the divisor is 4 and we know that the length of the string 0s to be appended is always one less than the length of the divisor.
- Now, the string becomes 11100000, and the resultant string is divided by the divisor 1001.
- The remainder generated from the binary division is known as CRC remainder. The generated value of the CRC remainder is 111.
- CRC remainder replaces the appended string of 0s at the end of the data unit, and the final string would be 11100111 which is sent across the network.



CRC Checker

- The functionality of the CRC checker is similar to the CRC generator.
- When the string 11100111 is received at the receiving end, then CRC checker performs the modulo-2 division.
- A string is divided by the same divisor, i.e., 1001.
- In this case, CRC checker generates the remainder of zero. Therefore, the data is accepted.



Error Correction

Error Correction codes are used to detect and correct the errors when data is transmitted from the sender to the receiver.

Error Correction can be handled in two ways:

- **Backward error correction:** Once the error is discovered, the receiver requests the sender to retransmit the entire data unit.
- **Forward error correction:** In this case, the receiver uses the error-correcting code which automatically corrects the errors.

A single additional bit can detect the error, but cannot correct it.

For correcting the errors, one has to know the exact position of the error. For example, If we want to calculate a single-bit error, the error correction code will determine which one of seven bits is in error. To achieve this, we have to add some additional redundant bits.

Suppose r is the number of redundant bits and d is the total number of the data bits. The number of redundant bits r can be calculated by using the formula:

$$2^r \geq d + r + 1$$

The value of r is calculated by using the above formula. For example, if the value of d is 4, then the possible smallest value that satisfies the above relation would be 3.

To determine the position of the bit which is in error, a technique developed by R.W Hamming is Hamming code which can be applied to any length of the data unit and uses the relationship between data units and redundant units.

Hamming Code

Parity bits: The bit which is appended to the original data of binary bits so that the total number of 1s is even or odd.

Even parity: To check for even parity, if the total number of 1s is even, then the value of the parity bit is 0. If the total number of 1s occurrences is odd, then the value of the parity bit is 1.

Odd Parity: To check for odd parity, if the total number of 1s is even, then the value of parity bit is 1. If the total number of 1s is odd, then the value of parity bit is 0.

Algorithm of Hamming code:

- An information of ' d ' bits are added to the redundant bits ' r ' to form $d+r$.
- The location of each of the $(d+r)$ digits is assigned a decimal value.
- The ' r ' bits are placed in the positions $1, 2, \dots, 2^{k-1}$.
- At the receiving end, the parity bits are recalculated. The decimal value of the parity bits determines the position of an error.

Relationship b/w Error position & binary number.

Error Position	Binary Number
0	000
1	001
2	010
3	011
4	100
5	101
6	110
7	111

Let's understand the concept of Hamming code through an example:

Suppose the original data is 1010 which is to be sent.

Total number of data bits ' d ' = 4

Number of redundant bits r : $2^r \geq d + r + 1$

$$2^r \geq 4 + r + 1$$

Therefore, the value of r is 3 that satisfies the above relation.

Total number of bits = $d + r = 4 + 3 = 7$;

Determining the position of the redundant bits

The number of redundant bits is 3. The three bits are represented by r_1 , r_2 , r_4 . The position of the redundant bits is calculated with corresponds to the raised power of 2. Therefore, their corresponding positions are 1 , 2^1 , 2^2 .

1. The position of $r_1 = 1$
2. The position of $r_2 = 2$
3. The position of $r_4 = 4$

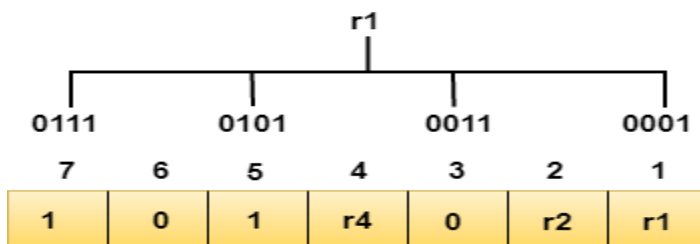
Representation of Data on the addition of parity bits:

7	6	5	4	3	2	1
1	0	1	r_4	0	r_2	r_1

Determining the Parity bits

Determining the r_1 bit

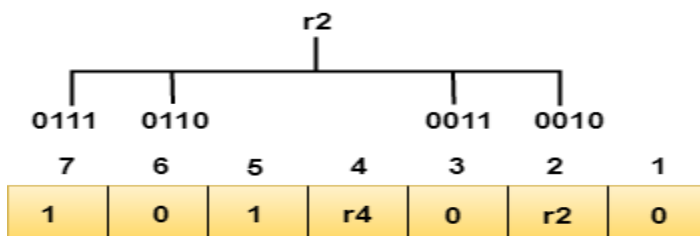
The r_1 bit is calculated by performing a parity check on the bit positions whose binary representation includes 1 in the first position.



We observe from the above figure that the bit positions that includes 1 in the first position are 1, 3, 5, 7. Now, we perform the even-parity check at these bit positions. The total number of 1 at these bit positions corresponding to r_1 is **even, therefore, the value of the r_1 bit is 0.**

Determining r_2 bit

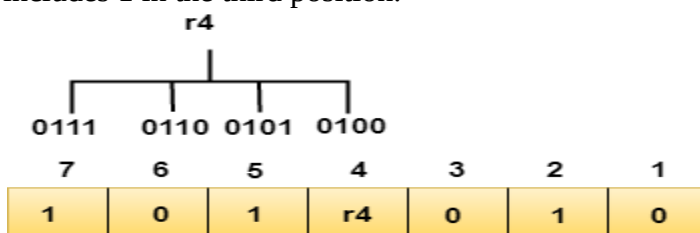
The r_2 bit is calculated by performing a parity check on the bit positions whose binary representation includes 1 in the second position.



We observe from the above figure that the bit positions that includes 1 in the second position are 2, 3, 6, 7. Now, we perform the even-parity check at these bit positions. The total number of 1 at these bit positions corresponding to r_2 is **odd, therefore, the value of the r_2 bit is 1.**

Determining r_4 bit

The r_4 bit is calculated by performing a parity check on the bit positions whose binary representation includes 1 in the third position.



We observe from the above figure that the bit positions that includes 1 in the third position are **4, 5, 6, 7**. Now, we perform the even-parity check at these bit positions. The total number of 1 at these bit positions corresponding to r4 is **even, therefore, the value of the r4 bit is 0**.

Data transferred is given below:

7	6	5	4	3	2	1
1	0	1	0	0	1	0

Suppose the 4th bit is changed from 0 to 1 at the receiving end, then parity bits are recalculated.

R1 bit

The bit positions of the r1 bit are 1,3,5,7

r1						
7	6	5	4	3	2	1
1	0	1	1	0	1	0

We observe from the above figure that the binary representation of r1 is 1100. Now, we perform the even-parity check, the total number of 1s appearing in the r1 bit is an even number. Therefore, the value of r1 is 0.

R2 bit

The bit positions of r2 bit are 2,3,6,7.

r2						
7	6	5	4	3	2	1
1	0	1	0	0	1	0

We observe from the above figure that the binary representation of r2 is 1001. Now, we perform the even-parity check, the total number of 1s appearing in the r2 bit is an even number. Therefore, the value of r2 is 0.

R4 bit

The bit positions of r4 bit are 4,5,6,7.

r4						
7	6	5	4	3	2	1
1	0	1	1	0	1	0

We observe from the above figure that the binary representation of r4 is 1011. Now, we perform the even-parity check, the total number of 1s appearing in the r4 bit is an odd number. Therefore, the value of r4 is 1.

- The binary representation of redundant bits, i.e., r4r2r1 is 100, and its corresponding decimal value is 4. Therefore, the error occurs in a 4th bit position. The bit value must be changed from 1 to 0 to correct the error.

➤ Data Link Control and Protocol

➤ Stop and Wait Protocol

Before understanding the stop and Wait protocol, we first know about the error control mechanism. The error control mechanism is used so that the received data should be exactly same whatever sender has sent the data. The error control mechanism is divided into two categories, i.e., Stop and Wait ARQ and sliding window. The sliding window is further divided into two categories, i.e., Go Back N, and Selective Repeat. Based on the usage, the people select the error control mechanism whether it is **stop and wait** or **sliding window**.

What is Stop and Wait protocol?

Here stop and wait means, whatever the data that sender wants to send, he sends the data to the receiver. After sending the data, he stops and waits until he receives the acknowledgment from the receiver. The stop and wait protocol is a flow control protocol where flow control is one of the services of the data link layer.

It is a data-link layer protocol which is used for transmitting the data over the noiseless channels. It provides unidirectional data transmission which means that either sending or receiving of data will take place at a time. It provides flow-control mechanism but does not provide any error control mechanism. The idea behind the usage of this frame is that when the sender sends the frame then he waits for the acknowledgment before sending the next frame.

Primitives of Stop and Wait Protocol

The primitives of stop and wait protocol are:

Sender side

Rule 1: Sender sends one data packet at a time.

Rule 2: Sender sends the next packet only when it receives the acknowledgment of the previous packet.

Therefore, the idea of stop and wait protocol in the sender's side is very simple, i.e., send one packet at a time, and do not send another packet before receiving the acknowledgment.

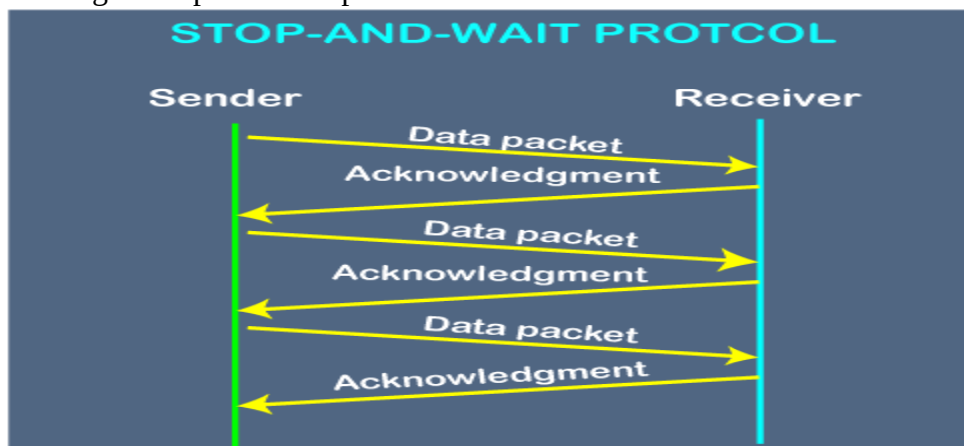
Receiver side

Rule 1: Receive and then consume the data packet.

Rule 2: When the data packet is consumed, receiver sends the acknowledgment to the sender.

Therefore, the idea of stop and wait protocol in the receiver's side is also very simple, i.e., consume the packet, and once the packet is consumed, the acknowledgment is sent. This is known as a flow control mechanism.

Working of Stop and Wait protocol



The above figure shows the working of the stop and waits protocol. If there is a sender and receiver, then sender sends the packet and that packet is known as a data packet. The sender will not send the second packet without receiving the acknowledgment of the first packet. The receiver sends the acknowledgment for the data packet that it has received. Once the acknowledgment is received, the sender sends the next packet. This process continues until all the packet are not sent. The main advantage of this protocol is its simplicity but it has some disadvantages also. For example, if there are 1000 data packets to be sent, then



all the 1000 packets cannot be sent at a time as in Stop and Wait protocol, one packet is sent at a time.

Disadvantages of Stop and Wait protocol

The following are the problems associated with a stop and wait

protocol: Problems occur due to lost data

Suppose the sender sends the data and the data is lost. The receiver is waiting for the data for a long time. Since the data is not received by the receiver, so it does not send any acknowledgment. Since the sender does not receive any acknowledgment so it will not send the next packet. This problem occurs due to the lost data.

In this case, two problems occur:

- Sender waits for an infinite amount of time for an acknowledgment.
- Receiver waits for an infinite amount of time for a data.

2. Problems occur due to lost acknowledgment



Suppose the sender sends the data and it has also been received by the receiver. On receiving the packet, the receiver sends the acknowledgment. In this case, the acknowledgment is lost in a network, so there is no chance for the sender to receive the acknowledgment. There is also no chance for the sender to send the next packet as in stop and wait protocol, the next packet cannot be sent until the acknowledgment of the previous packet is received.

In this case, one problem occurs:

- Sender waits for an infinite amount of time for an acknowledgment.

4. Problem due to the delayed data or acknowledgment



Suppose the sender sends the data and it has also been received by the receiver. The receiver then sends the acknowledgment but the acknowledgment is received after the timeout period on the sender's side. As the acknowledgment is received late, so acknowledgment can be wrongly considered as the acknowledgment of some other data packet

Error Control

Error Control is a technique of error detection and retransmission.

Categories of Error Control:



Stop-and-wait ARQ

Stop-and-wait ARQ is a technique used to retransmit the data in case of damaged or lost frames. This technique works on the principle that the sender will not transmit the next frame until it receives the acknowledgement of the last transmitted frame.

- **Go-Back-n ARQ:** In Go-Back-N ARQ protocol, if one frame is lost or damaged, then it retransmits all the frames after which it does not receive the positive ACK.

Selective-Reject ARQ

- Selective-Reject ARQ technique is more efficient than Go-Back-n ARQ.
- In this technique, only those frames are retransmitted for which negative acknowledgement (NAK) has been received.
- The receiver storage buffer keeps all the damaged frames on hold until the frame in error is correctly received.
- The receiver must have an appropriate logic for reinserting the frames in a correct order.
- The sender must consist of a searching mechanism that selects only the requested frame for retransmission.

➤ Go-Back-N

Before understanding the working of Go-Back-N ARQ, we first look at the sliding window protocol. As we know that the sliding window protocol is different from the stop-and-wait protocol. In the stop-and-wait protocol, the sender can send only one frame at a time and cannot send the next frame without receiving the acknowledgment of the previously sent frame, whereas, in the case of sliding window protocol, the multiple frames can be sent at a time. The variations of sliding window protocol are Go-Back-N ARQ and Selective Repeat ARQ.

What is Go-Back-N ARQ?(Automatic Repeat Request)

In Go-Back-N ARQ, **N** is the sender's window size. Suppose we say that Go-Back-3, which means that the three frames can be sent at a time before expecting the acknowledgment from the receiver.

It uses the principle of protocol pipelining in which the multiple frames can be sent before receiving the acknowledgment of the first frame. If we have five frames and the concept is Go-Back-3, which means that the three frames can be sent, i.e., frame no 1, frame no 2, frame no 3 can be sent before expecting the acknowledgment of frame no 1.

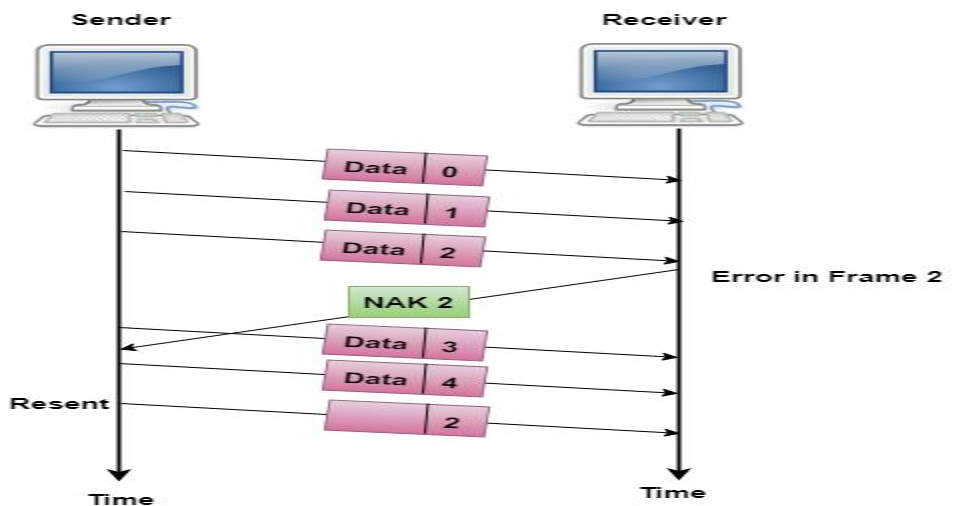
In Go-Back-N ARQ, the frames are numbered sequentially as Go-Back-N ARQ sends the multiple frames at a time that requires the numbering approach to distinguish the frame from another frame, and these numbers are known as the sequential numbers.

The number of frames that can be sent at a time totally depends on the size of the sender's window. So, we can say that 'N' is the number of frames that can be sent at a time before receiving the acknowledgment from the receiver.

If the acknowledgment of a frame is not received within an agreed-upon time period, then all the frames available in the current window will be retransmitted. Suppose we have sent the frame no 5, but we didn't receive the acknowledgment of frame no 5, and the current window is holding three frames, then these three frames will be retransmitted.

The sequence number of the outbound frames depends upon the size of the sender's window. Suppose the sender's window size is 2, and we have ten frames to send, then the sequence numbers will not be 1,2,3,4,5,6,7,8,9,10. Let's understand through an example.

- N is the sender's window size.

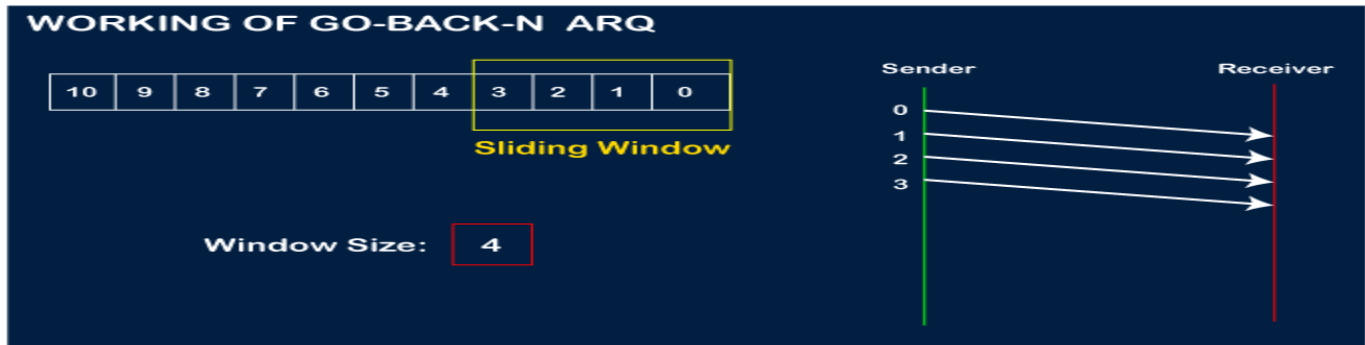


- If the size of the sender's window is 4 then the sequence number will be 0,1,2,3,0,1,2,3,0,1,2, and so on. The number of bits in the sequence number is 2 to generate the binary sequence 00,01,10,11.

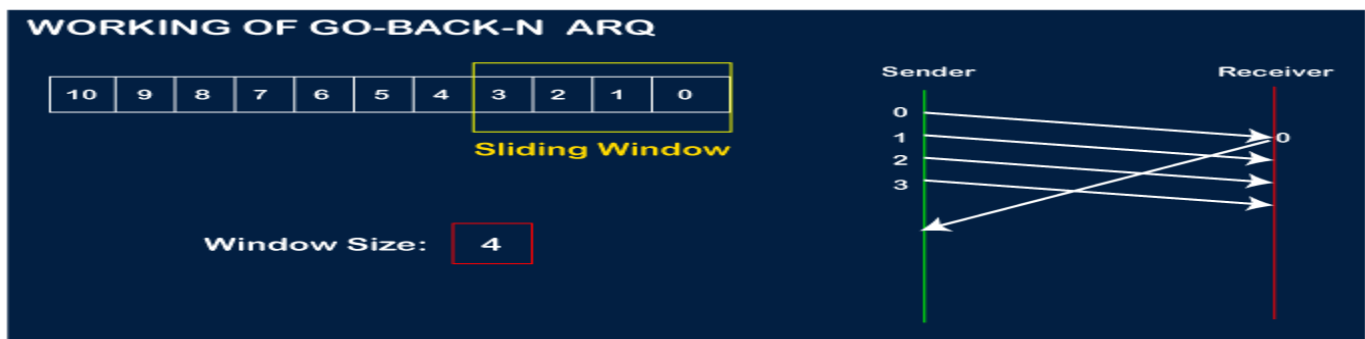
Working of Go-Back-N ARQ

Suppose there are a sender and a receiver, and let's assume that there are 11 frames to be sent. These frames are represented as 0,1,2,3,4,5,6,7,8,9,10, and these are the sequence numbers of the frames. Mainly, the sequence number is decided by the sender's window size. But, for the better understanding, we took the running sequence numbers, i.e., 0,1,2,3,4,5,6,7,8,9,10. Let's consider the window size as 4, which means that the four frames can be sent at a time before expecting the acknowledgment of the first frame.

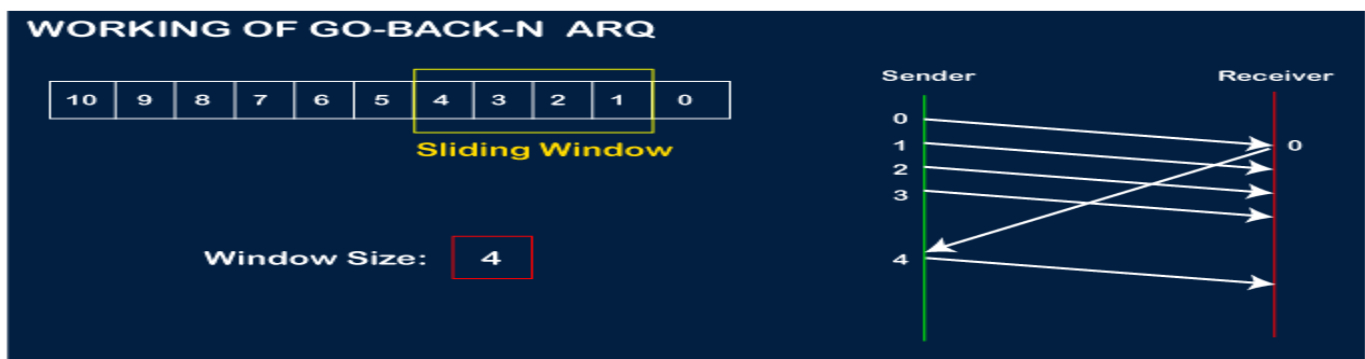
Step 1: Firstly, the sender will send the first four frames to the receiver, i.e., 0,1,2,3, and now the sender is expected to receive the acknowledgment of the 0th frame.



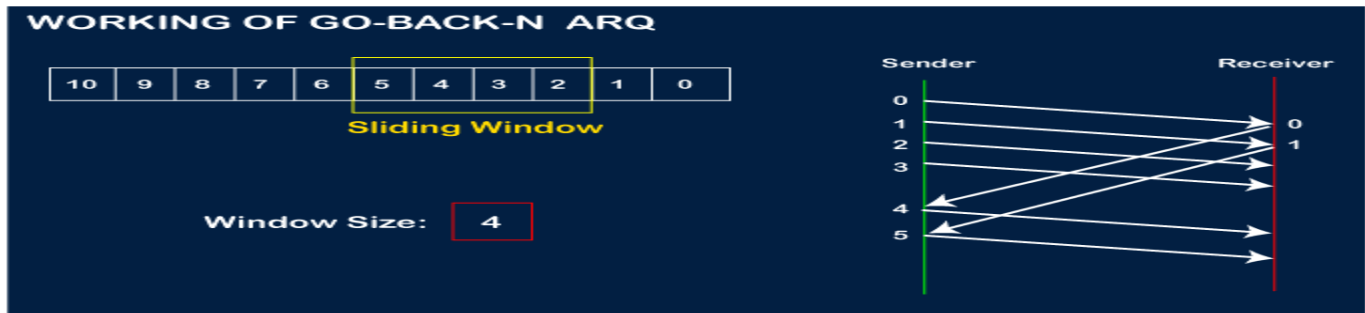
Let's assume that the receiver has sent the acknowledgment for the 0 frame, and the receiver has successfully received it.



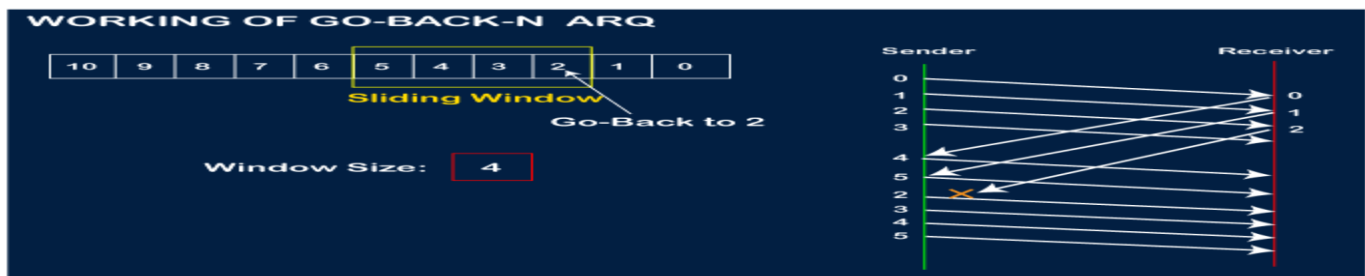
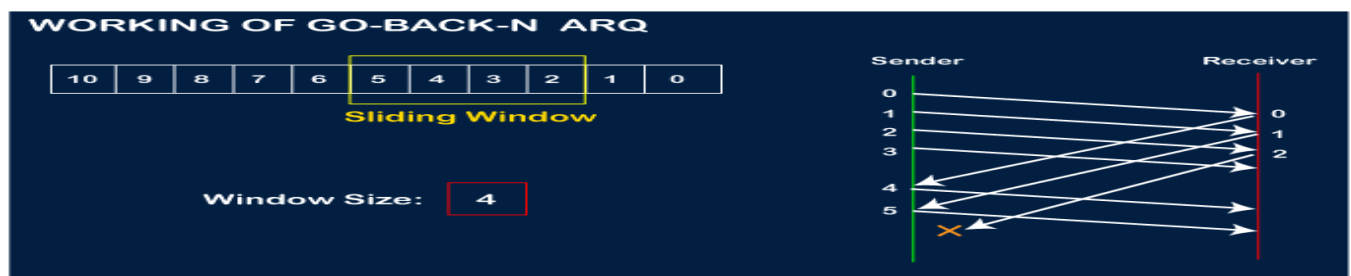
The sender will then send the next frame, i.e., 4, and the window slides containing four frames (1,2,3,4).



The receiver will then send the acknowledgment for the frame no 1. After receiving the acknowledgment, the sender will send the next frame, i.e., frame no 5, and the window will slide having four frames (2,3,4,5).



Now, let's assume that the receiver is not acknowledging the frame no 2, either the frame is lost, or the acknowledgment is lost. Instead of sending the frame no 6, the sender Go-Back to 2, which is the first frame of the current window, retransmits all the frames in the current window, i.e., 2,3,4,5.



Important points related to Go-Back-N ARQ:

- In Go-Back-N, N determines the sender's window size, and the size of the receiver's window is always 1.
- It does not consider the corrupted frames and simply discards them.
- It does not accept the frames which are out of order and discards them.
- If the sender does not receive the acknowledgment, it leads to the retransmission of all the current window frames.

➤ HDLC

High-level Data Link Control (HDLC) is a bit-oriented protocol for communication over point-to-point and multipoint links. It implements the Stop-and-Wait protocol we discussed earlier. Although this protocol is more a theoretical issue than practical, most of the concept defined in this protocol is the basis for other practical protocols such as PPP, which we discuss next, or the Ethernet protocol.

Configurations and Transfer Modes

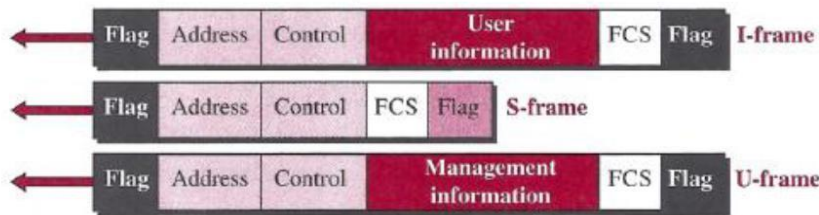
HDLC provides two common transfer modes that can be used in different configurations:

Normal response mode (NRM) and Asynchronous balanced mode (ABM)

In normal response mode (NRM), the station configuration is unbalanced. We have one primary station and multiple secondary stations. A primary station can send commands; a secondary station can only respond. The NRM is used for both point-to-point and multipoint links, as shown in below Figure. In ABM, the configuration is balanced. The link is point-to-point, and each station can function as a

primary and a secondary (acting as peers) this is the common mode today. link itself. Each frame in HDLC may contain up to six fields: a beginning flag field, an

HDLC frames



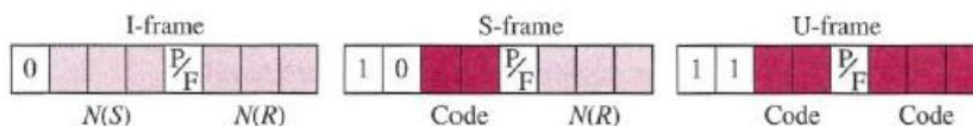
address field, a control field, an information field, a frame check sequence (FCS) field, and an ending flag field. In multiple-frame transmissions, the ending flag of one frame can serve as the beginning flag of the next

frame.

Let us now discuss the fields and their use in different frame types.

- D Flag field. This field contains synchronization pattern 01111110, which identifies both the beginning and the end of a frame.
- D Address field. This field contains the address of the secondary station. If a primary station created the frame, it contains a to address. If a secondary station creates the frame, it contains a from address. The address field can be one byte or several bytes long, depending on the needs of the network.
- Control field. The control field is one or two bytes used for flow and error control.
 - Information field. The information field contains the user's data from the network layer or management information. Its length can vary from one network to another.
 - FCS field. The frame check sequence (FCS) is the HDLC error detection field. It can contain either a 2- or 4-byte CRC.

Control field format for the different frame types



The control field determines the type of frame and defines its functionality. So let us discuss the format of this field in detail. The format is specific for the type of frame, as shown in below Figure.

Control Fieldfor I-Frames

I-frames are designed to carry user data from the network layer. In addition, they can include flow- and error-control information (piggybacking). The subfields in the control field are used to define these functions. The first bit defines the type. If the first bit of the control field is 0, this means the frame is an I-frame. The next 3 bits, called N(S), define the sequence number of the frame. Note that with 3 bits, we can define a sequence number between 0 and 7.

The last 3bits, called N(R), correspond to the acknowledgment number when piggybacking is used. The single bit between N(S) and N(R) is called the PIF bit. The PIF field is a single bit with a dual purpose. It has meaning only when it is set (bit = 1) and can mean poll or final. It means poll when the frame is sent by a primary station to a secondary (when the address field contains the address of the receiver). It means final when the frame is sent by a secondary to a primary (when the address field contains the address of the sender).

Control Field for S-Frames

Supervisory frames are used for flow and error control whenever piggybacking is either impossible or inappropriate. S-frames do not have information fields. If the first 2 bits of the control field are 10, this means the frame is an S-frame. The last 3 bits, called N(R), correspond to the acknowledgment number (ACK) or negative acknowledgment number (NAK), depending on the type of

S-frame. The 2 bits called code are used to define the type of S-frame itself. With 2 bits, we can have four types of S-frames, as described below:

Receive ready (RR)

If the value of the code subfield is 00, it is an RR S-frame. This kind of frame acknowledges the receipt of a safe and sound frame or group of frames. In this case, the value of the N(R) field defines the acknowledgment number.

□ **Receive not ready (RNR)** If the value of the code subfield is 10, it is an RNR S frame. This kind of frame is an RR frame with additional functions. It acknowledges the receipt of a frame or group of frames, and it announces that the receiver is busy and cannot receive more frames. It acts as a kind of congestion-control mechanism by asking the sender to slow down. The value of N(R) is the acknowledgment number.

□ **Reject (REJ)** If the value of the code subfield is 01, it is an REJ S-frame. This is a NAK frame, but not like the one used for Selective Repeat ARQ. It is a NAK that can be used in Go-Back-N ARQ to improve the efficiency of the process by informing the sender, before the sender timer expires, that the last frame is lost or damaged. The value of N(R) is the negative acknowledgment number.

□ **Selective reject (SREJ)** If the value of the code subfield is 11, it is an SREJ S frame. This is a NAK frame used in Selective Repeat ARQ. Note that the HDLC Protocol uses the term selective reject instead of selective repeat. The value of N(R) is the negative acknowledgment number.

Control Field or V-Frames

Unnumbered frames are used to exchange session management and control information between connected devices. Unlike S-frames, U-frames contain an information field, but one used for system management information, not user data. As with S-frames, however, much of the information carried by If-frames is contained in codes included in the control field. If-frame codes are divided into two sections: a 2-bit prefix before the P/F bit and a 3-bit suffix after the P/F bit. Together, these two segments (5 bits) can be used to create up to 32 different types of U-frames.

Control Field for V-Frames

Unnumbered frames are used to exchange session management and control information between connected devices. Unlike S-frames, U-frames contain an information field, but one used for system management information, not user data. As with S-frames, however, much of the information carried by U-frames is contained in codes included in the control field. U-frame codes are divided into two sections: a 2-bit prefix before the P/F bit and a 3-bit suffix after the P/F bit. Together, these two segments (5 bits) can be used to create up to 32 different types of If-frames.

➤ POINT-TO-POINT PROTOCOL (PPP)

One of the most common protocols for point-to-point access is the **Point-to-Point Protocol (PPP)**. Today, millions of Internet users who need to connect their home computers to the server of an Internet service provider use PPP. The majority of these users have a traditional modem; they are connected to the Internet through a telephone line, which provides the services of the physical layer. But to control and manage the transfer of data, there is a need for a point-to-point protocol at the data-link layer. PPP is by far the most common.

Services

The designers of PPP have included several services to make it suitable for a point-to-point protocol, but have ignored some traditional services to make it simple.

Services Provided by PPP

PPP defines the format of the frame to be exchanged between devices. It also defines how two devices can negotiate the establishment of the link and the exchange of data. PPP is designed to accept payloads from several network layers (not only IP). Authentication is also provided in the protocol, but it is optional. The new version of PPP, called Multilink PPP, provides connections over multiple links. One

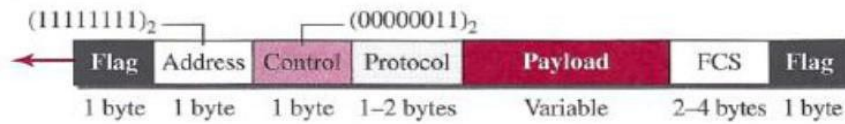
interesting feature of PPP is that it provides network address configuration. This is particularly useful when a home user needs a temporary network address to connect to the Internet.

Framing

PPP uses a character-oriented (or byte-oriented) frame. Below figure shows the format.

- **Flag** A PPP frame starts and ends with a 1-byte flag with the bit pattern 01111110.

PPP frame format



Address The address field in this protocol is a constant value and set to 11111111 (broadcast address).

□ **D Control** This field is set to the constant value 00000011 (imitating unnumbered frames in HDLC). As we will discuss later, PPP does not provide any flow control. Error control is also limited to error detection.

□ **Protocol** The protocol field defines what is being carried in the data field: either user data or other information. This field is by default 2 bytes long, but the two parties can agree to use only 1 byte.

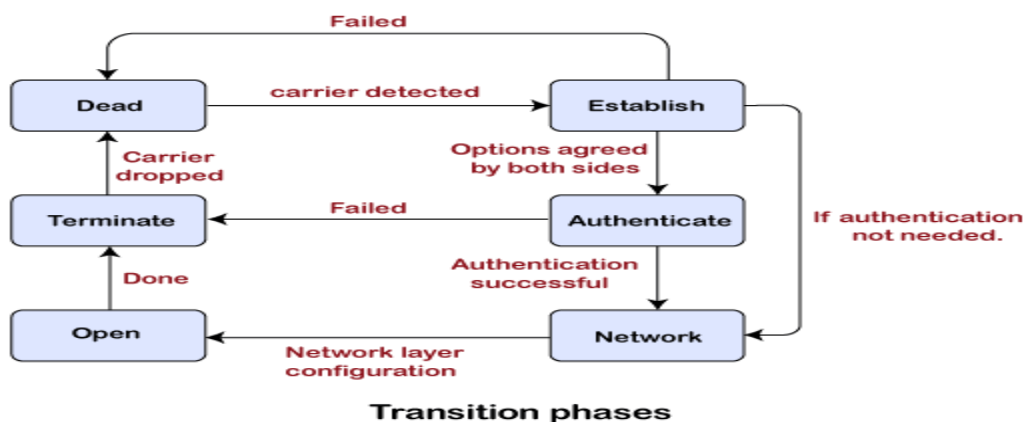
□ **Payload field** This field carries either the user data or other information that we will discuss shortly. The data field is a sequence of bytes with the default of a maximum of 1500 bytes; but this can be changed during negotiation. The data field is byte-stuffed if the flag byte pattern appears in this field. Because there is no field defining the size of the data field, padding is needed if the size is less than the maximum default value or the maximum negotiated value. The frame check sequence (FCS) is simply a 2-byte or 4-byte standard CRC.

Byte Stuffing

Since PPP is a byte-oriented protocol, the flag in PPP is a byte that needs to be escaped whenever it appears in the data section of the frame. The escape byte is 01111101, which means that every time the flag like pattern appears in the data, this extra byte is stuffed to tell the receiver that the next byte is not a flag. Obviously, the escape byte itself should be stuffed with another escape byte.

Transition phases of PPP protocol

The following are the transition phases of a PPP protocol:



- **Dead:** Dead is a transition phase which means that the link is not used or there is no active carrier at the physical layer.
- **Establish:** If one of the nodes starts working then the phase goes to the establish phase. In short, we can say that when the node starts communication or carrier is detected then it moves from the dead to the establish phase.

- **Authenticate:** It is an optional phase which means that the communication can also moves to the authenticate phase. The phase moves from the establish to the authenticate phase only when both the communicating nodes agree to make the communication authenticated.
- **Network:** Once the authentication is successful, the network is established or phase is network. In this phase, the negotiation of network layer protocols take place.
- **Open:** After the establishment of the network phase, it moves to the open phase. Here open phase means that the exchange of data takes place. Or we can say that it reaches to the open phase after the configuration of the network layer.
- **Terminate:** When all the work is done then the connection gets terminated, and it moves to the terminate phase.

On reaching the terminate phase, the link moves to the dead phase which indicates that the carrier is dropped which was earlier created.

There are two more possibilities that can exist in the transition phase:

- The link moves from the authenticate to the terminate phase when the authentication is failed.
- The link can also move from the establish to the dead state when the carrier is failed.

➤ LANS

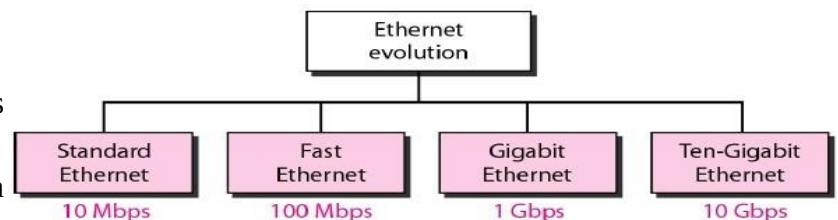
➤ INTRODUCTION

A local area network (LAN) is a computer network that is designed for a limited geographic area such as a building or a campus. Although a LAN can be used as an isolated network to connect computers in an organization for the sole purpose of sharing resources, most LANs today are also linked to a wide area network (WAN) or the Internet. The LAN market has seen several technologies such as Ethernet, Token Ring, Token Bus, FDDI, and ATM LAN. Some of these technologies survived for a while, but Ethernet is by far the dominant technology.

The original Ethernet was created in 1976 at Xerox's Palo Alto Research Center (PARC). Since then, it has gone through four generations:

- Standard Ethernet (10 Mbps),
- Fast Ethernet (100 Mbps),
- Gigabit Ethernet (1 Gbps), and
- Ten-Gigabit Ethernet (10 Gbps), as shown in Figure 13.1.

Figure 13.1 Ethernet evolution through four generations



➤ Standard Ethernet(IEEE 802.3)

Standard Ethernet also known as IEEE 802.3 was the LAN standard proposed by IEEE. Data rate for standard Ethernet is 10 Mbps.

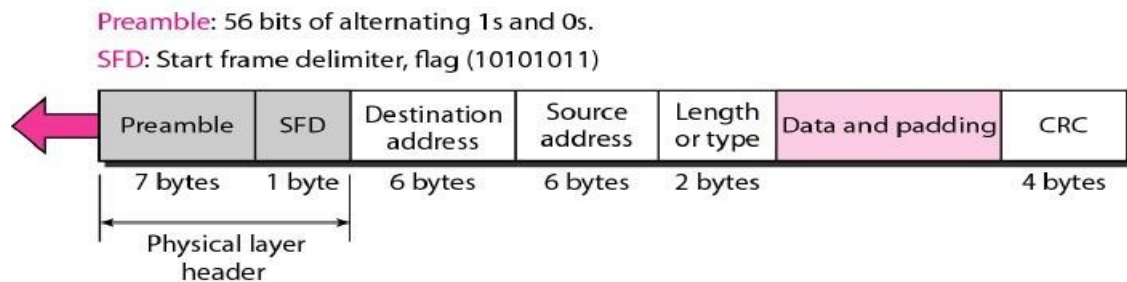
MAC Sublayer

In Standard Ethernet, the MAC sublayer governs the operation of the access method. It also frames the data received from the upper layer and passes them to the physical layer.

• Frame Format

The Ethernet frame contains seven fields: preamble, SFD, DA, SA, length or type of protocol data unit (PDU), upper-layer data, and the CRC. Ethernet does not provide any mechanism for acknowledging received frames, making it what is known as an unreliable medium. Acknowledgments must be implemented at the higher layers. The format of the MAC frame is shown in Figure 802.3 MAC frame

- i. **Preamble.** The first field of the 802.3 frame contains 7 bytes (56 bits) of alternating 0s and 1s that alerts the receiving system to the coming frame and enables it to synchronize its input timing. The pattern provides only an alert and a timing pulse. The 56-bit pattern allows the stations to miss some bits at the beginning of the frame. The preamble is actually added at the physical layer and is not (formally) part of the frame.



- ii. **Start frame delimiter (SFD).** The second field (1 byte: 10101011) signals the beginning of the frame. The SFD warns the station or stations that this is the last chance for synchronization. The last 2 bits is 11 and alerts the receiver that the next field is the destination address.
- iii. **Destination address (DA).** The DA field is 6 bytes and contains the physical address of the destination station or stations to receive the packet.
- iv. **Source address (SA).** The SA field is also 6 bytes and contains the physical address of the sender of the packet.
- v. **Length or type.** This field is defined as a type field or length field. The original Ethernet used this field as the type field to define the upper-layer protocol using the MAC frame. The IEEE standard used it as the length field to define the number of bytes in the data field. Both uses are common today.
- vi. **Data.** This field carries data encapsulated from the upper-layer protocols. It is a minimum of 46 and a maximum of 1500 bytes.
- vii. **CRC.** The last field contains error detection information, in this case a CRC-32.
- viii. **Frame Length** Ethernet has imposed restrictions on both the minimum and maximum lengths of a frame, as shown in Figure .

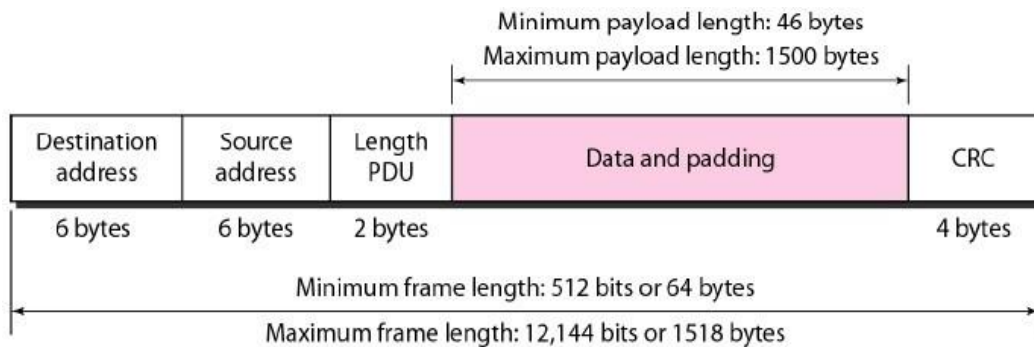


Figure Minimum and maximum lengths

The minimum length restriction is required for the correct operation of CSMA/CD as we will see shortly. An Ethernet frame needs to have a minimum length of 512 bits or 64 bytes. Part of this length is the header and the trailer. If we count 18 bytes of header and trailer (6 bytes of source address, 6 bytes of destination address, 2 bytes of length or type, and 4 bytes of CRC), then the minimum length of data from the upper layer is $64 - 18 = 46$ bytes. If the upper-layer packet is less than 46 bytes, padding is added to make up the difference.

The standard defines the maximum length of a frame (without preamble and SFD field) as 1518 bytes. If we subtract the 18 bytes of header and trailer, the maximum length of the payload is 1500 bytes. The maximum length restriction has two historical reasons. First, memory was very expensive when Ethernet was designed: a maximum length restriction helped to reduce the size of the buffer. Second, the maximum length restriction prevents one station from monopolizing the shared medium, blocking other stations that have data to send.

MAC Addressing

Each station on an Ethernet network (such as a PC, workstation, or printer) has its own network interface card (NIC). The NIC fits inside the station and provides the station with a 6-byte physical (MAC) address. As shown in Figure 13.4, the Ethernet address is 6 bytes (48 bits), normally written in hexadecimal notation, with a colon between the bytes.

Figure 13.4 Example of an Ethernet address in hexadecimal notation

06 : 01 : 02 : 01 : 2C : 4B

└────────────────────────────────┘
6 bytes = 12 hex digits = 48 bits

Unicast, Multicast, and Broadcast Addresses

Data is transmitted over a network by three simple methods i.e. Unicast, Broadcast, and Multicast Figure 13.5. So let's begin to summarize the difference between these three:

- **Unicast:** from one source to one destination i.e. One-to-One
- **Broadcast:** from one source to all possible destinations i.e. One-to-All
- **Multicast:** from one source to multiple destinations stating an interest in receiving the traffic i.e. One-to-Many

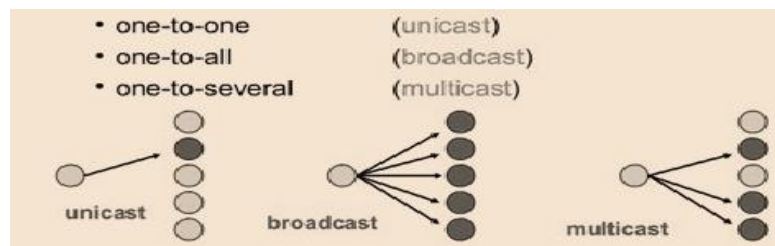


Figure Unicasting, Multicasting and Broadcasting

- A source address is always a unicast address as the frame comes from only one station.
- The destination address, however, can be unicast, multicast, or broadcast.

- Figure 13.6 shows how to distinguish a unicast address from a multicast address. If the least



significant bit of the first byte in a destination address is 0, the address is unicast; otherwise, it is multicast.

Figure Unicast and multicast MAC addresses

- A unicast destination address defines only one recipient; the relationship between the sender and the receiver is one-to-one.
- The broadcast address is a special case of the multicast address; the recipients are all the stations on the LAN. A broadcast destination address is forty-eight 1s.

Categories of Standard Ethernet

The Standard Ethernet defines several physical layer implementations; four of the most common, are shown in Figure 13.7.

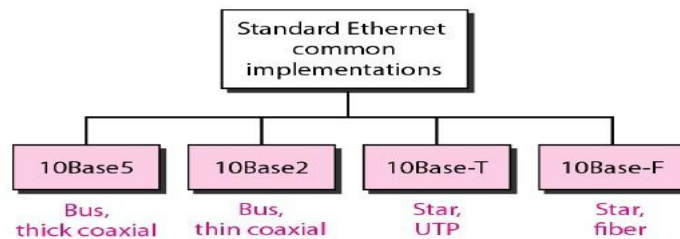


Figure Categories of Standard Ethernet

Encoding and Decoding

All standard implementations use digital signaling (baseband) at 10 Mbps. At the sender, data are converted to a digital signal using the Manchester scheme; at the receiver, the received signal is interpreted as Manchester and decoded into data.

- 10Base5: Thick Ethernet**

The first implementation is called **10Base5, thick Ethernet, or Thicknet**. The nickname derives from the cable, which is roughly the size of a garden hose and too stiff to bend with your hands. 10Base5 was the first Ethernet specification to use a bus topology with an external **transceiver** (transmitter/receiver) connected via a transceiver cable to a thick coaxial cable.



The transceiver is responsible for transmitting, receiving, and detecting collisions. The transceiver is connected to the station via a transceiver cable that provides separate paths for sending and receiving. This means that collision can only happen in the coaxial cable. The maximum length of the coaxial cable must not exceed 500 m, otherwise, there is excessive degradation of the signal. If a length of more than 500 m is needed, up to five segments, each a maximum of 500-meter, can be connected using repeaters.

exceed 500 m, otherwise, there is excessive degradation of the signal. If a length of more than 500 m is needed, up to five segments, each a maximum of 500-meter, can be connected using repeaters.

- **10Base2: Thin Ethernet**

The second implementation is called 10Base2, **thin** Ethernet, or Cheaper net. 10Base2 also uses a bus topology but the cable is much thinner and more flexible. The cable can be bent to pass very close to the stations. In this case, the transceiver is normally part of the network interface card (NIC), which is installed inside the station. Note that the collision here occurs in the thin coaxial cable. This implementation is more cost effective than 10Base5 because thin coaxial cable is less expensive than thick coaxial and the tee connections are much cheaper than taps. Installation is simpler because the thin coaxial cable is very flexible. However, the length of each segment cannot exceed 185 m (close to 200 m) due to the high level of attenuation in thin coaxial cable.

- **10Base-T: Twisted-Pair Ethernet**

The third implementation is called 10Base-T or twisted-pair Ethernet. 10Base-T uses a physical star topology. The stations are connected to a hub via two pairs of twisted cable. Note that two pairs of twisted cable create two paths (one for sending and one for receiving) between the station and the hub. Any collision here happens in the hub. Compared to 10Base5 or 10Base2, we can see that the hub actually replaces the coaxial cable as far as a collision is concerned. The maximum length of the twisted cable here is defined as 100 m, to minimize the effect of attenuation in the twisted cable.

- **10Base-F: Fiber Ethernet**

Although there are several types of optical fiber 10-Mbps Ethernet, the most common is called 10Base-F. 10Base-F uses a star topology to connect stations to a hub. The stations are connected to the hub using two fiber-optic cables.

Changes in the Standard

The 10-Mbps Standard Ethernet has gone through several changes before moving to the higher data rates. These changes actually opened the road to the evolution of the Ethernet to become compatible with other high-data-rate LANs. We discuss some of these changes in this section.

Bridged Ethernet

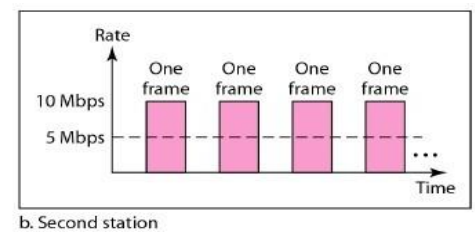
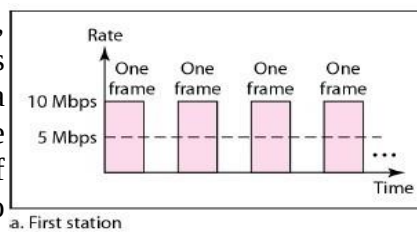
The first step in the Ethernet evolution was the division of a LAN by bridges. A Bridge is a two port switch used to connect two segments of a LAN. Bridges have two effects on an Ethernet LAN:

- They **raise the bandwidth** and
- They **separate collision domains**.

Raising the Bandwidth

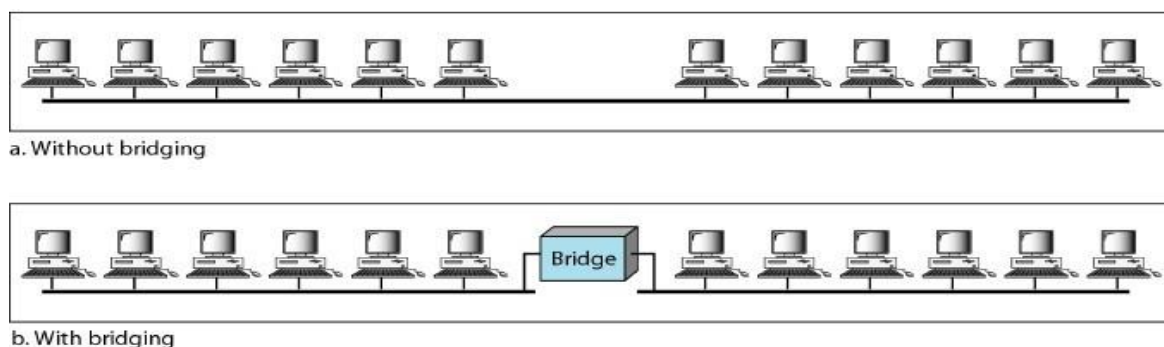
In an unbridged Ethernet network, the total capacity (10 Mbps) is shared among all stations with a frame to send; the stations share the bandwidth of the network. If only one station has frames to send, it benefits from the total capacity (10 Mbps). But if more than one station needs to use the network, the capacity is shared. For example, if two stations have a lot of frames to send, they probably alternate in usage. When one station is sending, the other one refrains from sending. We can say that, in this case, each station on average, sends at a rate of 5 Mbps. Figure 13.8 shows the situation.

Figure Sharing bandwidth



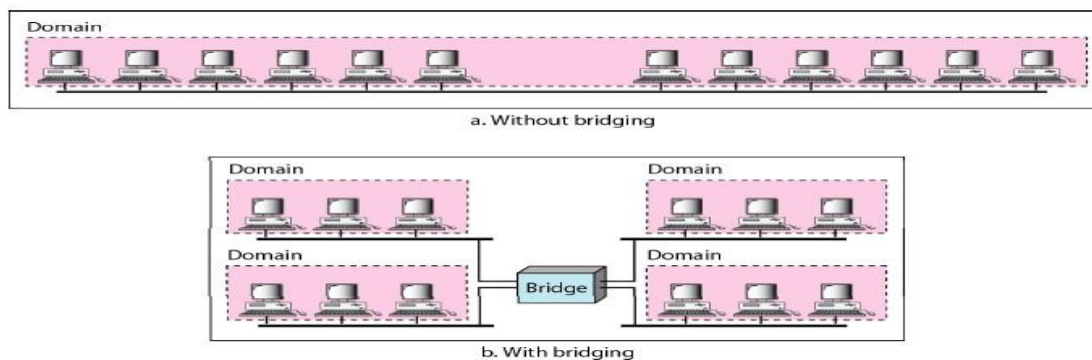
A bridge divides the network into two or more networks. Bandwidth-wise, each network is independent. For example, in Figure 13.9, a network with 12 stations is divided into two networks, each with 6 stations. Now each network has a capacity of 10 Mbps. The 10-Mbps capacity in each segment is now shared between 6 stations (actually 7 because the bridge acts as a station in each segment), not 12 stations. In a network with a heavy load, each station theoretically is offered 10/6 Mbps instead of 10/12 Mbps, assuming that the traffic is not going through the bridge. It is obvious that if we further divide the network, we can gain more bandwidth for each segment. For example, if we use a four-port bridge, each station is now offered 10/3 Mbps, which is 4 times more than an unbridged network.

Figure A network with and without a bridge



Separating Collision Domains

Another advantage of a bridge is the separation of the collision domain. Figure 13.10 shows the collision domains for an unbridged and a bridged network. You can see that

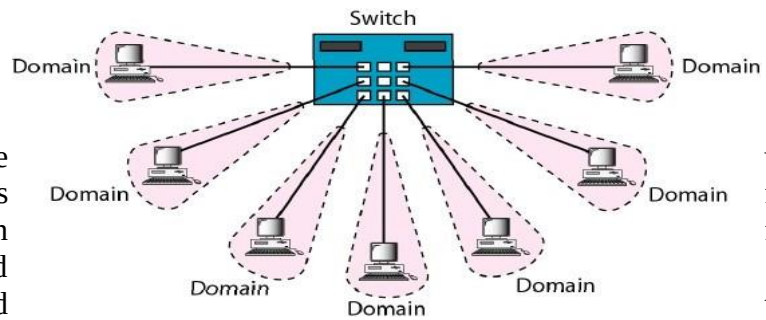


the collision domain becomes much smaller and the probability of collision is reduced tremendously. Without bridging, 12 stations contend for access to the medium; with bridging only 3 stations contend for access to the medium.

Figure Collision domains in an unbridged network and a bridged network

Switched Ethernet

The idea of a bridged LAN can be extended to a switched LAN. Instead of having two to four networks, why not have N networks, where N is the number of stations on the LAN? In other words, if we can have a multiple-port bridge, why not have an N-port switch? In this way, the bandwidth is shared only between the station and the switch (5 Mbps each). In addition, the collision domain is divided into N domains.



A layer 2 switch is an N-port bridge with additional sophistication that allows handling of the packets. Evolution bridged Ethernet to a switched Ethernet was a big step that opened way to an even faster Ethernet, as we will see. Figure 13.11 shows a switched LAN.

with faster from a the

Full-Duplex Ethernet

One of the limitations of 10Base5 and 10Base2 is that communication is half-duplex (10Base-T is always full-duplex); a station can either send or receive, but may not do both at the same time. The next step in the evolution was to move from switched Ethernet to full-duplex switched Ethernet. The full-duplex mode increases the capacity of each domain from 10 to 20 Mbps. Figure 13.12 shows a switched Ethernet in full-duplex mode. Note that instead of using one link between the station and the switch, the configuration uses two links: one to transmit and one to receive.

➤ Fast Ethernet(IEEE 802.3u)

Fast Ethernet was designed to compete with LAN protocols such as FDDI or Fiber Channel (or Fibre Channel, as it is sometimes spelled). IEEE created Fast Ethernet under the name 802.3u. Fast Ethernet is backward-compatible with Standard Ethernet, but it can transmit data 10 times faster at a rate of 100 Mbps. The goals of Fast Ethernet can be summarized as follows:

- Upgrade the data rate to 100 Mbps.
- Make it compatible with Standard Ethernet.
- Keep the same 48-bit address.
- Keep the same frame format.
- Keep the same minimum and maximum frame lengths.

Topology

Fast Ethernet is designed to connect two or more stations together. If there are only two stations, they can be connected point-to-point. Three or more stations need to be connected in a star topology with a hub or a switch at the center, as shown in Figure 13.13.

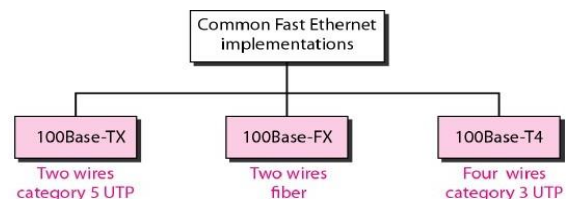
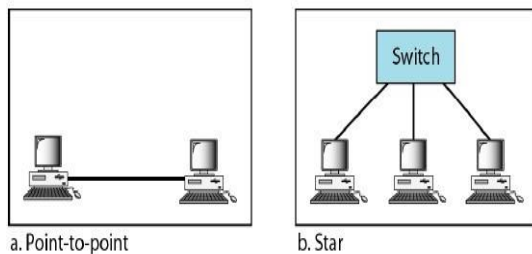


Figure Fast Ethernet topology

Implementation

Fast Ethernet implementation at the physical layer can be categorized as either two-wire or four-wire. The two-wire implementation can be either category 5 UTP (100Base-TX) or fiber-optic cable (100Base-FX). The four-wire implementation is designed only for category 3 UTP (100Base-T4). See Figure 13.14.

Figure 13.14 Fast Ethernet implementations Table 13.2 Summary of Fast Ethernet implementations

➤ WireLessLans

A **wireless LAN (WLAN)** is a wireless computer network that links two or more devices using wireless communication to form a local area network (LAN) within a limited area such as a home, school, computer laboratory, campus, or office building. This gives users the ability to move around within the area and remain connected to the network.

Wireless LANs based on the IEEE 802.11 standards are the most widely used computer networks in the world. These are commonly called Wi-Fi, which is a trademark belonging to the Wi-Fi Alliance. They are used for home and small office networks that link together laptop computers, printers, smartphones, Web TVs and gaming devices with a wireless router which links them to the internet. Hotspots provided by routers at restaurants, coffee shops, hotels, libraries, and airports allow consumers to access the internet with portable wireless devices

Applications

Wireless LANs have a great deal of applications. Modern implementations of WLANs range from small in-home networks to large, campus-sized ones to completely mobile networks on airplanes and trains.

Users can access the Internet from WLAN hotspots in restaurants, hotels, and now with portable devices that connect to 3G or 4G networks. Oftentimes these types of public access points require no registration or password to join the network. Others can be accessed once registration has occurred or a fee is paid.

Existing Wireless LAN infrastructures can also be used to work as indoor positioning systems with no modification to the existing hardware.

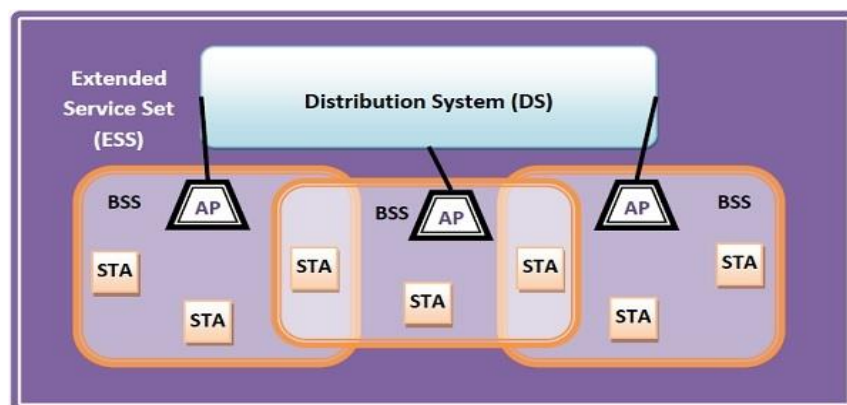
➤ IEEE 802.11

IEEE 802.11 standard, popularly known as WiFi, lays down the architecture and specifications of wireless LANs (WLANs). WiFi or WLAN uses high-frequency radio waves instead of cables for connecting the devices in LAN. Users connected by WLANs can move around within the area of network coverage.

IEEE 802.11 Architecture

The components of an IEEE 802.11 architecture are as follows –

- **Stations (STA)** – Stations comprises of all devices and equipment that are connected to the wireless LAN. A station can be of two types–
 - Wireless Access Point (WAP) – WAPs or simply access points (AP) are generally wireless routers that form the base stations or access.
 - Client. Clients are workstations, computers, printers, smartphones, etc.
- Each station has a wireless network interface controller.
- **Basic Service Set (BSS)** – A service set is a group of communicating at the physical layer level. BSS can two categories depending the mode of operation–
 - Infrastructure BSS – Here, devices communicate with other devices through access points.
 - Independent BSS – Here, the devices communicate in a peer-to-peer basis in an ad hoc manner.
- **Extended Service Set (ESS)** – It is a set of all connected BSS.
- **Distribution System (DS)** – It connects access points in ESS.



laptops,

basic
stations

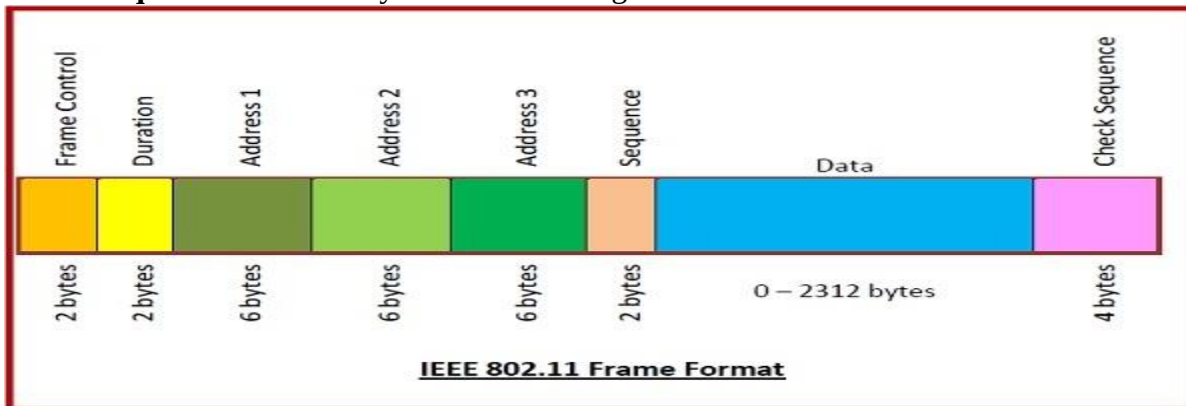
be of
upon

the

Frame Format of IEEE 802.11

The main fields of a frame of wireless LANs as laid down by IEEE 802.11 are –

- **Frame Control** – It is a 2 bytes starting field composed of 11 subfields. It contains control information of the frame.
- **Duration** – It is a 2-byte field that specifies the time period for which the frame and its acknowledgment occupy the channel.
- **Address fields** – There are three 6-byte address fields containing addresses of source, immediate destination, and final endpoint respectively.
- **Sequence** – It is a 2 bytes field that stores the frame numbers.
- **Data** – This is a variable-sized field that carries the data from the upper layers. The maximum size of the data field is 2312 bytes.
- **Check Sequence** – It is a 4-byte field containing error detection information.



➤ Bluetooth

Bluetooth technology is a high speed and low powered wireless technology designed to connect phones or other portable equipment for communication or file transmissions. Bluetooth wireless technology is a short range communications technology intended to replace the cables connecting portable unit and maintaining high levels of security. Bluetooth technology is based on **Ad-hoc technology** also known as **Ad-hoc Pico nets**, which is a local area network with a very limited coverage.

History of Bluetooth

WLAN technology enables device connectivity to infrastructure based services through a wireless carrier provider. The need for personal devices to communicate wirelessly with one another without an established infrastructure has led to the emergence of **Personal Area Networks (PANs)**.

- Ericsson's Bluetooth project in 1994 defines the standard for PANs to enable communication between mobile phones using low power and low cost radio interfaces.
- In May 1988, Companies such as IBM, Intel, Nokia and Toshiba joined Ericsson to form the Bluetooth Special Interest Group (SIG) whose aim was to develop a defacto standard for PANs.
- IEEE has approved a Bluetooth based standard named IEEE 802.15.1 for Wireless Personal Area Networks (WPANs). IEEE standard covers MAC and Physical layer applications.

Bluetooth specification details the entire protocol stack. Bluetooth employs Radio Frequency (RF) for communication. It makes use of **frequency modulation** to generate radio waves in the **ISM** band.

The usage of Bluetooth has widely increased for its special features.

- Bluetooth offers a uniform structure for a wide range of devices to connect and communicate with each other.
- Bluetooth technology has achieved global acceptance such that any Bluetooth enabled device, almost everywhere in the world, can be connected with Bluetooth enabled devices.

- Low power consumption of Bluetooth technology and an offered range of up to ten meters has paved the way for several usage models.
- Bluetooth offers interactive conference by establishing an adhoc network of laptops.
- Bluetooth usage model includes cordless computer, intercom, cordless phone and mobile phones.

Bluetooth defines two types of networks: Piconets and Scatternets

Bluetooth enabled electronic devices connect and communicate wirelessly through shortrange devices known as **Piconets**. Bluetooth devices exist in small ad-hoc configurations with the ability to act either as master or slave the specification allows a mechanism for **master** and **slave** to switch their roles. Point to point configuration with one master and one slave is the simplest configuration.

When more than two Bluetooth devices communicate with one another, this is called a **PICONET**. A Piconet can contain up to seven slaves clustered around a single master. The device that initializes establishment of the Piconet becomes the **master**.

The master is responsible for transmission control by dividing the network into a series of time slots amongst the network members, as a part of **time division multiplexing** scheme which is shown below.

The features of Piconets are as follows –

- Within a Piconet, the timing of various devices and the frequency hopping sequence of individual devices is determined by the clock and unique **48-bit address** of master.
- Each device can communicate simultaneously with up to seven other devices within a single Piconet.
- Each device can communicate with several piconets simultaneously.
- Piconets are established dynamically and automatically as Bluetooth enabled devices enter and leave piconets.
- There is no direct connection between the slaves and all the connections are essentially master-to-slave or slave-to-master.
- Slaves are allowed to transmit once these have been polled by the master.
- Transmission starts in the slave-to-master time slot immediately following a polling packet from the master.
- A device can be a member of two or more piconets, jumping from one piconet to another by adjusting the transmission regime-timing and frequency hopping sequence dictated by the master device of the second piconet.
- It can be a slave in one piconet and master in another. It however cannot be a master in more than once piconet.
- Devices resident in adjacent piconets provide a bridge to support inner-piconet connections, allowing assemblies of linked piconets to form a physically extensible communication infrastructure known as **Scatternet**.

Spectrum

Bluetooth technology operates in the unlicensed industrial, scientific and medical (ISM) band at 2.4 to 2.485 GHZ, using a spread spectrum hopping, full-duplex signal at a nominal rate of 1600 hops/sec. the 2.4 GHZ ISM band is available and unlicensed in most countries.

Range

Bluetooth operating range depends on the device Class 3 radios have a range of up to 1 meter or 3 feet Class 2 radios are most commonly found in mobile devices have a range of 10 meters or 30 feet Class 1 radios are used primarily in industrial use cases have a range of 100 meters or 300 feet.

Data rate

Bluetooth supports 1Mbps data rate for version 1.2 and 3Mbps data rate for Version 2.0 combined with **Error Data Rate**.

The Architecture of Bluetooth Technology

- In Bluetooth technology, the network of Bluetooth consists of a Personal Area Network or a

- Bluetooth's architecture is also called a "Piconet" because it is made of multiple networks.
- It contains a minimum of 2 to a maximum of 8 Bluetooth peer devices.
- It usually contains a single master and up to 7 slaves.
- Piconet provides the technology which facilitates data transmission based on its nodes, i.e., Master node and Slave Nodes.
- The master node is responsible for sending the data while the slave nodes are used to receive the data.
- In Bluetooth technology, data transmission occurs through Ultra-High frequency and short-wavelength radio waves.
- The Piconet uses the concept of multiplexing and spread spectrum. It is a combination of code division multiple access (CDMA) and frequency hopping spread spectrum (FHSS) technique.

How does Bluetooth work?

As we stated that there is one master and up to 7 slaves may exist for a Bluetooth connection. The master is the device that initiates communication with other devices. The master device handles the communications link and traffic between itself and the slave devices associated with it. The slave devices have to respond to the master device and synchronize their transmit/receive timing with the master device's specified time.



18M

294

Hell



Symbol of Bluetooth



An example of a Bluetooth device

for Beginner

**Advantages
Bluetooth**

of Technology

Following is a list of some advantages of the Bluetooth technology:

- Bluetooth Technology is based on Wireless technology. That's why it is cheap because it doesn't need any transmission wire that reduces the cost.
- It is very simple to form a Piconet in Bluetooth technology.
- It removes the problem of radio interference by using the Speed Frequency Hopping technique.
- The energy or power consumption is very low, about 0.3mW. It makes it possible for the least utilization of battery life.
- It is robust because it guarantees security at a bit level. The authentication is controlled using a 128bit key.
- You can use it for transferring the data, and verbal communication as Bluetooth can support data channels of up to 3 similar voice channels.
- It doesn't require line of sight and one to one communication as used in other modes of wireless communications such as infrared.

Disadvantages of Bluetooth Technology

Following is a list of some disadvantages of the Bluetooth technology:

- In Bluetooth technology, the bandwidth is low.
- The data transmission range may also be an issue because it is also less.

Applications of Bluetooth Technology

Bluetooth technology is used in many communicational and entertainment devices. The following are some most used applications of the Bluetooth technology:

- Bluetooth technology is used in cordless desktop. It means the peripheral devices such as a mouse, keyboard, printer, speakers, etc. are connected to the desktop without a wire.

- It



is used in the multimedia transfer, such as exchanging multimedia data like songs, videos, pictures etc. that can be transferred among devices using Bluetooth.

- This technology is also used in the following devices: i.e.
 - Bluetooth Speakers.
 - Bluetooth Headphones.

- Bluetooth Headsets for calling purposes.
- Bluetooth gaming consoles etc.

➤ Connecting LANs

LANs do not normally operate in isolation. They are connected to one another or to the Internet. To connect LANs, or segments of LANs, we use connecting devices. Connecting devices can operate in different layers of the Internet model.

➤ CONNECTING DEVICES

In this section, we divide connecting devices into five different categories based on the layer in which they operate in a network, as shown in Figure 15.1. The five categories contain devices which can be defined as in Table 1:

Table 1: Connecting Devices

S. No.	Device	Layer(s)
1.	Passive hub	below the physical layer
2.	Repeater/Active hub	at the physical layer
3.	Bridge/Two-layer switch	at the physical and data link layers
4.	Router/Three-layer switch	at the physical, data link, and network layers
5.	Gateway	at all five layers

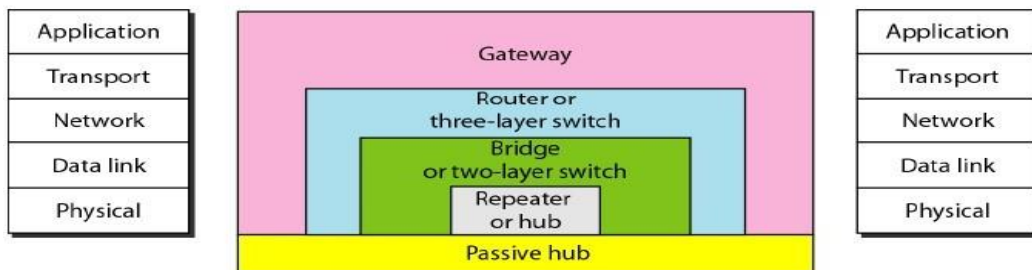


Figure Five categories of connecting devices

Passive Hubs

A passive hub is just a **connector**. It connects the wires coming from different branches. In a star-topology Ethernet LAN, a passive hub is just a point where the signals coming from different stations collide; the hub is the collision point. This type of a hub is part of the media; its location in the Internet model is below the physical layer.

Repeaters

A repeater is a device that operates only in the physical layer. Signals that carry information within a network can travel a fixed distance before attenuation endangers the integrity of the data. A repeater receives a signal and, before it becomes too weak or corrupted, regenerates the original bit pattern. The repeater then sends the refreshed signal. A repeater can extend the physical length of a LAN, as shown in Figure 15.2.

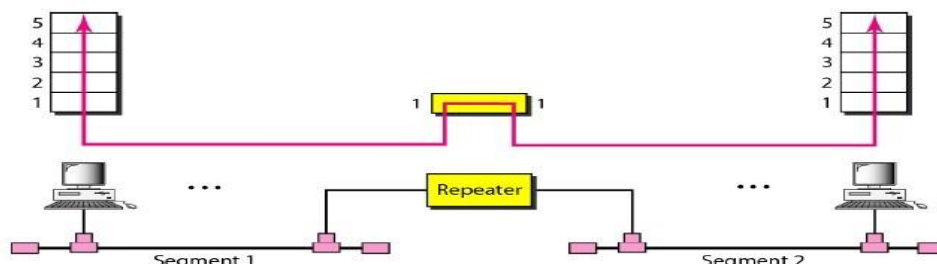


Figure 15.2 A repeater connecting two segments of a LAN

- A repeater does not actually connect two LANs; it connects two segments of the same LAN. The segments connected are still part of one single LAN. A repeater is not a device that can connect two LANs of different protocols.
- A repeater can overcome the 10Base5 Ethernet length restriction. In this standard, the length of the cable is limited to 500 m. To extend this length, we divide the cable into segments and install repeaters between segments. Note that the whole network is still considered one LAN, but the portions of the network separated by repeaters are called segments.
- The repeater acts as a two-port node, but operates only in the physical layer. When it receives a frame from any of the ports, it regenerates and forwards it to the other port.
- A repeater forwards every frame; it has no filtering capability.

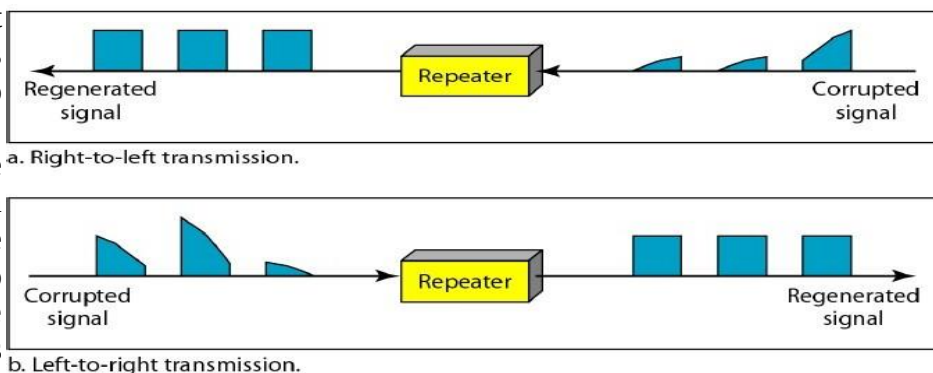


Figure Function of a repeater

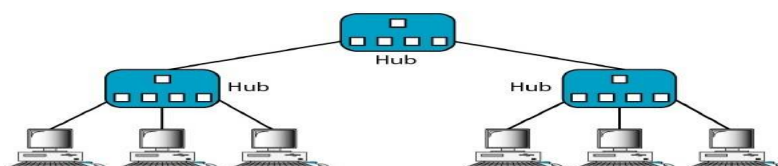
Difference between repeater and amplifier

- An amplifier cannot discriminate between the intended signal and noise; it amplifies equally everything fed into it.
- A repeater does not amplify the signal; it regenerates the signal. When it receives a weakened or corrupted signal, it creates a copy, bit for bit, at the original strength.
- Thus A repeater is a regenerator, not an amplifier.

Active Hubs

An active hub is actually a multipart repeater. It is normally used to create connections between stations in a physical star topology. We have seen examples of hubs in some Ethernet implementations (10Base-T, for example). However, hubs can also be used to create multiple levels of hierarchy, as shown in Figure 15.4. The hierarchical use of hubs removes the length limitation of 10Base-T (100 m).

Figure A hierarchy of hubs



Bridges

A bridge operates in both the physical and the data link layer. As a physical layer device, it regenerates the signal it receives. As a data link layer device, the bridge can check the physical (MAC) addresses (source and destination) contained in the frame.

What is the difference in functionality between a bridge and a repeater?

A bridge has **filtering capability**. It can check the destination address of a frame and decide if the frame should be forwarded or dropped. If the frame is to be forwarded, the decision must specify the port. A bridge has a table that maps addresses to ports. Let us give an example. In Figure 15.5, two LANs are connected by a bridge. If a frame destined for station 712B13456142 arrives at port 1, the bridge consults its table to find the departing port. According to its table, frames for 71:2B:13:45:61:42 leave through port 1; therefore, there is no need for forwarding, and the frame is dropped. On the other hand, if a frame for 71:2B:13:45:61:41 arrives at port 2, the departing port is port 1 and the frame is forwarded. In the first case, LAN 2 remains free of traffic; in the second case, both LANs have traffic. In our example, we show a two-port bridge; in reality a bridge usually has more ports.

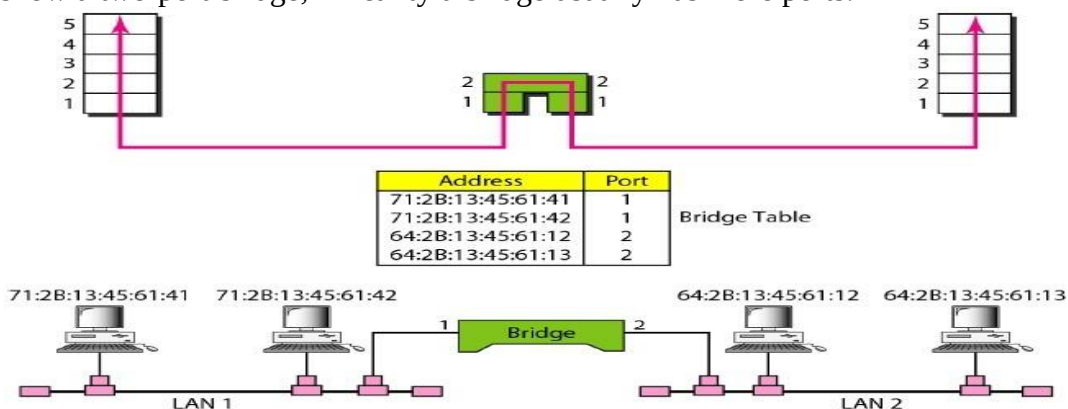


Figure A bridge connecting two LANs

Note: A bridge does not change the physical addresses contained in the frame.

Transparent Bridges

A transparent bridge is a bridge in which the stations are completely unaware of the bridge's existence. If a bridge is added or deleted from the system, reconfiguration of the stations is unnecessary. According to the IEEE 802.1d specification, a system equipped with transparent bridges must meet three criteria:

- Frames must be forwarded from one station to another.
- The forwarding table is automatically made by learning frame movements in the network.
- Loops in the system must be prevented.

Learning

The earliest bridges had forwarding tables that were **static**. The systems administrator would manually enter each table entry during bridge setup. Although the process was simple, it was not practical. If a station was added or deleted, the table had to be modified manually. The same was true if a station's MAC address changed, which is not a rare event. For example, putting in a new network card means a new MAC address.

A better solution to the static table is a **dynamic table** that maps addresses to ports automatically. To make a table dynamic, we need a bridge that gradually learns from the frame movements. To do this, the bridge inspects both the destination and the source addresses. The destination address is used for the forwarding decision (table lookup); the source address is used for adding entries to the table and for updating purposes. Let us elaborate on this process by using Figure

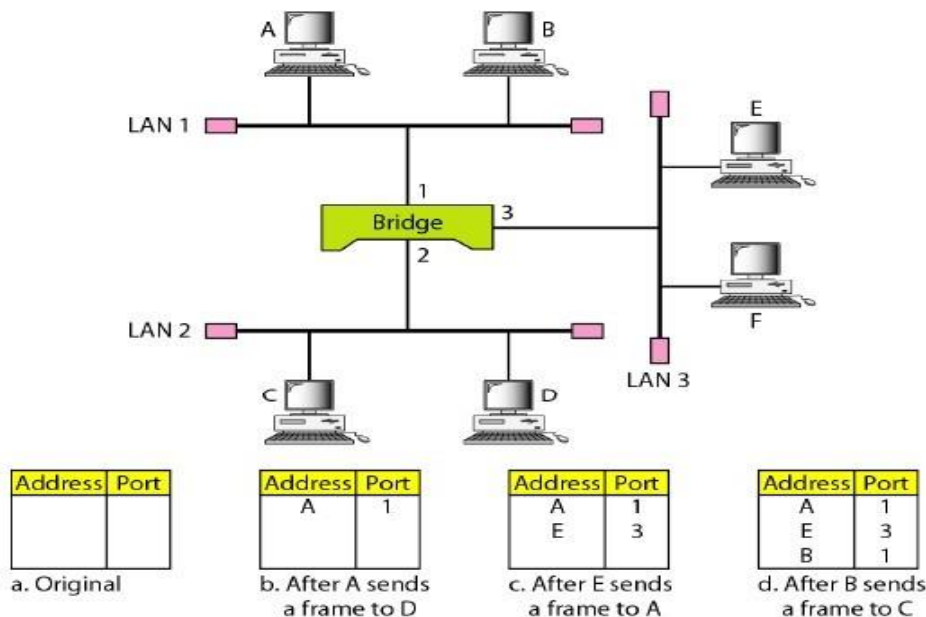


Figure A learning bridge and the process of learning

1. When station A sends a frame to station D, the bridge does not have an entry for either D or A. The frame goes out from all three ports; the frame floods the network. However, by looking at the source address, the bridge learns that station A must be located on the LAN connected to port 1. This means that frames destined for A, in the future, must be sent out through port 1. The bridge adds this entry to its table. The table has its first entry now.
2. When station E sends a frame to station A, the bridge has an entry for A, so it forwards the frame only to port 1. There is no flooding. In addition, it uses the source address of the frame, E, to add a second entry to the table.
3. When station B sends a frame to C, the bridge has no entry for C, so once again it floods the network and adds one more entry to the table.
4. The process of learning continues as the bridge forwards frames.

Loop Problem

Transparent bridges work fine as long as there are no redundant bridges in the system. Systems administrators, however, like to have redundant bridges (more than one bridge between a pair of LANs) to make the system more reliable. If a bridge fails, another bridge takes over until the failed one is repaired or replaced.

Redundancy can create loops in the system, which is very undesirable. Figure 15.7 shows a very simple example of a loop created in a system with two LANs connected by two bridges.

1. Station A sends a frame to station D. The tables of both bridges are empty. Both forward the frame and update their tables based on the source address A.
2. Now there are two copies of the frame on LAN 2. The copy sent out by bridge 1 is received by bridge 2, which does not have any information about the destination address D; it floods the bridge. The copy sent out by bridge 2 is received by bridge 1 and is sent out for lack of information about D. Note that each frame is handled separately because bridges, as two nodes on a network sharing the medium, use an access method such as CSMA/CD. The tables of both bridges are updated, but still there is no information for destination D.
3. Now there are two copies of the frame on LAN 1. Step 2 is repeated, and both copies flood the network.
4. The process continues on and on. Note that bridges are also repeaters and regenerate frames. So in each

iteration, there are newly generated fresh copies of the frames.

Solution of Loop Problem

To solve the looping problem, the IEEE specification requires that bridges use the spanning tree algorithm to create a loopless topology.

Spanning Tree

In graph theory, a **spanning tree** is a graph in which there is no loop. In a bridged LAN, this means creating a topology in which each LAN can be reached from any other LAN through one path only (no loop). We cannot change the physical topology of the system because of physical connections between cables and bridges, but we can create a logical topology that overlays the physical one. Figure 15.8 shows a system with four LANs and five bridges.

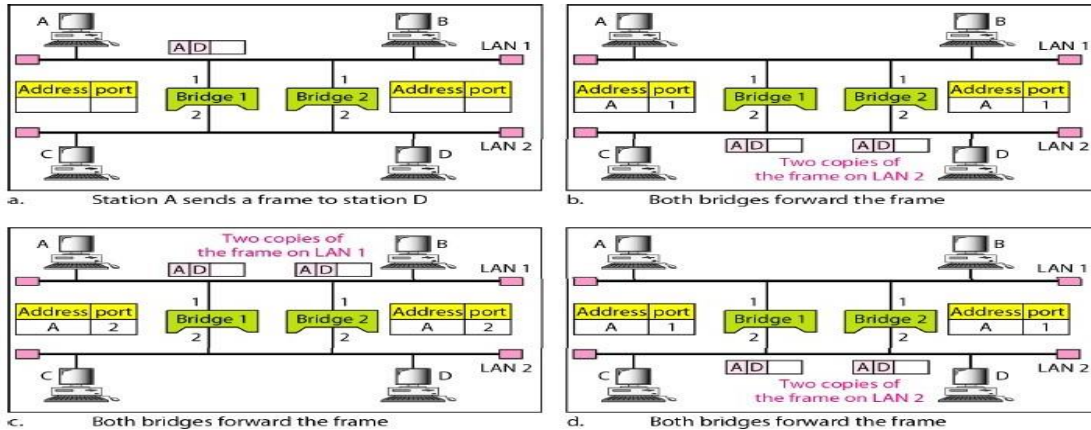
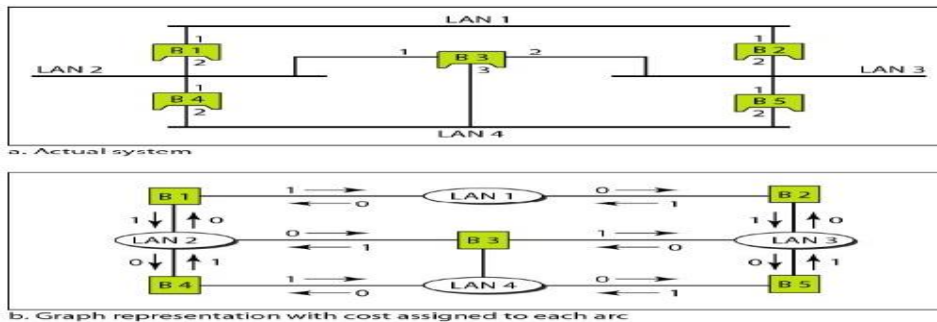


Figure Loop problem in a learning bridge

We have shown the physical system and its representation in graph theory. We have shown both LANs and bridges as nodes. The connecting arcs show the connection of a LAN to a bridge and vice versa.



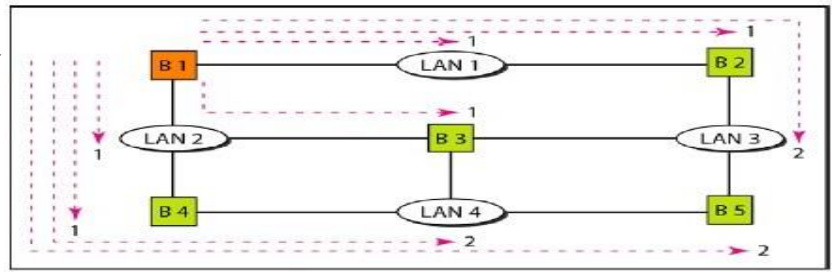
- To find the spanning tree, we need to assign a cost (metric) to each arc. The interpretation of the cost is left up to the systems administrator.
- It may be the path with minimum hops (nodes), the path with minimum delay, or the path with maximum bandwidth.
- If two ports have the same shortest value, the systems administrator just chooses one. We have chosen the minimum hops. However, as we will see in Chapter 22, the hop count is normally 1 from a bridge to the LAN and 0 in the reverse direction.

Figure A system of connected LANs and its graph representation The

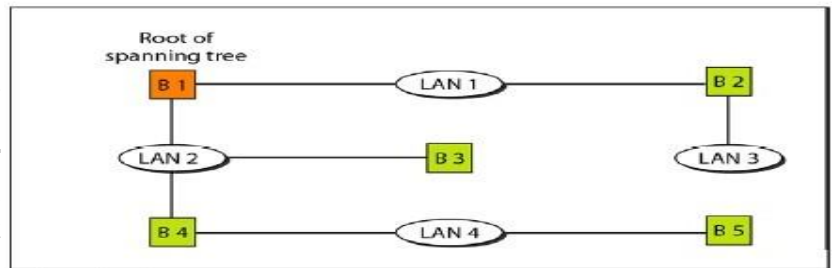
process to find the spanning

tree involves three steps:

1. Every bridge has a built-in ID (normally the serial number, which is unique). Each bridge broadcasts this ID so that all bridges know which one has the smallest ID. The bridge with the smallest ID is selected as the **root bridge** (root of the tree). We assume that bridge B1 has the smallest ID. It is, therefore, selected as the root bridge.
2. The algorithm tries to find the shortest path (a path with the shortest cost) from the root bridge to every other bridge or LAN. The shortest path can be found by examining the total cost from the root bridge to the destination.
3. The combination of the shortest paths creates the shortest tree, which is also shown in Figure Based on the spanning tree, we mark the ports that are part of the spanning tree, the forwarding ports, which forward a frame that the bridge receives. We also mark those ports that are not part of the spanning tree, the blocking ports, which block the frames received by the bridge. Figure 15.10 shows the physical systems of LANs with forwarding points (solid lines) and blocking ports (broken lines).



a. Shortest paths



b. Spanning tree

Figure Finding the shortest paths and the spanning tree in a system of bridges

Note that there is only one single path from any LAN to any other LAN in the spanning tree system. This means there is only one single path from one LAN to any other LAN. No loops are created. You can prove to yourself that there is only one path from LAN 1 to LAN 2, LAN 3, or LAN 4. Similarly, there is only one path from LAN 2 to LAN 1, LAN 3, and LAN 4. The same is true for LAN 3 and LAN 4.

Dynamic Algorithm

Each bridge is equipped with a software package that carries out this process dynamically. The bridges send special messages to one another, called bridge protocol data units (BPDUs), to update the spanning tree. The spanning tree is updated when there is a change in the system such as a failure of a bridge or an addition or deletion of bridges.

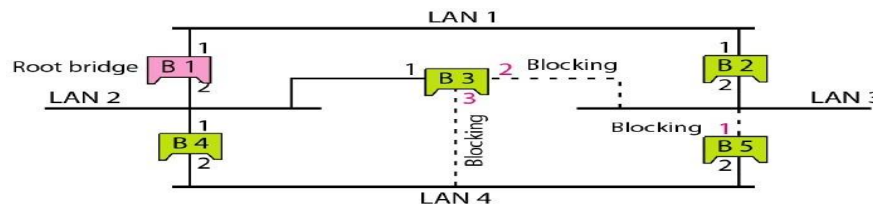


Figure Forwarding and blocking ports after using spanning tree algorithm

Ports 2 and 3 of bridge B3 are blocking ports (no frame is sent out of these ports). Port 1 of bridge B5 is also a blocking port (no frame is sent out of this port).

Switches

When we use the term **switch**, we must be careful because a switch can mean two different things. We must clarify the term by adding the level at which the device operates. We can have:

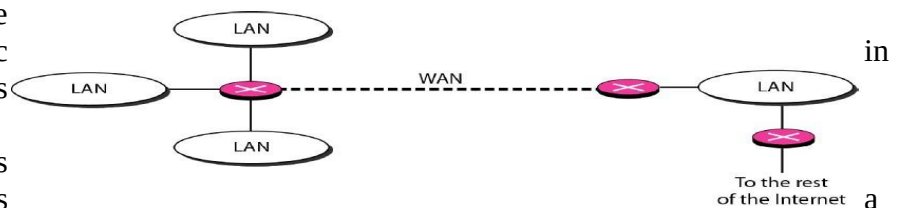
- **Two-layer switch:** performs at the physical and data link layers.
- **Three-layer switch.** is used at the network layer; it is a kind of router.

A **two-layer switch** is a **bridge**, a bridge with many ports and a design that allows better (faster) performance. A bridge with a few ports can connect a few LANs together. A bridge with many ports may be able to allocate a unique port to each station, with each station on its own independent entity. This means no competing traffic (no collision, as we saw in Ethernet).

- A two-layer switch, as a bridge does, makes a filtering decision based on the MAC address of the frame it received.
- However, a two-layer switch can be more sophisticated. It can have a buffer to hold the frames for processing.
- It can have a switching factor that forwards the frames faster. Some new two-layer switches, called cut-through switches, have been designed to forward the frame as soon as they check the MAC addresses in the header of the frame.

A **three-layer switch** is a **Router**, a router is a three-layer device that routes packets based on their logical addresses (host-to-host addressing).

- Router is faster and more sophisticated. The switching fabric a router(three-layer switch) allows faster table lookup and forwarding.
- A router normally connects LANs and WANs in the Internet and has routing table that is used for making decisions about the route.



• The routing tables are normally dynamic and are updated using routing protocols. Figure 15.11 shows a part of the Internet that uses routers to connect LANs and WANs.

Figure Routers connecting independent LANs and WANs

Gateway

Although some textbooks use the terms gateway and router interchangeably, most of the literature distinguishes between the two.

- A gateway is normally a computer that **operates in all five layers of the Internet or seven layers of OSI model.**
- A gateway takes an application message, reads it, and interprets it.
- This means that it can be used as a connecting device between two internetworks that use different models. For example, a network designed to use the OSI model can be connected to another network using the Internet model.
- The gateway connecting the two systems can take a frame as it arrives from the first system, move it up to the OSI application layer, and remove the message.
- Gateways can provide security. The gateway is also used to filter unwanted application-layer messages.

➤ BACKBONE NETWORKS

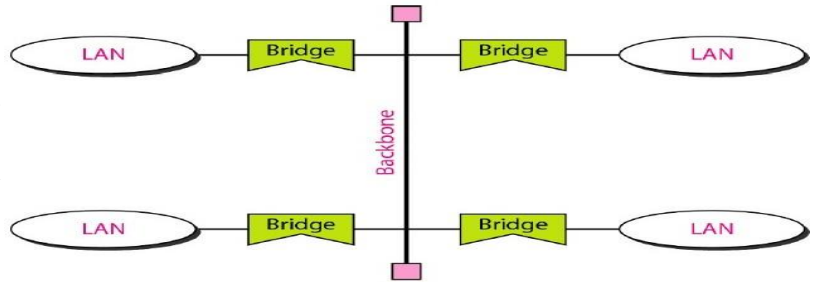
A backbone network allows several LANs to be connected using some connecting devices.

- In a backbone network, no station is directly connected to the backbone; the stations are part of a LAN, and the backbone connects the LANs.
- The backbone is itself a LAN that uses a LAN protocol such as Ethernet; each connection to the backbone is itself another LAN.
- Although many different architectures can be used for a backbone, the two most common are: the bus and the star.

Bus Backbone

In a bus backbone, the topology of the backbone is a bus. The backbone itself can use one of the protocols that support a bus topology such as 10Base5 or 10Base2.

- Bus backbones are normally used as a distribution backbone to connect different buildings in an organization. Each building can comprise either a single LAN or another backbone (normally a star backbone).
- A good example of a bus backbone is one that connects single or multiple-floor buildings on a campus. Each single-floor building usually has a single LAN. Each multiple-floor building has a backbone (usually a star) that connects each LAN on a floor.
- A bus backbone can interconnect these LANs and backbones. Figure shows an example of a bridge-based backbone with four LANs.

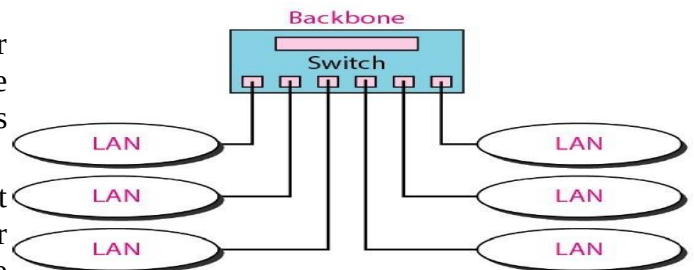


In Figure if a station in a LAN needs to send a frame to another station in the same LAN, the corresponding bridge blocks the frame; the frame never reaches the backbone. However, if a station needs to send a frame to a station in another LAN, the bridge passes the frame to the backbone, which is received by the appropriate bridge and is delivered to the destination LAN.

Each bridge connected to the backbone has a table that shows the stations on the LAN side of the bridge. The blocking or delivery of a frame is based on the contents of this table.

Star Backbone

In a star backbone, sometimes called a collapsed or switched backbone, the topology of the backbone is a star. In this configuration, the backbone is just one switch (that is why it is called, erroneously, a collapsed backbone) that connects the LANs. Figure 15.13 shows a star backbone. Note that, in this configuration, the switch does the job of the backbone and at the same time connects the LANs.



- Star backbones are mostly used as a distribution backbone inside a building. In a multifloor building, we usually find one LAN that serves each particular floor. A star backbone connects these LANs.
- The backbone network, which is just a switch, can be installed in the basement or the first floor, and separate cables can run from the switch to each LAN.
- If the individual LANs have a physical star topology, either the hubs (or switches) can be installed in a closet on the corresponding floor, or all can be installed close to the switch.
- We often find a rack or chassis in the basement where the backbone switch and all hubs or switches are installed.

Connecting Remote LANs

Another common application for a backbone network is to connect remote LANs. This type of backbone network is useful when a company has several offices with LANs and needs to connect them.

- The connection can be done through bridges, sometimes called remote bridges.
- The bridges act as connecting devices connecting LANs and point-to-point networks, such as leased telephone lines or ADSL lines.

- The point-to-point network in this case is considered a LAN without stations. The point-to-point link can use a protocol such as PPP.

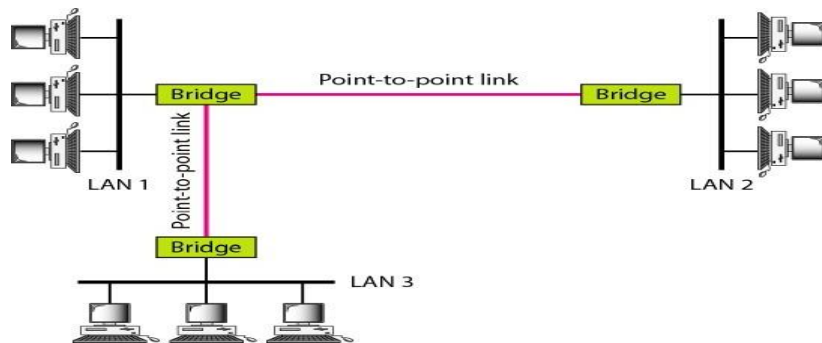


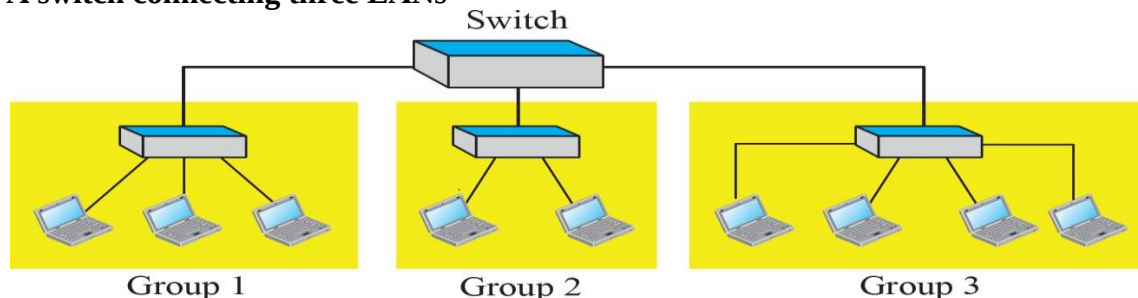
Figure Connecting remote LANs with bridges

➤ Virtual LANs

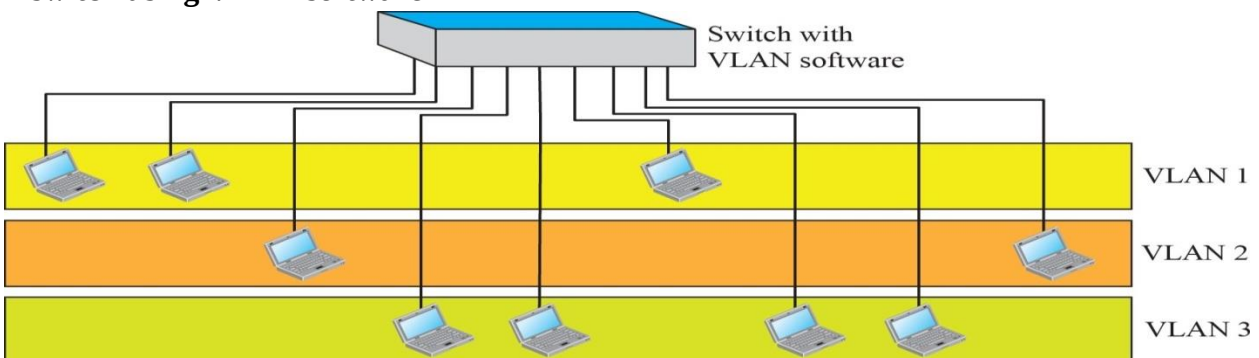
VLANs have the primary role to enable easier configuration and management of large corporate networks built around many bridges. There are several implementation strategies for these virtual networks. Virtual LAN is software that is employed to provide multiple networks in single hub by grouping terminals connected to switching hubs. It is a LANs that is grouped together by logical addresses into a virtual LAN instead of a physical LAN through a switch. The switch can support many virtual LANs that operate with having different network addresses or as subnets. Users within a virtual LAN are grouped either by IP address or by port address, with each node attached to the switch via a dedicated circuit. Users also can be assigned to more than one virtual LAN.

The VLAN can be defined as a broadcast domain in which the broadcast address reaches all stations belonging to the VLAN. Communications within the VLAN can be secured, and between those two controlled separate VLANs

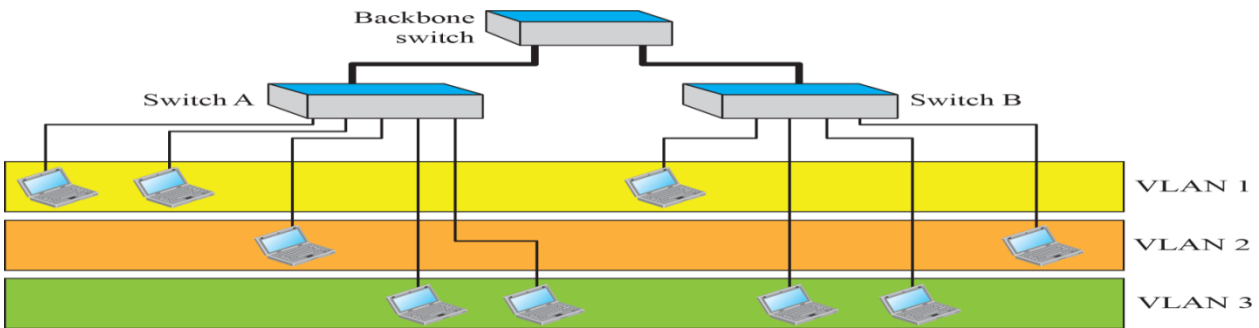
A switch connecting three LANs



A switch using VLAN software



Two switches in a backbone using VLAN software



Membership

What characteristic can be used to group stations in a VLAN?

- Vendors use different characteristics such as interface numbers, port numbers, MAC addresses, IP addresses, IP multicast addresses, or a combination of two or more of these.

Port Numbers

Some VLAN vendors use switch port numbers as a membership characteristic. For example, the administrator can define that stations connecting to ports 1, 2, 3, and 7 belong to VLAN 1; stations connecting to ports 4, 10, and 12 belong to VLAN 2; and so on

MAC Addresses

Some VLAN vendors use the 48-bit MAC address as a membership characteristic. For example, the administrator can stipulate that stations having MAC addresses E21342A12334 and F2A123BCD341 belong to VLAN 1.

IP Addresses Some VLAN vendors use the 32-bit IP address as a membership characteristic. For example, the administrator can stipulate that stations having IP addresses 181.34.23.67, 181.34.23.72, 181.34.23.98, and 181.34.23.112 belong to VLAN 1.

Multicast IP Addresses

Some VLAN vendors use the multicast IP address as a membership characteristic. Multicasting at the IP layer is now translated to multicasting at the data link layer.

Combination

Recently, the software available from some vendors allows all these characteristics to be combined. The administrator can choose one or more characteristics when installing the software. In addition, the software can be reconfigured to change the settings.

Configuration

How are the stations grouped into different VLANs? Stations are configured in one of three ways: manual, semiautomatic, and automatic.

Manual Configuration

In a manual configuration, the network administrator uses the VLAN software to manually assign the stations into different VLANs at setup. Later migration from one VLAN to another is also done manually. Note that this is not a physical configuration; it is a logical configuration. The term manually here means that the administrator types the port numbers, the IP addresses, or other characteristics, using the VLAN software.

Automatic Configuration

In an automatic configuration, the stations are automatically connected or disconnected from a VLAN using criteria defined by the administrator. For example, the administrator can define the project number as the criterion for being a member of a group. When a user changes the project, he or she automatically migrates to a new VLAN.

Semiautomatic Configuration

A semiautomatic configuration is somewhere between a manual configuration and an automatic configuration. Usually, the initializing is done manually, with migrations done automatically.

Communication Between Switches

In a multiswitched backbone, each switch must know not only which station belongs to which VLAN, but also the membership of stations connected to other switches. For example, in Figure 15.17, switch A must know the membership status of stations connected to switch B, and switch B must know the same about switch A. Three methods have been devised for this purpose: table maintenance, frame tagging, and time-division multiplexing.

Frame Tagging

In this method, when a frame is traveling between switches, an extra header is added to the MAC frame to define the destination VLAN. The frame tag is used by the receiving switches to determine the VLANs to be receiving the broadcast message. Time-Division Multiplexing (TDM)

In this method, the connection (trunk) between switches is divided into timeshared channels. For example, if the total number of VLANs in a backbone is five, each trunk is divided into five channels. The traffic destined for VLAN 1 travels in channel 1, the traffic destined for VLAN 2 travels in channel 2, and so on. The receiving switch determines the destination VLAN by checking the channel from which the frame arrived.

Advantages

There are several advantages to using VLANs.

Cost and Time Reduction VLANs can reduce the migration cost of stations going from one group to another. Physical reconfiguration takes time and is costly. Instead of physically moving one station to another segment or even to another switch, it is much easier and quicker to move it by using software.

Creating Virtual Work Groups VLANs can be used to create virtual work groups. For example, in a campus environment, professors working on the same project can send broadcast messages to one another without the necessity of belonging to the same department. This can reduce traffic if the multicasting capability of IP was previously used.

Security VLANs provide an extra measure of security. People belonging to the same group can send broadcast messages with the guaranteed assurance that users in other groups will not receive these messages.

➤ 2G,3G,4G,5G Wireless technologies

Simply, the "G" stands for "GENERATION". While connected to the internet, the speed of the connection depends upon the signal strength that is shown in abbreviations like 2G, 3G, 4G, 5G, etc. on any mobile device. Each generation of wireless broadband is defined as a set of telephone network standards that describe the technological implementation of the system.

The aim of wireless communication is to provide high quality, reliable communication just like wired communication and each new generation represents a big leap in that direction. Mobile communication has become more popular in the last few years due to fast reform in mobile technology. For the comparison of 2G, 3G, 4G, and 5G we first need to understand the key features of all these technologies.

SECOND GENERATION (2G)

- 2G refers to the second generation of mobile networks based on GSM. The radio signals used by the 1G network were analog, while 2G networks were digital. 2G capabilities were achieved by allowing multiple users on a single channel via multiplexing. During 2G, cellular phones were used for data along with voice. Some of the key features of 2G were:
 - Data speeds of up to 64 kbps

- Use of digital signals instead of analog
- Enabled services such as SMS and MMS (Multimedia Message)
- Provided better quality voice calls
- It used a bandwidth of 30 to 200 KHz

- **THIRD GENERATION (3G)**

- The 3G standard utilises Universal Mobile Telecommunications System (UMTS) as its core network architecture. 3G network combines aspects of the 2G network with new technologies and protocols to deliver a significantly faster data rate. By using packet switching, the original technology was improved to allow speeds up to 14 Mbps. It used Wide Band Wireless Network that increased clarity. It operates at a range of 2100 MHz and has a bandwidth of 15-20 MHz. Some of the main features of 3G are:

- Speed of up to 2 Mbps
- Increased bandwidth and data transfer rates
- Send/receive large email messages
- Large capacities and broadband capabilities

International Mobile Telecommunications-2000 (IMT-2000) were the specifications by the International Telecommunication Union for the 3G network; theoretically, 21.6 Mbps is the max speed of HSPA+.

FOURTH GENERATION (4G)

- The main difference between 3G and 4G is the data rate. There is also a huge difference between 3G and 4G technology. The key technologies that have made 4G possible are MIMO (Multiple Input Multiple Output) and OFDM (Orthogonal Frequency Division Multiplexing). The most important 4G standards are WiMAX and LTE. While 4G LTE is a major improvement over 3G speeds, it is technically not 4G. What is the difference between 4G and LTE?
- Even after it was widely available, many networks were not up to the required speed of 4G. 4G LTE is a “fourth generation long term evolution”, capable of delivering a very fast and secure internet connection. Basically, 4G is the predetermined standard for mobile network connections. 4G LTE is the term given to the path which has to be followed to achieve those predefined standards. Some of the features of 4G LTE are:
 - Support interactive multimedia, voice, video.
 - High speed, high capacity and low cost per bit (Speeds of up to 20 Mbps or more.)
 - Global and scalable mobile networks.
 - Ad hoc and multi-hop networks.

Following is the comparison between 4G and 5G speeds:

We conducted a comparison test campaign of 4G and 5G with our RantCell app. The result displays 4G and 5G data points with peak throughput in a particular location for ‘23415’ mobile network operator i.e. Mobile network code for Vodafone as shown in the image below:

FIFTH GENERATION (5G)

- 5G networks operate on rarely used radio millimeter bands in the 30 GHz to 300 GHz range. Testing of 5G range in mmWave has produced results approximately 500 meters from the tower. Using small cells, the deployment of 5G with millimetre wave based carriers can improve overall coverage area. Combined with beamforming, small cells can deliver extremely fast coverage with low latency.

Comparison	2G	3G	4G	5G
Introduced in year	1993	2001	2009	2018
Technology	GSM	WCDMA	LTE, WiMAX	MIMO, mm Waves
Access system	TDMA, CDMA	CDMA	CDMA	OFDM, BDMA
Switching type	Circuit switching for voice and packet switching for data	Packet switching except for air interference	Packet switching	Packet switching
Internet service	Narrowband	Broadband	Ultra broadband	Wireless World Wide Web
Bandwidth	25 MHz	25 MHz	100 MHz	30 GHz to 300 GHz
Advantage	Multimedia features (SMS, MMS), internet access and SIM introduced	High security, international roaming	Speed, high speed handoffs, global mobility	Extremely high speeds, low latency
Applications	Voice calls, short messages	Video conferencing, mobile TV, GPS	High speed applications, mobile TV, wearable devices	High resolution video streaming, remote control of vehicles, robots, and medical procedures

- Low latency is one of 5G's most important features. 5G uses a scalable orthogonal frequency-division multiplexing (OFDM) framework. 5G benefits greatly from this and can have latency as low as one millisecond with realistic estimates to be around 1 – 10 seconds. 5G is estimated to be 60 to 120 times faster than the average 4G latency.
- Active antenna 5G encapsulated with 5G massive MIMO is used for providing better connections and enhanced user experience. Big 5G array antennas are deployed to gain additional beamforming information and knock out propagation challenges that are experienced at mmWave frequency ranges.
- Further, 5G networks clubbed with network slicing architecture enables telecom operators to offer on-demand tailored connectivity to their users that is adhered to Service Level Agreement (SLA). Such customised network capabilities comprise latency, data speed, latency, reliability, quality, services, and security. With speeds of up to 10 Gbps, 5G is set to be as much as 10 times faster than 4G. Following is a brief comparison of 2G, 3G, 4G, and 5G. Each generation in some way has improved over its predecessor. There is a lot of ground to compare the cell networks over. Following is the comparison between 2G, 3G, 4G, 5G. The comparison of 2G, 3G, 4G, and 5G clearly shows the differences in the technologies. The comparison of 2G, 3G, 4G, and 5G also makes it evident that 5G is going to be one of the most ambitious leaps in the history of cell network technologies.

➤ SATELLITE NETWORKS

A satellite network is a combination of nodes, some of which are satellites, that provides communication from one point on the Earth to another. A node in the network can be a satellite, an Earth station, or an end-user terminal or telephone.

Satellite networks are like cellular networks in that they divide the planet into cells. Satellites can provide transmission capability to and from any location on Earth, no matter how remote. This advantage makes high-quality communication available to undeveloped parts of the world without

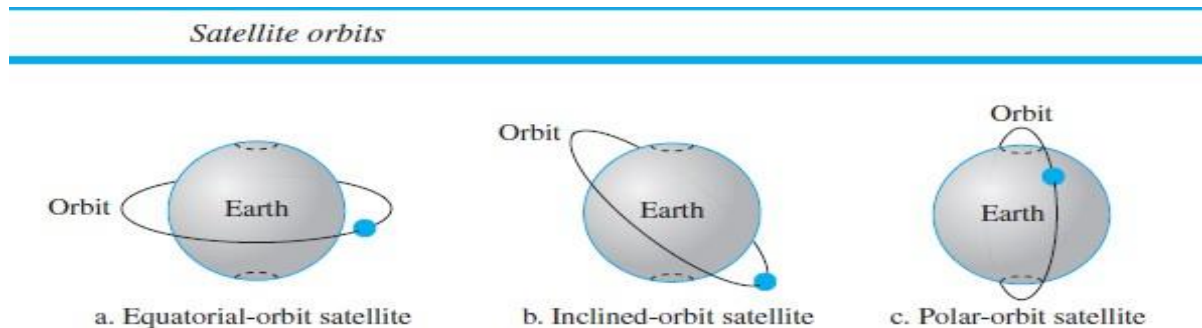
requiring a huge investment in ground-based infrastructure.

*Operation

Some general issues related to the operation of satellites.

*Orbits

An artificial satellite needs to have an orbit, the path in which it travels around the Earth. The orbit can be equatorial, inclined, or polar.



*Footprint

Satellites process microwaves with bidirectional antennas (line-of-sight). Therefore, the signal from a satellite is normally aimed at a specific area called the footprint. The signal power at the center of the footprint is maximum. The power decreases as we move out from the footprint center.

*Frequency Bands for Satellite Communication

The frequencies reserved for satellite microwave communication are in the gigahertz (GHz) range. Each satellite sends and receives over two different bands. Transmission from the Earth to the satellite is called the uplink. Transmission from the satellite to the Earth is called the downlink.

- **Three Categories of Satellites**

Based on the location of the orbit, satellites can be divided into three categories: geostationary Earth orbit (GEO), low-Earth-orbit (LEO), and medium-Earth-orbit (MEO).

- **GEO Satellites**

Line-of-sight propagation requires that the sending and receiving antennas be locked onto each other's location at all times (one antenna must have the other in sight). For this reason, a satellite that moves faster or slower than the Earth's rotation is useful only for short periods. To ensure constant communication, the satellite must move at the same

speed as the Earth so that it seems to remain fixed above a certain spot. Such satellites are called geostationary.

Because orbital speed is based on the distance from the planet, only one orbit can be geostationary. This orbit occurs at the equatorial plane and is approximately 22,000 mi from the surface of the Earth.

- **MEO Satellites**

Medium-Earth-orbit (MEO) satellites are positioned between the two Van Allen belts. A satellite at this orbit takes approximately 6 to 8 hours to circle the Earth.

Global Positioning System

One example of a MEO satellite system is the Global Positioning System (GPS), contracted and operated by the U.S. Department of Defense, orbiting at an altitude about 18,000 km (11,000 mi) above the Earth. The system consists of 24 satellites and is used for land, sea, and air navigation to provide time and location for vehicles and ships.

- **LEO Satellites**

Low-Earth-orbit (LEO) satellites have polar orbits. The altitude is between 500 and 2000 km, with a rotation period of 90 to 120 min. The satellite has a speed of 20,000 to 25,000 km/h. A LEO system usually has a cellular type of access, similar to the cellular telephone system. The footprint normally has a diameter of 8000 km. Because LEO satellites are close to Earth, the round-trip time propagation delay is normally less than 20 ms, which is acceptable for audio communication.

A LEO system is made of a constellation of satellites that work together as a network; each satellite acts as a switch. Satellites that are close to each other are connected through intersatellite links (ISLs). A mobile system communicates with the satellite through a user mobile link (UML). A satellite can also communicate with an Earth station (gateway) through a gateway link (GWL).

LEO satellite system



➤ **Virtual circuit switching**

Two common WAN technologies use virtual-circuit switching. Frame Relay is a relatively high-speed protocol that can provide some services not available in other WAN technologies such as DSL, cable TV, and T lines. ATM, as a high-speed protocol, can be the superhighway of communication when it deploys physical layer carriers such as SONET.

Frame Relay:

Frame Relay is a virtual-circuit wide-area network that was designed in response to demands for a new type of WAN.

1. Prior to Frame Relay, some organizations were using a virtual-circuit switching network called X.25 that performed switching at the network layer.

Drawbacks of X.25:

1. X.25 has a low 64-kbps data rate. There was a need for higher data-rate WANs.
2. X.25 has extensive flow and error control at both the data link layer and the network layer.
3. Originally X.25 was designed for private use, not for the Internet.
4. Disappointed with X.25, some organizations started their own private WAN by leasing T-1 or T-3 lines from public service providers.

Drawbacks:

a) If an organization has n branches spread over an area, it needs $n(n - 1)/2$ T-1 or T-3 lines. The organization pays for all these lines although it may use the lines only 10 percent of the time. This can be very costly:

b) The services provided by T-1 and T-3 lines assume that the user has fixed-rate data all the time. For example, a T-1 line is designed for a user who wants to use the line at a consistent 1.544 Mbps. This type of service is not suitable for the many users today that need to send bursty data.

In response to the above drawbacks, Frame Relay was designed.

Features:

- a) Frame Relay operates at a higher speed (1.544 Mbps and recently 44.376 Mbps). This means that it can easily be used instead of a mesh of T-1 or T-3 lines.
- b) Frame Relay operates in just the physical and data link layers. This means it can easily be used as a backbone network to provide services to protocols that already have a network layer protocol, such as the Internet.
- c) Frame Relay allows bursty data.
- d) Frame Relay allows a frame size of 9000 bytes, which can accommodate all local area network frame sizes.
- e) Frame Relay is less expensive than other traditional WANs.
- f) Frame Relay has error detection at the data link layer only. There is no flow control or error control.

1. Architecture

Frame Relay provides permanent virtual circuits and switched virtual circuits. The routers are used, to connect LANs and WANs in the Internet. In the figure, the Frame Relay WAN is used as one link in the global Internet.

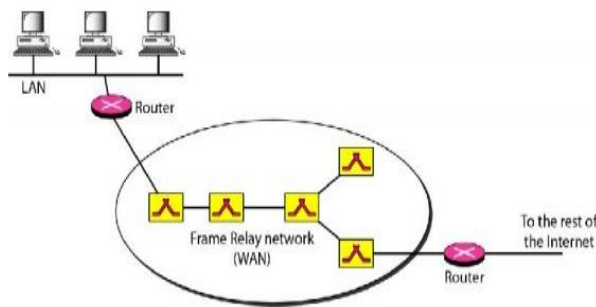


Figure 2.86 Frame Relay network

Virtual Circuits

Frame Relay is a virtual circuit network. A virtual circuit in Frame Relay is identified by a number called a data link connection identifier (DLCI). **VCIs in Frame Relay are called DLCIs.**

Switches

Each switch in a Frame Relay network has a table to route frames. The table matches an incoming port-DLCI combination with an outgoing port-DLCI combination. The only difference is that

VCIs are replaced by DLCIs.

2. Frame Relay Layers

Frame Relay has only physical and data link layers.

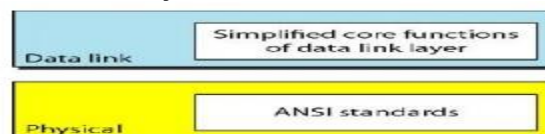


Figure 2.87 Frame Relay layers

Physical Layer

No specific protocol is defined for the physical layer in Frame Relay. Instead, it is left to the implementer to use whatever is available. Frame Relay supports any of the protocols recognized by ANSI.

Data Link Layer

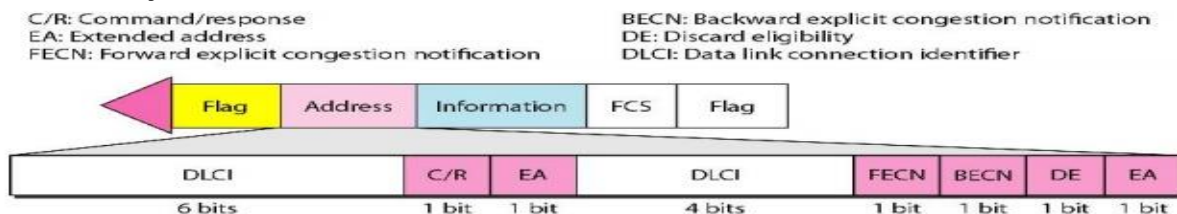


Figure 2.88 Frame Relay frame

At the data link layer, Frame Relay uses a simple protocol that does not support flow or error control. It only has an error detection mechanism. The address field defines the DLCI as well as some bits used to control congestion.

The descriptions of the fields are as follows:

- **Address (DLCI) field.** The first 6 bits of the first byte makes up the first part of the DLCI. The second part of the DLCI uses the first 4 bits of the second byte. These bits are part of the 10-bit data link connection identifier defined by the standard.
- **Command/response (C/R).** The command/response (C/R) bit is provided to allow upper layers to identify a frame as either a command or a response. It is not used by the Frame Relay protocol.
- **Extended address (EA).** The extended address (EA) bit indicates whether the current byte is the final byte of the address. An EA of 0 means that another address byte is to follow. An EA of 1 means that the current byte is the final one.
- **Forward explicit congestion notification (FECN).** The forward explicit congestion notification (FECN) bit can be set by any switch to indicate that traffic is congested. This bit informs the destination that congestion has occurred.
- **Backward explicit congestion notification (BECN).** The backward explicit congestion notification (BECN) bit is set to indicate a congestion problem in the network. This bit informs the sender that congestion has occurred.
- **Discard eligibility (DE).** The discard eligibility (DE) bit indicates the priority level of the frame. In emergency situations, switches may have to discard frames to relieve bottlenecks and keep the network from collapsing due to overload.

➤ **Asynchronous Transfer Mode (ATM)**

Asynchronous Transfer Mode (ATM) is the cell relay protocol designed by the ATM Forum and adopted by the ITU-T. The combination of ATM and SONET will allow high-speed interconnection of all the world's networks. In fact, ATM can be thought of as the "highway" of the information superhighway.

1. Design Goals

Among the challenges faced by the designers of ATM, six stand out.

- a) Foremost is the need for a transmission system to optimize the use of high-data-rate transmission media, in particular optical fiber. In addition to offering large bandwidths, newer transmission media and equipment are dramatically less susceptible to noise degradation. A technology is needed to take advantage of both factors and thereby maximize data rates.
- b) The system must interface with existing systems and provide wide-area interconnectivity between them without lowering their effectiveness or requiring their replacement.
- c) The design must be implemented inexpensively so that cost would not be a barrier to adoption. If ATM is to become the backbone of international communications, as intended, it must be available at low cost to every user who wants it.
- d) The new system must be able to work with and support the existing telecommunications hierarchies (local loops, local providers, long-distance carriers, and so on)
- e) The new system must be connection-oriented to ensure accurate and predictable delivery.
- f) Last but not least, one objective is to move as many of the functions to hardware as possible (for speed) and eliminate as many software functions as possible (again for speed).

2. Frame Networks

Before ATM, data communications at the data link layer had been based on frame switching and frame networks. Different protocols use frames of varying size and intricacy. As networks become more complex, the information that must be carried in the header becomes more extensive. The result is larger and larger headers relative to the size of the data unit. In response, some protocols have enlarged the size of the data unit to make

header use more efficient. Unfortunately, large data fields create waste. If there is not much information to transmit, much of the field goes unused. To improve utilization, some protocols provide variable frame sizes to users.

3. Mixed Network Traffic

The variety of frame sizes makes traffic unpredictable. Switches, multiplexers, and routers must incorporate elaborate software systems to manage the various sizes of frames. A great deal of header information

must be read, and each bit counted and evaluated to ensure the integrity of every frame. Internetworking among the different frame networks is slow and expensive at best, and impossible at worst.

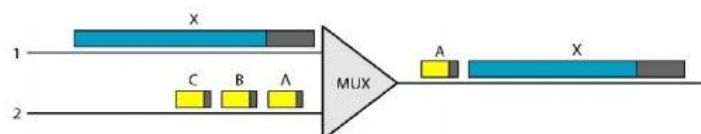


Figure 2.89 Multiplexing using different frame sizes

4. Cell Networks

Many of the problems associated with frame internetworking are solved by adopting a concept called cell networking. A cell is a small data unit of fixed size. In a **cell** network, which uses the **cell** as the basic unit of data exchange, all data are loaded into identical cells that can be transmitted with complete predictability and uniformity. As frames of different sizes and formats reach the cell network from a tributary network, they are split into multiple small data units of equal length and are loaded into cells. The cells are then multiplexed with other cells and routed through the cell network.

5. Asynchronous TDM

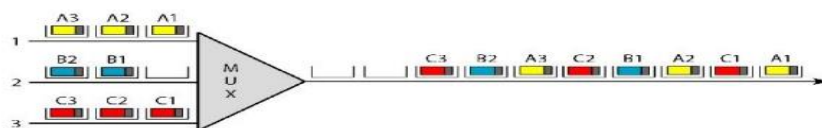


Figure 2.90 ATM Multiplexing

ATM uses asynchronous time-division multiplexing—that is why it is called Asynchronous Transfer Mode—to multiplex cells coming from different channels. It uses fixed-size slots (size of a cell). ATM multiplexers fill a slot with a cell from any input channel that has a cell; the slot is empty if none of the channels has a cell to send

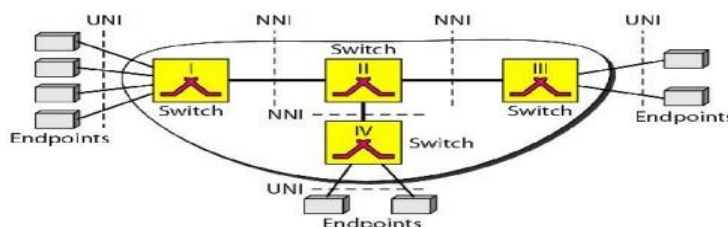


Figure 2.91 Architecture of an ATM network

6. Architecture

ATM is a cell-switched network. The user access devices, called the endpoints, are connected through a user-to-network interface (UNI) to the switches inside the network. The switches are connected through network-to-network interfaces (NNIs).

Virtual Connection

Connection between two endpoints is accomplished through transmission paths (TPs), virtual paths (YPs), and virtual circuits (YCs). A transmission path (TP) is the physical connection (wire, cable, satellite, and so on) between an endpoint and a switch or

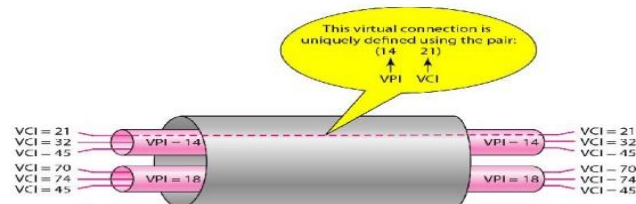


Figure 2.92 Connection identifiers

between two switches. Think of two switches as two cities. A transmission path is the set of all highways that directly connect the two cities.

A transmission path is divided into several virtual paths. A virtual path (VP) provides a connection or a set of connections between two switches. Think of a virtual path as a highway that connects two cities. Each highway is a virtual path; the set of all highways is the transmission path.

Cell networks are based on virtual circuits (VCs). All cells belonging to a single message follow the same virtual circuit and remain in their original order until they reach their destination. Think of a virtual circuit as the lanes of a highway (virtual path).

Identifiers

In a virtual circuit network, to route data from one endpoint to another, the virtual connections need to be identified. For this purpose, the designers of ATM created a hierarchical identifier with two levels: a virtual path identifier (VPI) and a virtual-circuit identifier (VCI). The VPI defines the specific VP, and the VCI defines a particular VC inside the VP. The VPI is the same for all virtual connections that are bundled (logically) into one VP.

7. ATM Layers

The ATM standard defines three layers. They are, from top to bottom, the application adaptation layer, the ATM layer, and the physical layer. The endpoints use all three layers while the switches use only the two bottom layers.

Physical Layer

Like Ethernet and wireless LANs, ATM cells can be carried by any physical layer carrier.

The ATM layer provides routing, traffic management, switching, and multiplexing services. It processes outgoing traffic by accepting 48-byte segments from the AAL sublayers and transforming them into 53-byte cells by the addition of a 5-byte header.

Header Format:

ATM uses two formats for this header, one for user-to-network interface (UNI) cells and another for network-to-network interface (NNI) cells.

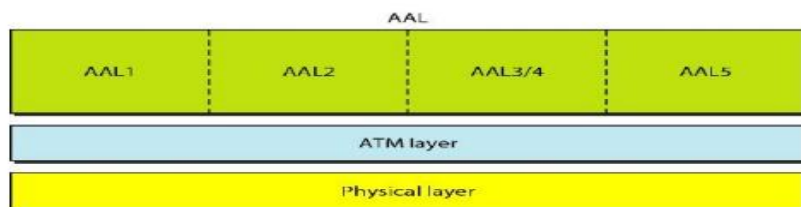


Figure 2.93 ATM layers

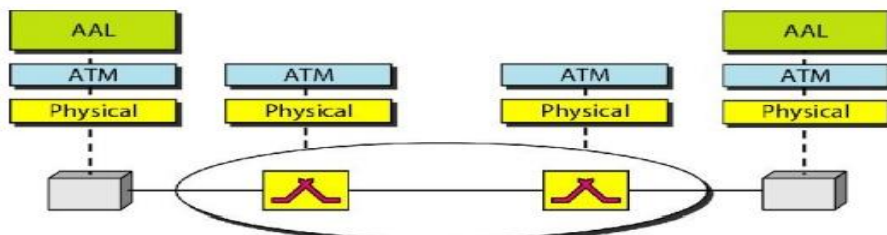


Figure 2.94 ATM layers in endpoint devices and switches



Figure 2.95 ATM layer

UNIT-III

➤ Network Layer

- The Network Layer is the third layer of the OSI model.
- It handles the service requests from the transport layer and further forwards the service request to the data link layer.
- The network layer translates the logical addresses into physical addresses
- It determines the route from the source to the destination and also manages the traffic problems such as switching, routing and controls the congestion of data packets.
- The main role of the network layer is to move the packets from sending host to the receiving host. The main functions performed by the network layer are:
 - **Routing:** When a packet reaches the router's input link, the router will move the packets to the router's output link. For example, a packet from S1 to R1 must be forwarded to the next router on the path to S2.
 - **Logical Addressing:** The data link layer implements the physical addressing and network layer implements the logical addressing. Logical addressing is also used to distinguish between source and destination system. The network layer adds a header to the packet which includes the logical addresses of both the sender and the receiver.
- **Internetworking:** This is the main role of the network layer that it provides the logical connection between different types of networks.
- **Fragmentation:** The fragmentation is a process of breaking the packets into the smallest individual data units that travel through different networks.

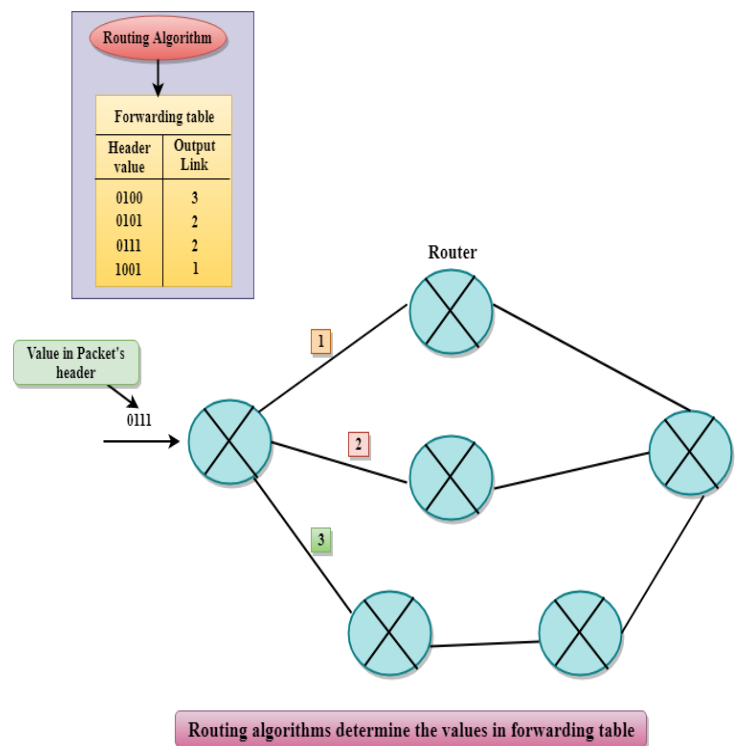
Forwarding & Routing

In Network layer, a router is used to forward the packets. Every router has a forwarding table. A router forwards a packet by examining a packet's header field and then using the header field value to index into the forwarding table. The value stored in the forwarding table corresponding to the header field value indicates the router's outgoing interface link to which the packet is to be forwarded.

For example, the router with a header field value of 0111 arrives at a router, and then router indexes this header value into the forwarding table that determines the output link interface is 2. The router forwards the packet to the interface 2. The routing algorithm determines the values that are inserted in the forwarding table. The routing algorithm can be centralized or decentralized.

Services Provided by the Network Layer

- **Guaranteed delivery:** This layer provides the service which guarantees that the packet will arrive at its destination.
- **Guaranteed delivery with bounded delay:** This service guarantees that the packet will be delivered within a specified host-to-host delay bound.
- **In-Order packets:** This service ensures that the packet arrives at the destination in the order in which they are sent.
- **Guaranteed max jitter:** This service ensures that



the amount of time taken between two successive transmissions at the sender is equal to the time between their receipt at the destination.

- Security services: The network layer provides security by using a session key between the source and destination host. The network layer in the source host encrypts the payloads of datagrams being sent to the destination host. The network layer in the destination host would then decrypt the payload. In such a way, the network layer maintains the data integrity and source authentication services.

➤ Network Addressing

- There are two types of addressing schemes:
 - Classful Addressing
 - Classless Addressing

Classful Addressing

An IP address is 32-bit long. An IP address is divided into sub-classes:

- **Class A**

- **Class B**

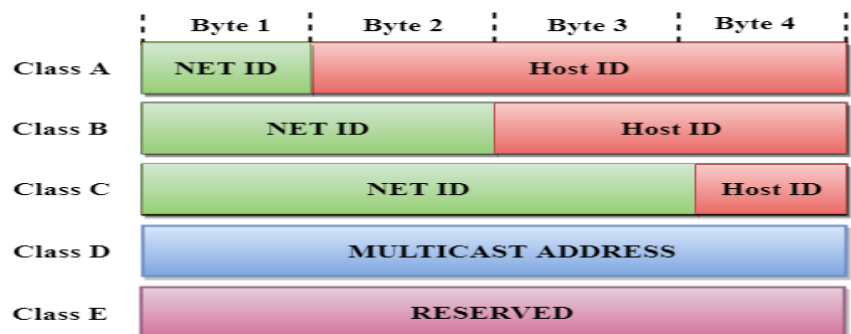
- **Class C**

- **Class D**

- **Class E**

An ip address is divided into two parts:

- Network ID: It represents the number of networks.
- Host ID: It represents the number of hosts.



In the above diagram, we observe that each class have a specific range of IP addresses. The class of IP address is used to determine the number of bits used in a class and number of networks and hosts available in the class.

Class A In Class A, an IP address is assigned to those networks that contain a large number of hosts.

Exception Hand lThe network ID is 8 bits long.

- The host ID is 24 bits long.

In Class A, the first bit in higher order bits of the first octet is always set to 0 and the remaining 7 bits determine the network ID. The 24 bits determine the host ID in any network.

The total number of networks in Class A = $2^7 = 128$ network address

The total number of hosts in Class A = $2^{24} - 2 = 16,777,214$ host address

Class B

In Class B, an IP address is assigned to those networks that range from small-sized to large-sized networks.

- The Network ID is 16 bits long.
- The Host ID is 16 bits long.

In Class B, the higher order bits of the first octet is always set to 10, and the remaining 14 bits determine the network ID. The other 16 bits determine the Host ID.

The total number of networks in Class B = $2^{14} = 16384$ network address

The total number of hosts in Class B = $2^{16} - 2 = 65534$ host address

Class C

In Class C, an IP address is assigned to only small-sized networks.



- The Network ID is 24 bits long.
- The host ID is 8 bits long.

In Class C, the higher order bits of the first octet is always set to 110, and the remaining 21 bits determine the network ID. The 8 bits of the host ID determine the host in a network.



The total number of networks = 2^{21} = 2097152 network address

The total number of hosts = $2^8 - 2$ = 254 host address

Class D

In Class D, an IP address is reserved for multicast addresses. It does not possess subnetting. The higher order bits of the first octet is always set to 1110, and the remaining bits determines the host ID in any network.



Class E

In Class E, an IP address is used for the future use or for the research and development purposes. It does not possess any subnetting. The higher order bits of the first octet is always set to 1111, and the remaining bits determines the host ID in any network.



Classful Network Architecture

Class	Higher bits	NET ID bits	HOST ID bits	No.of networks	No.of hosts per network	Range
A	0	8	24	2^7	2^{24}	0.0.0.0 to 127.255.255.255
B	10	16	16	2^{14}	2^{16}	128.0.0.0 to 191.255.255.255
C	110	24	8	2^{21}	2^8	192.0.0.0 to 223.255.255.255
D	1110	Not Defined	Not Defined	Not Defined	Not Defined	224.0.0.0 to 239.255.255.255
E	1111	Not Defined	Not Defined	Not Defined	Not Defined	240.0.0.0 to 255.255.255.255

Classless Addressing

To overcome address depletion and give more organizations access to the Internet, classless addressing was designed and implemented. In this scheme, there are no classes, but the addresses are still granted in blocks. Address Blocks In classless addressing, when an entity, small or large, needs to be connected to the Internet, it is granted a block (range) of addresses. The size of the block (the number of addresses) varies based on the nature and size of the entity. For example, a household may be given only two addresses; a large organization may be given thousands of addresses.

An ISP, as the Internet service provider, may be given thousands or hundreds of thousands based on the number of customers it may serve. The Internet authorities impose three restrictions on classless address blocks:

1. The addresses in a block must be contiguous, one after another.
2. The number of addresses in a block must be a power of 2 (1, 2, 4, 8, ...).
3. The first address must be evenly divisible by the number of addresses.

➤ Routing

- A Router is a process of selecting path along which the data can be transferred from source to the destination. Routing is performed by a special device known as a router.
- A Router works at the network layer in the OSI model and internet layer in TCP/IP model
- A router is a networking device that forwards the packet based on the information available in the packet header and forwarding table.
- The routing algorithms are used for routing the packets. The routing algorithm is nothing but a software responsible for deciding the optimal path through which packet can be transmitted.
- The routing protocols use the metric to determine the best path for the packet delivery. The metric is the standard of measurement such as hop count, bandwidth, delay, current load on the path, etc. used by the routing algorithm to determine the optimal path to the destination.
- The routing algorithm initializes and maintains the routing table for the process of path determination.

Routing Metrics and Costs

Routing metrics and costs are used for determining the best route to the destination. The factors used by the protocols to determine the shortest path, these factors are known as a metric.

Metrics are the network variables used to determine the best route to the destination. For some protocols use the static metrics means that their value cannot be changed and for some other routing protocols use the dynamic metrics means that their value can be assigned by the system administrator.

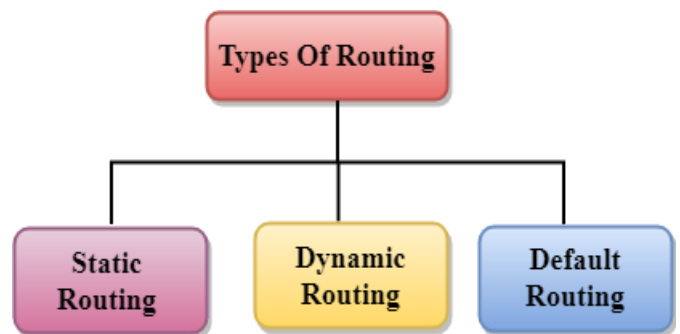
The most common metric values are given below:

- Hop count: Hop count is defined as a metric that specifies the number of passes through internetworking devices such as a router, a packet must travel in a route to move from source to the destination. If the routing protocol considers the hop as a primary metric value, then the path with the least hop count will be considered as the best path to move from source to the destination.
- Delay: It is a time taken by the router to process, queue and transmit a datagram to an interface. The protocols use this metric to determine the delay values for all the links along the path end-to-end. The path having the lowest delay value will be considered as the best path.
- Bandwidth: The capacity of the link is known as a bandwidth of the link. The bandwidth is measured in terms of bits per second. The link that has a higher transfer rate like gigabit is preferred over the link that has the lower capacity like 56 kb. The protocol will determine the bandwidth capacity for all the links along the path, and the overall higher bandwidth will be considered as the best route.
- Load: Load refers to the degree to which the network resource such as a router or network link is busy. A Load can be calculated in a variety of ways such as CPU utilization, packets processed per second. If the traffic increases, then the load value will also be increased. The load value changes with respect to the change in the traffic.
- Reliability: Reliability is a metric factor may be composed of a fixed value. It depends on the network links, and its value is measured dynamically. Some networks go down more often than others. After network failure, some network links repaired more easily than other network links. Any reliability factor can be considered for the assignment of reliability ratings, which are generally numeric values assigned by the system administrator.

Types of Routing

Routing can be classified into three categories:

- Static Routing



- Default Routing
- Dynamic Routing
- Static Routing
- **Static Routing** is also known as Non adaptive Routing.
- It is a technique in which the administrator manually adds the routes in a routing table.
- A Router can send the packets for the destination along the route defined by the administrator.
- In this technique, routing decisions are not made based on the condition or topology of the networks

Advantages Of Static Routing

Following are the advantages of Static Routing:

- **No Overhead:** It has no overhead on the CPU usage of the router. Therefore, the cheaper router can be used to obtain static routing.
- **Bandwidth:** It has not bandwidth usage between the routers.
- **Security:** It provides security as the system administrator is allowed only to have control over the routing to a particular network.

Disadvantages of Static Routing:

Following are the disadvantages of Static Routing:

- For a large network, it becomes a very difficult task to add each route manually to the routing table.
- The system administrator should have a good knowledge of a topology as he has to add each route manually.

Default Routing

- Default Routing is a technique in which a router is configured to send all the packets to the same hop device, and it doesn't matter whether it belongs to a particular network or not. A Packet is transmitted to the device for which it is configured in default routing.
- Default Routing is used when networks deal with the single exit point.
- It is also useful when the bulk of transmission networks have to transmit the data to the same hp device.
- When a specific route is mentioned in the routing table, the router will choose the specific route rather than the default route. The default route is chosen only when a specific route is not mentioned in the routing table.

Dynamic Routing

- It is also known as Adaptive Routing.
- It is a technique in which a router adds a new route in the routing table for each packet in response to the changes in the condition or topology of the network.
- Dynamic protocols are used to discover the new routes to reach the destination.
- In Dynamic Routing, RIP and OSPF are the protocols used to discover the new routes.
- If any route goes down, then the automatic adjustment will be made to reach the destination.

The Dynamic protocol should have the following features:

- All the routers must have the same dynamic routing protocol in order to exchange the routes.
- If the router discovers any change in the condition or topology, then router broadcast this information to all other routers.

Advantages of Dynamic Routing:

- It is easier to configure.
- It is more effective in selecting the best route in response to the changes in the condition or topology.

Disadvantages of Dynamic Routing:

- It is more expensive in terms of CPU and bandwidth usage.
- It is less secure as compared to default and static routing

➤ Network Layer Protocols

Network protocols are a set of rules, conventions, and data structures that dictate how devices exchange data across networks. In other words, network protocols can be equated to languages that two devices must understand for seamless communication of information, regardless of their infrastructure and design disparities.

Address Resolution Protocol(ARP)

While communicating, a host needs Layer-2 (MAC) address of the destination machine which belongs to the same broadcast domain or network. A MAC address is physically burnt into the Network Interface Card (NIC) of a machine and it never changes.

To know the MAC address of remote host on a broadcast domain, a computer wishing to initiate communication sends out an ARP broadcast message asking, “Who has this IP address?”. ARP packet contains the IP address of destination host, the sending host wishes to talk to. When a host receives an ARP packet destined to it, it replies back with its own MAC address.

Internet Control Message Protocol (ICMP)

ICMP is network diagnostic and error reporting protocol. ICMP belongs to IP protocol suite and uses IP as carrier protocol. After constructing ICMP packet, it is encapsulated in IP packet. Because IP itself is a best-effort non-reliable protocol, so is ICMP.

Internet Control Message Protocol (ICMP)

ICMP is network diagnostic and error reporting protocol. ICMP belongs to IP protocol suite and uses IP as carrier protocol. After constructing ICMP packet, it is encapsulated in IP packet. Because IP itself is a best-effort non-reliable protocol, so is ICMP.

Any feedback about network is sent back to the originating host. If some error in the network occurs, it is reported by means of ICMP. ICMP contains dozens of diagnostic and error reporting messages.

ICMP-echo and ICMP-echo-reply are the most commonly used ICMP messages to check the reachability of end-to-end hosts. When a host receives an ICMP-echo request, it is bound to send back an ICMP-echo-reply. If there is any problem in the transit network, the ICMP will report that problem.

Internet Protocol Version 4 (IPv4)

IPv4 is 32-bit addressing scheme used as TCP/IP host addressing mechanism. IP addressing enables every host on the TCP/IP network to be uniquely identifiable.

Internet Protocol Version 6 (IPv6)

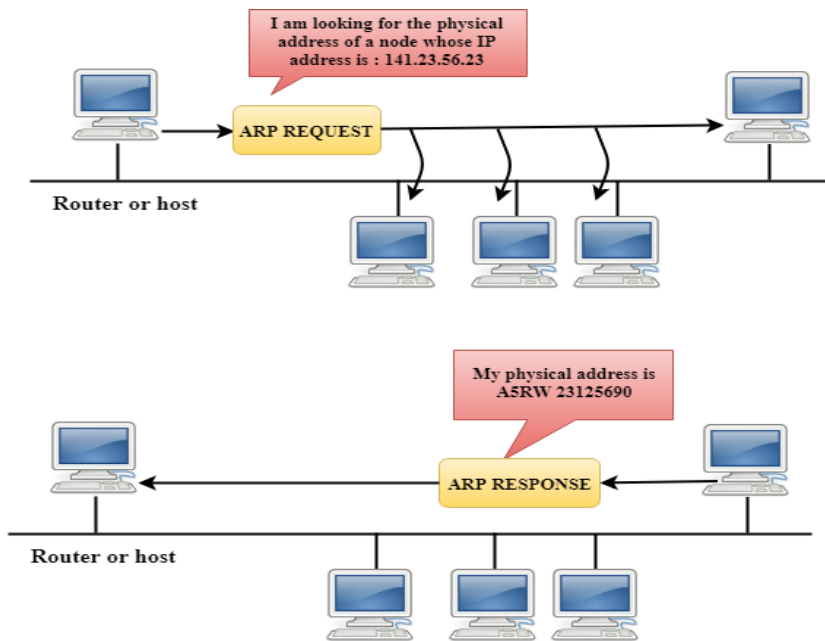
Exhaustion of IPv4 addresses gave birth to a next generation Internet Protocol version 6. IPv6 addresses its nodes with 128-bit wide address providing plenty of address space for future to be used on entire planet or beyond.

➤ ARP stands for Address Resolution Protocol.

- It is used to associate an IP address with the MAC address.
- Each device on the network is recognized by the MAC address imprinted on the NIC. Therefore, we can say that devices need the MAC address for communication on a local area network. MAC address can be changed easily. For example, if the NIC on a particular machine fails, the MAC address changes but IP address does not change. ARP is used to find the MAC address of the node when an internet address is known.

How ARP works

If the host wants to know the physical address of another host on its network, then it sends an ARP query packet that includes the IP address and broadcast it over the network. Every host on the network receives and processes the ARP packet, but only the intended recipient recognizes the IP address and sends back the physical address. The host holding the datagram adds the physical address to the cache memory and to the datagram header, then sends back to the sender.



Steps taken by ARP protocol

If a device wants to communicate with another device, the following steps are taken by the device:

- The device will first look at its internet list, called the ARP cache to check whether an IP address contains a matching MAC address or not. It will check the ARP cache in command prompt by using a command arp-a.

```

C:\WINDOWS\system32\cmd.exe

C:\>arp -a
No ARP Entries Found
  
```

- If ARP cache is empty, then device broadcast the message to the entire network asking each device for a matching MAC address.
- The device that has the matching IP address will then respond back to the sender with its MAC address
- Once the MAC address is received by the device, then the communication can take place between two devices.
- If the device receives the MAC address, then the MAC address gets stored in the ARP cache. We can check the ARP cache in command prompt by using a command arp -a.

```

C:\Users\admin>arp -a

Interface: 192.168.1.10 --- 0x3
Internet Address      Physical Address      Type
192.168.1.1           74-da-da-db-f7-67    dynamic
192.168.1.11          fc-aa-14-ee-cc-c2    dynamic
192.168.1.14          18-60-24-bd-3d-1d    dynamic
192.168.1.32          1c-1b-0d-bd-d2-7e    dynamic
192.168.1.41          58-20-b1-40-b7-74    dynamic
192.168.1.55          fc-aa-14-a5-67-7a    dynamic
192.168.1.255         ff-ff-ff-ff-ff-ff    static
224.0.0.22            01-00-5e-00-00-16    static
224.0.0.251           01-00-5e-00-00-fb    static
224.0.0.252           01-00-5e-00-00-fc    static
239.255.255.250       01-00-5e-7f-ff-fa    static
255.255.255.255       ff-ff-ff-ff-ff-ff    static
  
```

In the above screenshot, we observe the association of IP address to the MAC address. There are two types of ARP entries:

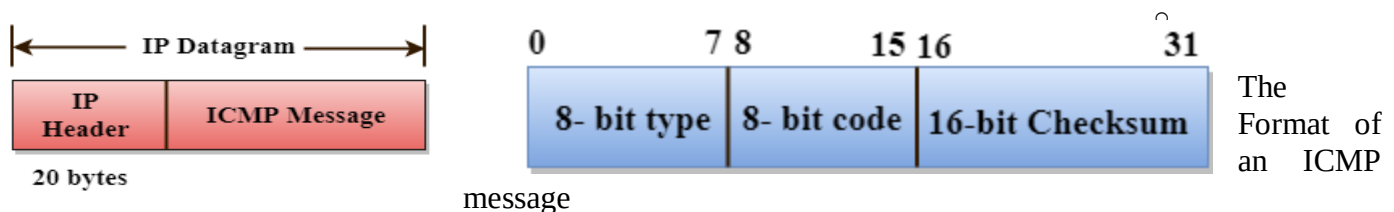
- **Dynamic entry:** It is an entry which is created automatically when the sender broadcast its message to the entire network. Dynamic entries are not permanent, and they are removed periodically.
- **Static entry:** It is an entry where someone manually enters the IP to MAC address association by using the ARP command utility.

RARP

- RARP stands for Reverse Address Resolution Protocol.
- If the host wants to know its IP address, then it broadcast the RARP query packet that contains its physical address to the entire network. A RARP server on the network recognizes the RARP packet and responds back with the host IP address.
- The protocol which is used to obtain the IP address from a server is known as Reverse Address Resolution Protocol.
- The message format of the RARP protocol is similar to the ARP protocol.
- Like ARP frame, RARP frame is sent from one machine to another encapsulated in the data portion of a frame.

➤ ICMP

- ICMP stands for Internet Control Message Protocol.
- The ICMP is a network layer protocol used by hosts and routers to send the notifications of IP datagram problems back to the sender.
- ICMP uses echo test/reply to check whether the destination is reachable and responding.
- ICMP handles both control and error messages, but its main function is to report the error but not to correct them.
- An IP datagram contains the addresses of both source and destination, but it does not know the address of the previous router through which it has been passed. Due to this reason, ICMP can only send the messages to the source, but not to the immediate routers.
- ICMP protocol communicates the error messages to the sender. ICMP messages cause the errors to be returned back to the user processes.
- ICMP messages are transmitted within IP datagram



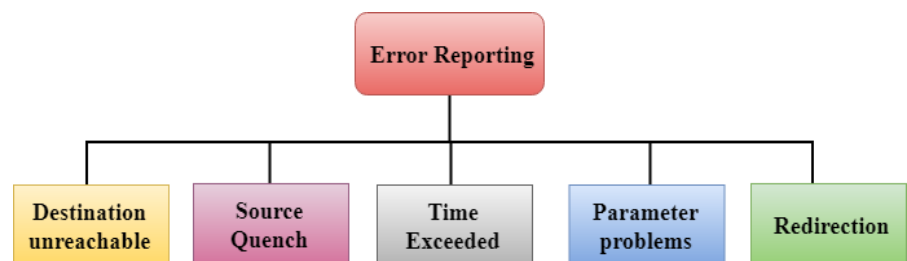
- The first field specifies the type of the message.
- The second field specifies the reason for a particular message type.
- The checksum field covers the entire ICMP message.

Error Reporting

ICMP protocol reports the error messages to the sender.

Five types of errors are handled by the ICMP protocol:

- Destination unreachable
- Source Quench
- Time Exceeded
- Parameter problems
- Redirection



- **Destination unreachable:** The message of "Destination Unreachable" is sent from receiver to the sender when destination cannot be reached, or packet is discarded when the destination is not reachable.
- **Source Quench:** The purpose of the source quench message is congestion control. The message sent from the congested router to the source host to reduce the transmission rate. ICMP will take the IP of the discarded packet and then add the source quench message to the IP datagram to inform the source host to reduce its transmission rate. The source host will reduce the transmission rate so that the router will be free from congestion.
- **Time Exceeded:** Time Exceeded is also known as "Time-To-Live". It is a parameter that defines how long a packet should live before it would be discarded.

There are two ways when Time Exceeded message can be generated:

Sometimes packet discarded due to some bad routing implementation, and this causes the looping issue and network congestion. Due to the looping issue, the value of TTL keeps on decrementing, and when it reaches zero, the router discards the datagram. However, when the datagram is discarded by the router, the time exceeded message will be sent by the router to the source host.

When destination host does not receive all the fragments in a certain time limit, then the received fragments are also discarded, and the destination host sends time Exceeded message to the source host.

- **Parameter problems:** When a router or host discovers any missing value in the IP datagram, the router discards the datagram, and the "parameter problem" message is sent back to the source host.
- **Redirection:** Redirection message is generated when host consists of a small routing table. When the host consists of a limited number of entries due to which it sends the datagram to a wrong router. The router that receives a datagram will forward a datagram to a correct router and also sends the "Redirection message" to the host to update its routing table.

➤ IPV4

What is IP?

An IP stands for internet protocol. An IP address is assigned to each device connected to a network. Each device uses an IP address for communication. It also behaves as an identifier as this address is used to identify the device on a network. It defines the technical format of the packets. Mainly, both the networks, i.e., IP and TCP, are combined together, so together, they are referred to as a [TCP/IP](#). It creates a virtual connection between the source and the destination.

What is IPv4?

IPv4 is a version 4 of IP. It is a current version and the most commonly used IP address. It is a 32-bit address written in four numbers separated by 'dot', i.e., periods. This address is unique for each device.

For example, 66.94.29.13

The above example represents the IP address in which each group of numbers separated by periods is called an Octet. Each number in an octet is in the range from 0-255. This address can produce 4,294,967,296 possible unique addresses.

In today's computer network world, computers do not understand the IP addresses in the standard numeric format as the computers understand the numbers in binary form only. The binary number can be either 1 or 0. The IPv4 consists of four sets, and these sets represent the octet. The bits in each octet represent a number. Each bit in an octet can be either 1 or 0. If the bit the 1, then the number it represents will count, and if the bit is 0, then the number it represents does not count.

Representation of 8 Bit Octet

128	64	32	16	8	4	2	1
-----	----	----	----	---	---	---	---

The above representation shows the structure of 8-bit octet.

Now, we will see how to obtain the binary representation of the above IP address, i.e., 66.94

Step 1: First, we find the binary number of 66.

128	64	32	16	8	4	2	1
0	1	0	0	0	0	1	0

To obtain 66, we put 1 under 64 and 2 as the sum of 64 and 2 is equal to 66 ($64+2=66$), and the remaining bits will be zero, as shown above. Therefore, the binary bit version of 66 is 01000010.

Step 2: Now, we calculate the binary number of 94.

128	64	32	16	8	4	2	1
0	1	0	1	1	1	1	0

To obtain 94, we put 1 under 64, 16, 8, 4, and 2 as the sum of these numbers is equal to 94, and the remaining bits will be zero. Therefore, the binary bit version of 94 is 01011110.

Drawback of IPv4

Currently, the population of the world is 7.6 billion. Every user is having more than one device connected with the internet, and private companies also rely on the internet. As we know that IPv4 produces 4 billion addresses, which are not enough for each device connected to the internet on a planet. Although the various techniques were invented, such as variable-length mask, network address translation, port address translation, classes, inter-domain translation, to conserve the bandwidth of IP address and slow down the depletion of an IP address. In these techniques, public IP is converted into a private IP due to which the user having public IP can also use the internet. But still, this was not so efficient, so it gave rise to the development of the next generation of IP addresses, i.e., IPv6.

➤ IPv6

IPv6 is the next generation of IP addresses. The main difference between IPv4 and IPv6 is the address size of IP addresses. The IPv4 is a 32-bit address, whereas IPv6 is a 128-bit hexadecimal address. IPv6 provides a large address space, and it contains a simple header as compared to IPv4.

It provides transition strategies that convert IPv4 into IPv6, and these strategies are as follows:

Dual stacking: It allows us to have both the versions, i.e., IPv4 and IPv6, on the same device.

Tunneling: In this approach, all the users have IPv6 communicates with an IPv4 network to reach IPv6.

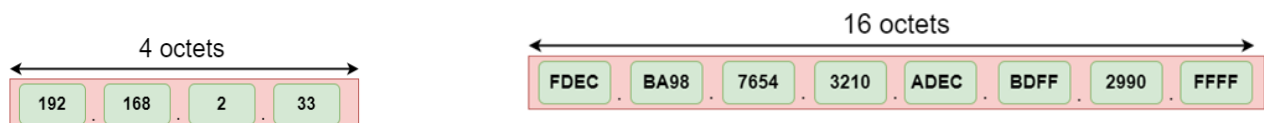
Network Address Translation: The translation allows the communication between the hosts having a different version of IP.

This hexadecimal address contains both numbers and alphabets. Due to the usage of both the numbers and alphabets, IPv6 is capable of producing over 340 undecillion (3.4×10^{38}) addresses.

IPv6 is a 128-bit hexadecimal address made up of 8 sets of 16 bits each, and these 8 sets are separated by a colon. In IPv6, each hexadecimal character represents 4 bits. So, we need to convert 4 bits to a hexadecimal number at a time

Address format

The address format of IPv4:



The address format of IPv6:

The above diagram shows the address format of IPv4 and IPv6. An IPv4 is a 32-bit decimal address. It contains 4 octets or fields separated by 'dot', and each field is 8-bit in size. The number that each field contains should be in the range of 0-255. Whereas an IPv6 is a 128-bit hexadecimal address. It contains 8 fields separated by a colon, and each field is 16-bit in size.

Differences between IPv4 and IPv6

	Ipv4	Ipv6
Address length	IPv4 is a 32-bit address.	IPv6 is a 128-bit address.
Fields	IPv4 is a numeric address that consists of 4 fields which are separated by dot (.).	IPv6 is an alphanumeric address that consists of 8 fields, which are separated by colon.
Classes	IPv4 has 5 different classes of IP address that includes Class A, Class B, Class C, Class D, and Class E.	IPv6 does not contain classes of IP addresses.
Number of IP address	IPv4 has a limited number of IP addresses.	IPv6 has a large number of IP addresses.
VLSM	It supports VLSM (Virtual Length Subnet Mask). Here, VLSM means that Ipv4 converts IP addresses into a subnet of different sizes.	It does not support VLSM.
Address configuration	It supports manual and DHCP configuration.	It supports manual, DHCP, Auto-configuration, and renumbering.
Address space	It generates 4 billion unique addresses	It generates 340 undecillion unique addresses.
End-to-end connection integrity	In IPv4, end-to-end connection integrity is unachievable.	In the case of IPv6, end-to-end connection integrity is achievable.
Security features	In IPv4, security depends on the application. This IP address is not developed in keeping the security feature in mind.	In IPv6, IPSEC is developed for security purposes.
Address representation	In IPv4, the IP address is represented in decimal.	In IPv6, the representation of the IP address in hexadecimal.
Fragmentation	Fragmentation is done by the senders and the forwarding routers.	Fragmentation is done by the senders only.
Packet flow identification	It does not provide any mechanism for packet flow identification.	It uses flow label field in the header for the packet flow identification.
Checksum field	The checksum field is available in IPv4.	The checksum field is not available in IPv6.
Transmission scheme	IPv4 is broadcasting.	On the other hand, IPv6 is multicasting, which provides efficient network operations.
Encryption and Authentication	It does not provide encryption and authentication.	It provides encryption and authentication.
Number of octets	It consists of 4 octets.	It consists of 8 fields, and each field contains 2 octets. Therefore, the total number of octets in IPv6 is 16.

➤ Routing algorithm

In order to transfer the packets from source to the destination, the network layer must determine the best route through which packets can be transmitted.

Whether the network layer provides datagram service or virtual circuit service, the main job of the network layer is to provide the best route. The routing protocol provides this job. The routing protocol is a routing algorithm that provides the best path from the source to the destination. The best path is the path that has the "least-cost path" from source to the destination.

Routing is the process of forwarding the packets from source to the destination but the best route to send the packets is determined by the routing algorithm.

Classification of a Routing algorithm

The Routing algorithm is divided into two categories:

Adaptive Routing algorithm

Non-adaptive Routing algorithm

Adaptive Routing algorithm

An adaptive routing algorithm is also known as dynamic routing algorithm.

This algorithm makes the routing decisions based on the topology and network traffic.

The main parameters related to this algorithm are hop count, distance and estimated transit time.

Non-Adaptive Routing algorithm

Non Adaptive routing algorithm is also known as a static routing algorithm.

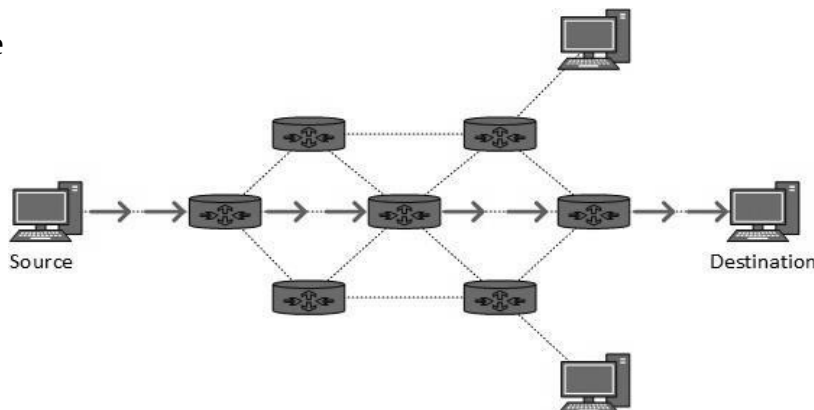
When booting up the network, the routing information stores to the routers.

Non Adaptive routing algorithms do not take the routing decision based on the network topology or network traffic.

Unicast routing

Most of the traffic on the internet and intranets known as unicast data or unicast traffic is sent with specified destination. Routing unicast data over the internet is called unicast routing. It is the simplest form of routing because the destination is already known.

Hence the router just has to look up the table and forward the packet to next hop.



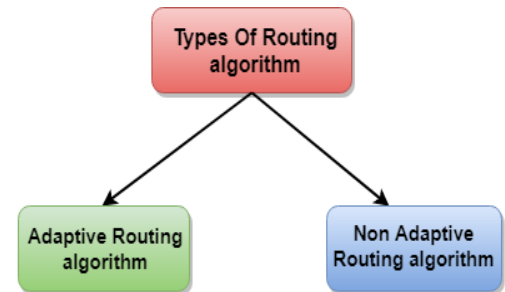
the destination is already known. router just has to look up the table and forward the packet to

Unicast Routing Protocols

There are two kinds of routing protocols available to route unicast packets:

Distance Vector Routing Protocol

Distance Vector is simple routing protocol which takes routing decision on the number of hops between source and destination. A route with less number of hops is considered as the best route. Every router advertises its set best routes to other routers. Ultimately, all routers build up their network topology based on the advertisements of their peer routers, for example, Routing Information Protocol (RIP).



Distance Vector Routing Algorithm

- **The Distance vector algorithm is iterative, asynchronous and distributed.**
 - **Distributed:** It is distributed in that each node receives information from one or more of its directly attached neighbors, performs calculation and then distributes the result back to its neighbors.
 - **Iterative:** It is iterative in that its process continues until no more information is available to be exchanged between neighbors.
 - **Asynchronous:** It does not require that all of its nodes operate in the lock step with each other.
- The Distance vector algorithm is a dynamic algorithm.
- It is mainly used in ARPANET, and RIP.
- Each router maintains a distance table known as **Vector**.

Three Keys to understand the working of Distance Vector Routing Algorithm:

- **Knowledge about the whole network:** Each router shares its knowledge through the entire network. The Router sends its collected knowledge about the network to its neighbors.
- **Routing only to neighbors:** The router sends its knowledge about the network to only those routers which have direct links. The router sends whatever it has about the network through the ports. The information is received by the router and uses the information to update its own routing table.
- **Information sharing at regular intervals:** Within 30 seconds, the router sends the information to the neighboring routers.

Distance Vector Routing Algorithm

Algorithm

At each node x,

Initialization

for all destinations y in N:

$D_x(y) = c(x,y)$ // If y is not a neighbor then $c(x,y) = \infty$

for each neighbor w

$D_w(y) = ?$ for all destination y in N.

for each neighbor w

send distance vector $D_x = [D_x(y) : y \text{ in } N]$ to w

loop

wait(until I receive any distance vector from some neighbor w)

for each y in N:

$D_x(y) = \min_v \{ c(x,v) + D_v(y) \}$

If $D_x(y)$ is changed for any destination y

Send distance vector $D_x = [D_x(y) : y \text{ in } N]$ to all neighbors **forever**

Link State Routing Protocol

Link State protocol is slightly complicated protocol than Distance Vector. It takes into account the states of links of all the routers in a network. This technique helps routes build a common graph of the entire network. All routers then calculate their best path for routing purposes, for example, Open Shortest Path First (OSPF) and Intermediate System to Intermediate System (ISIS).

Link State Routing

Link state routing is a technique in which each router shares the knowledge of its neighborhood with every other router in the internetwork.

The three keys to understand the Link State Routing algorithm:

- **Knowledge about the neighborhood:** Instead of sending its routing table, a router sends the information about its neighborhood only. A router broadcast its identities and cost of the directly attached links to other routers.
- **Flooding:** Each router sends the information to every other router on the internetwork except its neighbors. This process is known as Flooding. Every router that receives the packet sends the copies to all its neighbors. Finally, each and every router receives a copy of the same information.
- **Information sharing:** A router sends the information to every other router only when the change occurs in the information.

Link State Routing has two phases:

Reliable Flooding

- **Initial state:** Each node knows the cost of its neighbors.
- **Final state:** Each node knows the entire graph.

Route Calculation

Each node uses Dijkstra's algorithm on the graph to calculate the optimal routes to all nodes.

- The Link state routing algorithm is also known as Dijkstra's algorithm which is used to find the shortest path from one node to every other node in the network.
- The Dijkstra's algorithm is an iterative, and it has the property that after k^{th} iteration of the algorithm, the least cost paths are well known for k destination nodes.

Algorithm

Initialization

$N = \{A\}$ // A is a root node.

for all nodes v

if v adjacent to A

then $D(v) = c(A,v)$

else $D(v) = \text{infinity}$

loop

find w not in N such that $D(w)$ is a minimum.

Add w to N

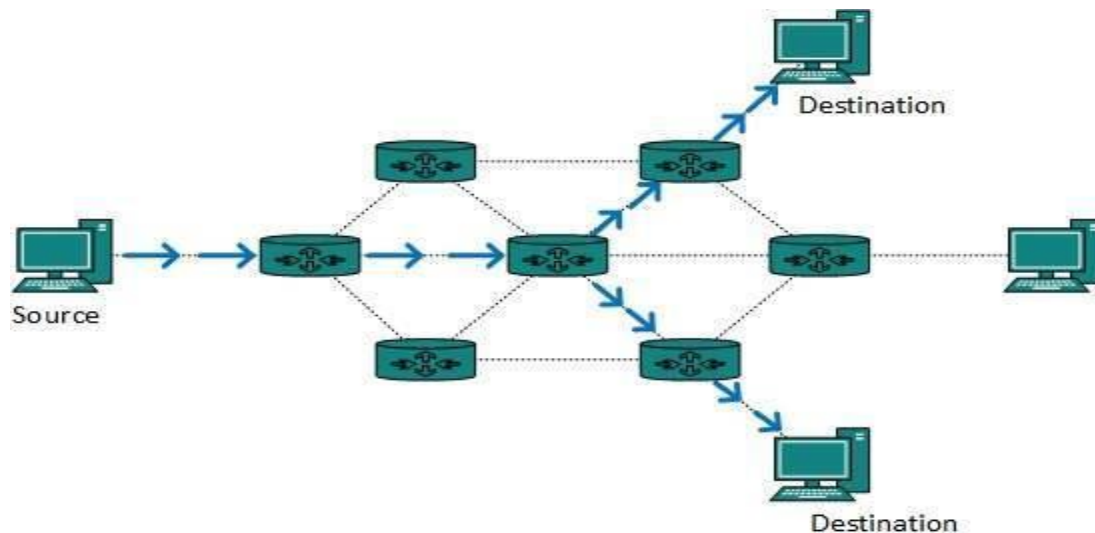
Update $D(v)$ for all v adjacent to w and not in N :

$D(v) = \min(D(v), D(w) + c(w,v))$

Until all nodes in N

In the above algorithm, an initialization step is followed by the loop. The number of times the loop is executed is equal to the total number of nodes available in the network.

➤ **Multicast Routing**



Multicast routing is special case of broadcast routing with significance difference and challenges. In broadcast routing, packets are sent to all nodes even if they do not want it. But in Multicast routing, the data is sent to only nodes which wants to receive the packets.

The router must know that there are nodes, which wish to receive multicast packets (or stream) then only it should forward. Multicast routing works spanning tree protocol to avoid looping.

➤ RIP Protocol

RIP stands for Routing Information Protocol. RIP is an intra-domain routing protocol used within an autonomous system. Here, intra-domain means routing the packets in a defined domain, for example, web browsing within an institutional area. To understand the RIP protocol, our main focus is to know the structure of the packet, how many fields it contains, and how these fields determine the routing table.

- RIP is based on the distance vector-based strategy, so we consider the entire structure as a graph where nodes are the routers, and the links are the networks.

RIP Message Format

Now, we look at the structure of the RIP message format. The message format is used to share information among different routers. The RIP contains the following fields in a message:

Command	Version	Reserved
Family		All 0s
Network address		
All 0s		
All 0s		
Distance		

Repeating the above structure for multiple destinations.

- Command: It is an 8-bit field that is used for request or reply. The value of the request is 1, and the value of the reply is 2.
- Version: Here, version means that which version of the protocol we are using. Suppose we are using the protocol of version1, then we put the 1 in this field.
- Reserved: This is a reserved field, so it is filled with zeroes.
- Family: It is a 16-bit field. As we are using the TCP/IP family, so we put 2 value in this field.

- Network Address: It is defined as 14 bytes field. If we use the IPv4 version, then we use 4 bytes, and the other 10 bytes are all zeroes.
- Distance: The distance field specifies the hop count, i.e., the number of hops used to reach the destination

Advantages of RIP

The following are the advantages of a RIP protocol:

- It is easy to configure
- It has less complexity
- The CPU utilization is less.

Disadvantages of RIP

The following are the disadvantages of RIP:

- The RIP is a classful routing protocol, so it does not support the VLSM (Variable Length Subnet Mask). The classful routing protocol is a protocol that does not include the subnet mask information in the routing updates.
- It broadcasts the routing updates to the entire network that creates a lot of traffic. In RIP, the routing table updates every 30 seconds. Whenever the updates occur, it sends the copy of the update to all the neighbors except the one that has caused the update. The sending of updates to all the neighbors creates a lot of traffic. This rule is known as a split-horizon rule.
- It faces a problem of Slow convergence. Whenever the router or link fails, then it often takes minutes to stabilize or take an alternative route; This problem is known as Slow convergence.
- RIP supports maximum 15 hops which means that the maximum 16 hops can be configured in a RIP
- The Administrative distance value is 120 (Ad value). If the Ad value is less, then the protocol is more reliable than the protocol with more Ad value.
- The RIP protocol has the highest Ad value, so it is not as reliable as the other routing protocols.

➤ OSPF Protocol

The OSPF stands for **Open Shortest Path First**. It is a widely used and supported routing protocol. It is an intradomain protocol, which means that it is used within an area or a network. It is an interior gateway protocol that has been designed within a single autonomous system. It is based on a link-state routing algorithm in which each router contains the information of every domain, and based on this information, it determines the shortest path. The goal of routing is to learn routes. The OSPF achieves by learning about every router and subnet within the entire network. Every router contains the same information about the network. The way the router learns this information by sending LSA (Link State Advertisements). These LSAs contain information about every router, subnet, and other networking information. Once the LSAs have been flooded, the OSPF stores the information in a link-state database known as LSDB. The main goal is to have the same information about every router in an LSDBs.

Types of links in OSPF

A link is basically a connection, so the connection between two routers is known as a link.

There are four types of links in OSPF:

1. **Point-to-point link:** The point-to-point link directly connects the two routers without any host or router in between.
2. **Transient link:** When several routers are attached in a network, they are known as a transient link. The transient link has two different implementations

Unrealistic topology: When all the routers are connected to each other, it is known as an unrealistic topology.

Realistic topology: When some designated router exists in a network then it is known as a realistic topology. Here designated router is a router to which all the routers are connected. All the packets sent by the routers will be passed through the designated router.

3. **Stub link:** It is a network that is connected to the single router. Data enters to the network through the single router and leaves the network through the same router.
4. **Virtual link:** If the link between the two routers is broken, the administration creates the virtual path between the routers, and that path could be a long one also.

OSPF Message Format

Version(8)	Type(8)	Message (16)
Source IP address		
Area Identification		
Chcek sum	Auth.Type	
Authentication (32)		

The following are the fields in an OSPF message format:

Version: It is an 8-bit field that specifies the OSPF protocol version.

- **Type:** It is an 8-bit field. It specifies the type of the OSPF packet.

- **Message:** It is a 16-bit field that defines the total length of the message, including the header. Therefore, the total length is equal to the sum of the

length of the message and header.

- **Source IP address:** It defines the address from which the packets are sent. It is a sending routing IP address.
- **Area identification:** It defines the area within which the routing takes place.
- **Checksum:** It is used for error correction and error detection.
- **Authentication type:** There are two types of authentication, i.e., 0 and 1. Here, 0 means for none that specifies no authentication is available and 1 means for pwd that specifies the password-based authentication.
- **Authentication:** It is a 32-bit field that contains the actual value of the authentication data.

➤ Border Gateway Protocol

It is an interdomain routing protocol, and it uses the path-vector routing. It is a gateway protocol that is used to exchange routing information among the autonomous system on the internet.

As we know that Border Gateway Protocol works on different autonomous systems, so we should know the history of BGP, types of autonomous systems, etc.

There are many versions of BGP, such as:

- BGP version 1: This version was released in 1989 and is defined in RFC 1105.
- BGP version 2: It was defined in RFC 1163.
- BGP version 3: It was defined in RFC 1267.
- BGP version 4: It is the current version of BGP defined in RFC 1771.

The following are the features of a BGP protocol:

- **Open standard**

It is a standard protocol which can run on any window device.

- **Exterior Gateway Protocol**

It is an exterior gateway protocol that is used to exchange the routing information between two or more autonomous system numbers.

- **InterAS-domain routing**

It is specially designed for inter-domain routing, where interAS-domain routing means exchanging the

routing information between two or more autonomous number system.

- **Supports internet**

It is the only protocol that operates on the internet backbone.

- **Classless**

It is a classless protocol.

- **Incremental and trigger updates**

Like IGP, BGP also supports incremental and trigger updates.

- **Path vector protocol**

The BGP is a path vector protocol. Here, path vector is a method of sending the routes along with routing information.

- **Configure neighborhood relationship**

It sends updates to configure the neighborhood relationship manually. Suppose there are two routers R1 and R2. Then, R1 has to send the configure command saying that you are my neighbor. On the other side, R2 also has to send the configure command to R1, saying that R1 is a neighbor of R1. If both the configure commands match, then the neighborhood relationship will get developed between these two routers.

- **Application layer protocol**

It is an application layer protocol and uses TCP protocol for reliability.

- **Metric**

It has lots of attributes like weight attribute, origin, etc. BGP supports a very rich number of attributes that can affect the path manipulation process.

- **Administrative distance**

If the information is coming from the external autonomous system, then it uses 20 administrative distance. If the information is coming from the same autonomous system, then it uses 200 administrative distance.

➤ Transport Layer

- The transport layer is a 4th layer from the top.
- The main role of the transport layer is to provide the communication services directly to the application processes running on different hosts.
- The transport layer provides a logical communication between application processes running on different hosts. Although the application processes on different hosts are not physically connected, application processes use the logical communication provided by the transport layer to send the messages to each other.
- The transport layer protocols are implemented in the end systems but not in the network routers.
- A computer network provides more than one protocol to the network applications. For example, TCP and UDP are two transport layer protocols that provide a different set of services.

- All transport layer protocols provide multiplexing/demultiplexing service. It also provides services such as reliable data transfer, bandwidth delay guarantees.
- Each of the applications in the application layer has to send a message by using TCP or UDP. The application communicates by using either of these two TCP and UDP will then communicate with the protocol in the internet layer. The applications can communicate to the transport layer. Therefore, we can say that this is a two-way process.

Services provided by the Transport Layer

The services provided by the transport layer are similar to those of the data link layer. The data link layer provides the services within a single network while the transport layer provides services across an internetwork made up of many networks. The data link layer controls the physical layer while the transport layer controls all the lower layers.

The services provided by the transport layer protocols can be divided into five categories:

- End-to-end delivery
- Addressing
- Reliable delivery
- Flow control
- Multiplexing

End-to-end delivery:

The transport layer transmits the entire message to the destination. Therefore, it ensures the end-to-end delivery of an entire message from a source to the destination.

Reliable delivery:

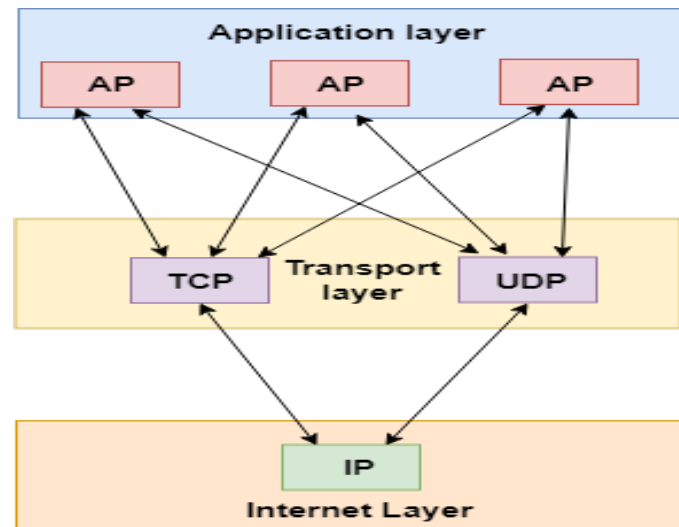
The transport layer provides reliability services by retransmitting the lost and damaged packets.

The reliable delivery has four aspects:

- Error control
- Sequence control
- Loss control
- Duplication control

Error Control

- The primary role of reliability is **Error Control**. In reality, no transmission will be 100 percent error-free delivery. Therefore, transport layer protocols are designed to provide error-free transmission.



- The data link layer also provides the error handling mechanism, but it ensures only node-to-node error-free delivery. However, node-to-node reliability does not ensure the end-to-end reliability.
- The data link layer checks for the error between each network. If an error is introduced inside one of the routers, the error will not be caught by the data link layer. It only detects those errors that have been introduced between the beginning and end of the link. Therefore, the transport layer performs the checking for the errors end-to-end to ensure that the packet has arrived correctly.

Sequence Control

- The second aspect of the reliability is sequence control which is implemented at the transport layer.
- On the sending end, the transport layer is responsible for ensuring that the packets received from the upper layers can be used by the lower layers. On the receiving end, it ensures that the various segments of a transmission can be correctly reassembled.

Loss Control

Loss Control is a third aspect of reliability. The transport layer ensures that all the fragments of a transmission arrive at the destination, not some of them. On the sending end, all the fragments of transmission are given sequence numbers by the transport layer. These sequence numbers allow the receiver's transport layer to identify the missing segment.

Duplication Control

Duplication Control is the fourth aspect of reliability. The transport layer guarantees that no duplicate data arrive at the destination. Sequence numbers are used to identify the lost packets; similarly, it allows the receiver to identify and discard duplicate segments.

Flow Control

Flow control is used to prevent the sender from overwhelming the receiver. If the receiver is overloaded with too much data, then the receiver discards the packets and asking for the retransmission of packets. This increases network congestion and thus, reducing the system performance. The transport layer is responsible for flow control. It uses the sliding window protocol that makes the data transmission more efficient as well as it controls the flow of data so that the receiver does not become overwhelmed. Sliding window protocol is byte oriented rather than frame oriented.

Multiplexing

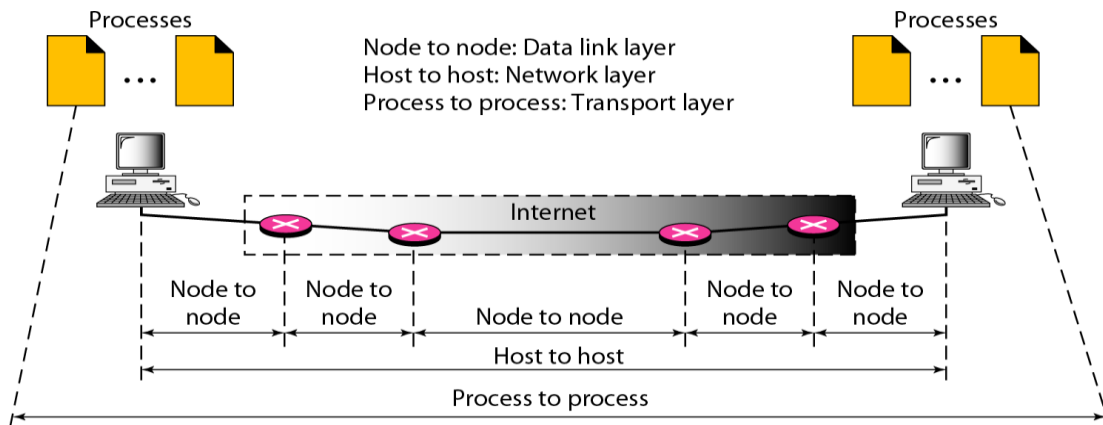
The transport layer uses the multiplexing to improve transmission efficiency.

Multiplexing can occur in two ways:

- **Upward multiplexing:** Upward multiplexing means multiple transport layer connections use the same network connection. To make more cost-effective, the transport layer sends several transmissions bound for the same destination along the same path; this is achieved through upward multiplexing.
- **Downward multiplexing:** Downward multiplexing means one transport layer connection uses the multiple network connections. Downward multiplexing allows the transport layer to split a connection among several paths to improve throughput. This type of multiplexing is used when networks have a low or slow capacity.

Process to Process Delivery:

The data link layer is responsible for delivery of frames between two neighboring nodes over a link. This is called node-to-node delivery. The network layer is responsible for delivery of datagrams between two hosts. This is called host-to-host delivery. Communication on the Internet is not defined as the exchange of data between two nodes or between two hosts. Real communication takes place between two processes (application programs). We need process-to-process delivery. However, at any moment, several processes may be running on the source host and several on the destination host. To complete the delivery, we need a mechanism to deliver data from one of these processes running on the source host to the corresponding process running on the destination host. The transport layer is responsible for process-to-process delivery. It ensures the delivery of a packet, part of a message, from one process to another. Two processes communicate in a client/server relationship, as we will see later. Figure 23.1 shows these three types of deliveries and their domains. The transport layer is responsible for process-to-process delivery.



Client/Server Paradigm

Although there are several ways to achieve process-to-process communication, the most common one is through the client/server paradigm. A process on the local host, called a client, needs services from a process usually on the remote host, called a server.

Both processes (client and server) have the same name. For example, to get the day and time from a remote machine, you need a Daytime client process running on the local host and a Daytime server process running on a remote machine.

Operating systems today support both multiuser and multiprogramming environments. A remote computer can run several server programs at the same time, just as local computers can run one or more client programs at the same time.

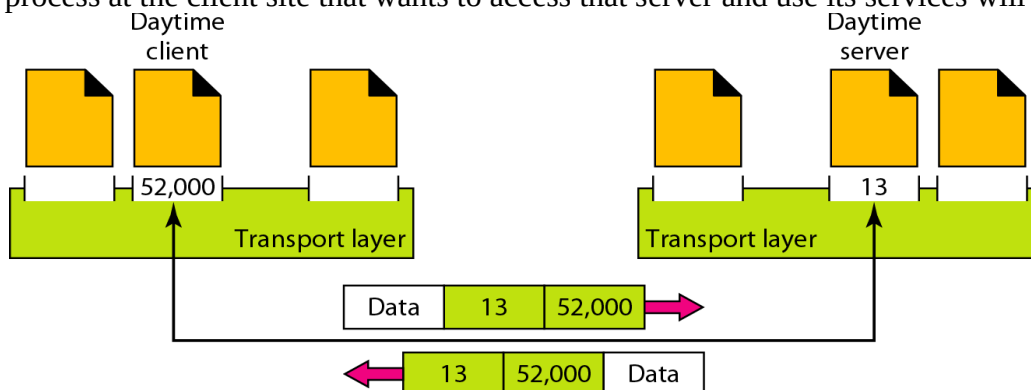
For communication, we must define the following:

1. Local host
2. Local process
3. Remote host
4. Remote process

Addressing Whenever we need to deliver something to one specific destination among many, we need an address. At the data link layer, we need a MAC address to choose one node among several nodes if the connection is not point-to-point. A frame in the data link layer needs a destination MAC address for delivery and a source address for the next node's reply.

At the network layer, we need an IP address to choose one host among millions. A datagram in the network layer needs a destination IP address for delivery and a source IP address for the destination's reply. At the transport layer, we need a transport layer address, called a port number, to choose among multiple processes running on the destination host. A destination port number is needed for delivery; the source port number is needed for the reply. In the Internet mode, port numbers are 16-bit integers between 0 and 65,535. The client program defines itself with a port number, chosen randomly by the transport layer software running on the client host. This is the ephemeral port number.

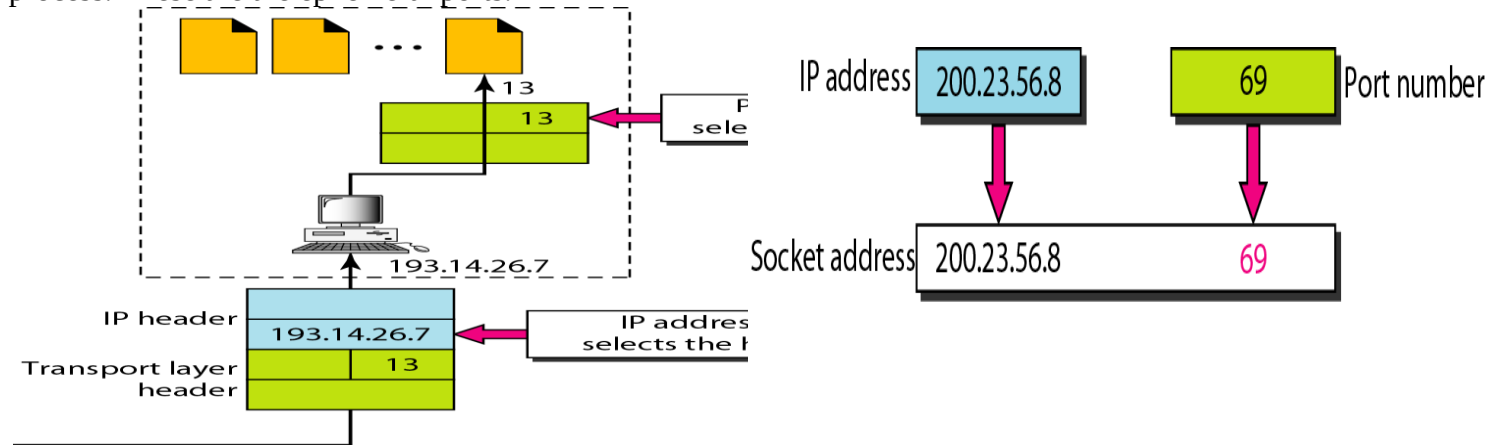
The server process must also define itself with a port number. This port number, however, cannot be chosen randomly. If the computer at the server site runs a server process and assigns a random number as the port number, a process at the client site that wants to access that server and use its services will not know the port number.



IANA Ranges

The IANA (Internet Assigned Number Authority) has divided the port numbers into three ranges: well known, registered and dynamic (or private), as shown in Figure 23.4.

- o Well-known ports. The ports ranging from 0 to 1023 are assigned and controlled by IANA. These are the well-known ports.
- o Registered ports. The ports ranging from 1024 to 49,151 are not assigned or controlled by IANA. They can only be registered with IANA to prevent duplication.
- o Dynamic ports. The ports ranging from 49,152 to 65,535 are neither controlled nor registered. They can be used by any process. These are the ephemeral ports.



Socket Addresses

Process-to-process delivery needs two identifiers, IP address and the port number, at each end to make a connection. The combination of an IP address and a port number is called a socket address. The client socket address defines the client process uniquely just as the server socket address defines the server process uniquely (see Figure 23.5). A transport layer protocol needs a pair of socket addresses: the client socket address and the server socket address. These four pieces of information are part of the IP header and the transport layer protocol header. The IP header contains the IP addresses and the UDP or TCP header contains the port numbers.

Multiplexing and Demultiplexing

Multiplexing

At the sender site, there may be several processes that need to send packets. However, there is only one transport layer protocol at any time. This is a many-to-one relationship and requires multiplexing. The protocol accepts messages from different processes, differentiated by their assigned port numbers. After adding the header, the transport layer puts the packet to the network layer.

Demultiplexing

At the receiver site, the relationship is one-to-many and requires demultiplexing. The transport layer receives datagrams from the network layer. After error checking and dropping of the header, the transport layer delivers the message to the appropriate process based on the port number.

Connectionless Versus Connection-Oriented Service

A transport layer protocol can either be connectionless or connection-oriented.

Connectionless Service In a connectionless service, the packets are sent from one party to another with no need for connection establishment or connection release. The packets are not numbered; they may be delayed or lost or may arrive out of sequence. There is no acknowledgment either. We will see shortly that one of the transport layer protocols in the Internet model, UDP, is connectionless.

Connection -Oriented Service In a connection-oriented service, a connection is first established between the sender and the receiver. Data are transferred. At the end, the connection is released. We will see shortly that TCP and SCTP are connection-oriented protocols.

Reliable Versus Unreliable

The transport layer service can be reliable or unreliable. If the application layer program needs reliability, we use a reliable transport layer protocol by implementing flow and error control at the transport layer. This means a slower and more complex service. On the other hand, if the application program does not need reliability because it uses its own flow and error control mechanism or it needs fast service or the nature of the service does not demand flow and error control (real-time applications), then an unreliable protocol can be used.

TCP

TCP stands for **Transmission Control Protocol**. It is a transport layer protocol that facilitates the transmission of packets from source to destination. It is a connection-oriented protocol that means it establishes the connection prior to the communication that occurs between the computing devices in a network. This protocol is used with an IP protocol, so together, they are referred to as a TCP/IP.

The main functionality of the TCP is to take the data from the application layer. Then it divides the data into several packets, provides numbering to these packets, and finally transmits these packets to the destination. The TCP, on the other side, will reassemble the packets and transmits them to the application layer. As we know that TCP is a connection-oriented protocol, so the connection will remain established until the communication is not completed between the sender and the receiver.

Features of TCP protocol

The following are the features of a TCP protocol:

- **Transport Layer Protocol**

TCP is a transport layer protocol as it is used in transmitting the data from the sender to the receiver.

- **Reliable**

TCP is a reliable protocol as it follows the flow and error control mechanism. It also supports the acknowledgment mechanism, which checks the state and sound arrival of the data. In the acknowledgment mechanism, the receiver sends either positive or negative acknowledgment to the sender so that the sender can get to know whether the data packet has been received or needs to resend.

- **Order of the data is maintained**

This protocol ensures that the data reaches the intended receiver in the same order in which it is sent. It orders and numbers each segment so that the TCP layer on the destination side can reassemble them based on their ordering.

- **Connection-oriented**

It is a connection-oriented service that means the data exchange occurs only after the connection establishment. When the data transfer is completed, then the connection will get terminated.

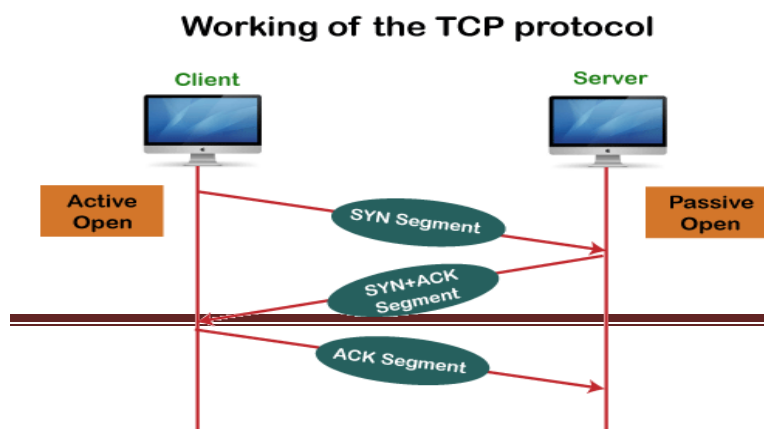
- **Full duplex**

It is a full-duplex means that the data can transfer in both directions at the same time.

- **Stream-oriented**

TCP is a stream-oriented protocol as it allows the sender to send the data in the form of a stream of bytes and also allows the receiver to accept the data in the form of a stream of bytes. TCP creates an environment in which both the sender and receiver are connected by an imaginary tube known as a virtual circuit. This virtual circuit carries the stream of bytes across the internet.

Working of TCP



In TCP, the connection is established by using three-way handshaking. The client sends the segment with its sequence number. The server, in return, sends its segment with its own sequence number as well as the acknowledgement sequence, which is one more than the client sequence number. When

the client receives the acknowledgment of its segment, then it sends the acknowledgment to the server. In this way, the connection is established between the client and the server.

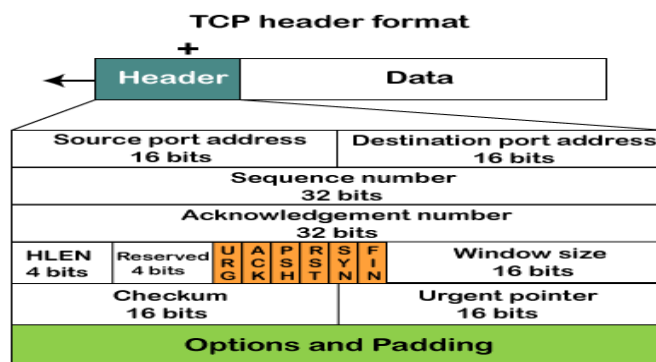
Advantages of TCP

- It provides a connection-oriented reliable service, which means that it guarantees the delivery of data packets. If the data packet is lost across the network, then the TCP will resend the lost packets.
- It provides a flow control mechanism using a sliding window protocol.
- It provides error detection by using checksum and error control by using Go Back or ARP protocol.
- It eliminates the congestion by using a network congestion avoidance algorithm that includes various schemes such as additive increase/multiplicative decrease (AIMD), slow start, and congestion window.

Disadvantage of TCP

It increases a large amount of overhead as each segment gets its own TCP header, so fragmentation by the router increases the overhead.

TCP Header format



- **Source port:** It defines the port of the application, which is sending the data. So, this field contains the source port address, which is 16 bits.
- **Destination port:** It defines the port of the application on the receiving side. So, this field contains the destination port address, which is 16 bits.
- **Sequence number:** This field contains the sequence number of data bytes in a particular session.
- **Acknowledgment number:** When the ACK flag is set, then this contains the next sequence number of the data byte and works as an acknowledgment for the previous data received. For example, if the receiver receives the segment number 'x', then it responds 'x+1' as an acknowledgment number.
- **HLEN:** It specifies the length of the header indicated by the 4-byte words in the header. The size of the header lies between 20 and 60 bytes. Therefore, the value of this field would lie between 5 and 15.
- **Reserved:** It is a 4-bit field reserved for future use, and by default, all are set to zero.

Flags

There are six control bits or flags:

1. **URG:** It represents an urgent pointer. If it is set, then the data is processed urgently.
2. **ACK:** If the ACK is set to 0, then it means that the data packet does not contain an acknowledgment.
3. **PSH:** If this field is set, then it requests the receiving device to push the data to the receiving application without buffering it.
4. **RST:** If it is set, then it requests to restart a connection.
5. **SYN:** It is used to establish a connection between the hosts.
6. **FIN:** It is used to release a connection, and no further data exchange will happen.

- **Window size**
It is a 16-bit field. It contains the size of data that the receiver can accept. This field is used for the flow control between the sender and receiver and also determines the amount of buffer allocated by the receiver for a segment. The value of this field is determined by the receiver.
- **Checksum**
It is a 16-bit field. This field is optional in UDP, but in the case of TCP/IP, this field is mandatory.
- **Urgent pointer**
It is a pointer that points to the urgent data byte if the URG flag is set to 1. It defines a value that will be added to the sequence number to get the sequence number of the last urgent byte.
- **Options**
It provides additional options. The optional field is represented in 32-bits. If this field contains the data less than 32-bit, then padding is required to obtain the remaining bits.

➤ **UDP Protocol**

In computer networking, the UDP stands for User Datagram Protocol. The David P. Reed developed the UDP protocol in 1980. It is defined in RFC 768, and it is a part of the TCP/IP protocol, so it is a standard protocol over the internet. The UDP protocol allows the computer applications to send the messages in the form of datagrams from one machine to another machine over the Internet Protocol (IP) network. The UDP is an alternative communication protocol to the TCP protocol (transmission control protocol). Like TCP, UDP provides a set of rules that governs how the data should be exchanged over the internet. The UDP works by encapsulating the data into the packet and providing its own header information to the packet. Then, this UDP packet is encapsulated to the IP packet and sent off to its destination. Both the TCP and UDP protocols send the data over the internet protocol network, so it is also known as TCP/IP and UDP/IP. There are many differences between these two protocols. UDP enables the process to process communication, whereas the TCP provides host to host communication. Since UDP sends the messages in the form of datagrams, it is considered the best-effort mode of communication. TCP sends the individual packets, so it is a reliable transport medium. Another difference is that the TCP is a connection-oriented protocol whereas, the UDP is a connectionless protocol as it does not require any virtual circuit to transfer the data.

UDP also provides a different port number to distinguish different user requests and also provides the checksum capability to verify whether the complete data has arrived or not; the IP layer does not provide these two services.

Features of UDP protocol

The following are the features of the UDP protocol:

- **Transport layer protocol**
UDP is the simplest transport layer communication protocol. It contains a minimum amount of communication mechanisms. It is considered an unreliable protocol, and it is based on best-effort delivery services. UDP provides no acknowledgment mechanism, which means that the receiver does not send the acknowledgment for the received packets and the sender also does not wait for the acknowledgment for the packet that it has sent.
- **Connectionless**
The UDP is a connectionless protocol as it does not create a virtual path to transfer the data. It does not use the virtual path, so packets are sent in different paths between the sender and the receiver, which leads to the loss of packets or received out of order.
Ordered delivery of data is not guaranteed.
In the case of UDP, the datagrams are sent in some order will be received in the same order is not guaranteed as the datagrams are not numbered.
- **Ports**
The UDP protocol uses different port numbers so that the data can be sent to the correct destination. The port numbers are defined between 0 and 1023.

- **Faster transmission**

UDP enables faster transmission as it is a connectionless protocol, i.e., no virtual path is required to transfer the data. But there is a chance that the individual packet is lost, which affects the transmission quality. On the other hand, if the packet is lost in TCP connection, that packet will be resent, so it guarantees the delivery of the data packets.

- **Acknowledgment mechanism**

The UDP does not have any acknowledgment mechanism, i.e., there is no handshaking between the UDP sender and UDP receiver. If the message is sent in TCP, then the receiver acknowledges that I am ready, then the sender sends the data. In the case of TCP, the handshaking occurs between the sender and the receiver, whereas in UDP, there is no handshaking between the sender and the receiver.

- **Segments are handled independently.**

Each UDP segment is handled individually of others as each segment takes different path to reach the destination. The UDP segments can be lost or delivered out of order to reach the destination as there is no connection setup between the sender and the receiver.

- **Stateless**

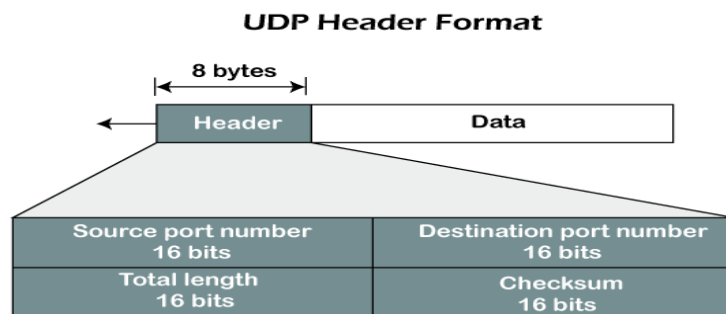
It is a stateless protocol that means that the sender does not get the acknowledgement for the packet which has been sent.

Why do we require the UDP protocol?

As we know that the UDP is an unreliable protocol, but we still require a UDP protocol in some cases. The UDP is deployed where the packets require a large amount of bandwidth along with the actual data. For example, in video streaming, acknowledging thousands of packets is troublesome and wastes a lot of bandwidth. In the case of video streaming, the loss of some packets couldn't create a problem, and it can also be ignored.

UDP Header Format

In UDP, the header size is 8 bytes, and the packet size is up to 65,535 bytes. But this packet size is not possible as the data needs to be encapsulated in the IP datagram, and an IP packet, the header size can be 20 bytes; therefore, the maximum of UDP would be 65,535 minus 20. The size of the data that the UDP packet can carry would be 65,535 minus 28 as 8 bytes for the header of the UDP packet



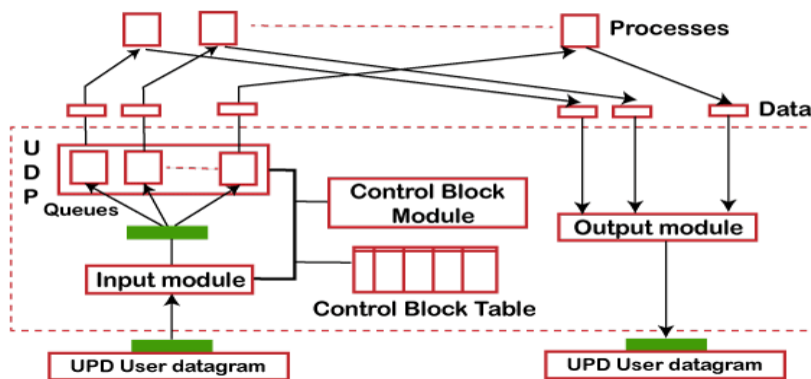
and 20 bytes for IP header.

The UDP header contains four fields:

- **Source port number:** It is 16-bit information that identifies which port is going to send the packet.
- **Destination port number:** It identifies which port is going to accept the information. It is 16-bit information which is used to identify application-level service on the destination machine.
- **Length:** It is 16-bit field that specifies the entire length of the UDP packet that includes the header also. The minimum value would be 8-byte as the size of the header is 8 bytes.
- **Checksum:** It is a 16-bits field, and it is an optional field. This checksum field checks whether the information is accurate or not as there is the possibility that the information can be corrupted while transmission. It is an optional field, which means that it depends upon the application, whether it wants to write the checksum or not. If it does not want to write the checksum, then all the 16 bits are zero; otherwise, it writes the checksum. In UDP, the checksum field is applied to the entire packet, i.e., header as well as data part whereas, in IP, the checksum field is applied to only the header field.

Concept of Queuing in UDP protocol

Concept of Queuing in UDP protocol



In UDP protocol, numbers are used to distinguish the different processes on a server and client. We know that UDP provides a process-to-process communication. The client generates the processes that need services while the server generates the processes that provide services. The queues are available for both the processes, i.e., two queues for each process. The first queue is the incoming queue that receives the messages, and the second one is the outgoing queue that sends the messages. The queue functions when the process is running. If the process is terminated then the queue will also get destroyed.

UDP handles the sending and receiving of the UDP packets with the help of the following components:

- **Input queue:** The UDP packets use a set of queues for each process.
- **Input module:** This module takes the user datagram from the IP, and then it finds the information from the control block table of the same port. If it finds the entry in the control block table with the same port as the user datagram, it enqueues the data.
- **Control Block Module:** It manages the control block table.
- **Control Block Table:** The control block table contains the entry of open ports.
- **Output module:** The output module creates and sends the user datagram.

Several processes want to use the services of UDP. The UDP multiplexes and demultiplexes the processes so that the multiple processes can run on a single host.

Limitations

- It provides an unreliable connection delivery service. It does not provide any services of IP except that it provides process-to-process communication.
- The UDP message can be lost, delayed, duplicated, or can be out of order.
- It does not provide a reliable transport delivery service. It does not provide any acknowledgment or flow control mechanism. However, it does provide error control to some extent.

Advantages

- It produces a minimal number of overheads.

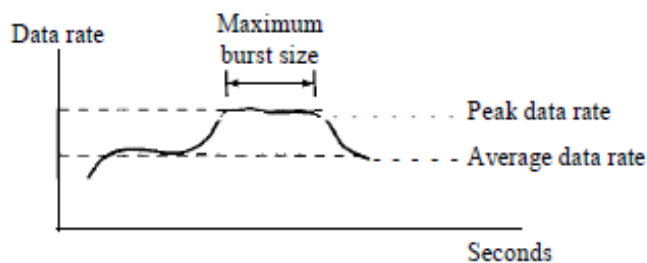
➤ Data Traffic

The main focus of congestion control and quality of service is data traffic. In congestion control we try to avoid congestion. In quality of service, we try to create an appropriate environment for the traffic. So, before talking about congestion control and quality of service, we discuss the data traffic itself.

Traffic Descriptor:

Traffic descriptors are qualitative values that represent a data flow. Following figure shows a traffic flow with some o

Traffic descriptors



values.

Average Data Rate

The average data rate is the number of bits sent during a period of time, divided by the number of seconds in that period. We use the following equation:

Average data rate = amount of data / time

The average data rate is a very useful characteristic of traffic because it indicates the average bandwidth needed by the traffic.

Peak Data Rate

The peak data rate defines the maximum data rate of the traffic. In the above figure it is the maximum y-axis value. The peak data rate is a very important measurement because

it indicates the peak bandwidth that the network needs for traffic to pass through without changing its data flow.

Maximum Burst Size

Although the peak data rate is a critical value for the network, it can usually be ignored if the duration of the peak is very short. For example, if data are flowing steadily at the rate of 1 Mbps with a sudden peak data rate of 2 Mbps for 10 ms, the network probably can handle the situation. However, if the peak data rate lasts 60 ms, there may be a problem for the network. The maximum burst size normally refers to the maximum length of time the traffic is generated at the peak data rate.

Effective Bandwidth

The effective bandwidth is the bandwidth that the network needs to allocate for the flow of traffic. The effective bandwidth is a function of three values: average data rate, peak data rate, and maximum burst size. The calculation of this value is very complex.

Traffic Profiles:

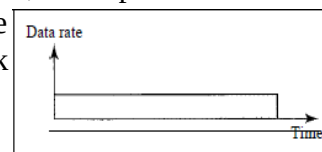
For our purposes, a data flow can have one of the following traffic profiles: constant bit rate, variable bit rate, or bursty as shown in the following figure.

Constant Bit Rate

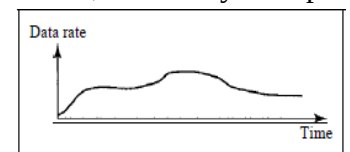
A constant-bit-rate (CBR), or a fixed-rate, traffic model has a data rate that does not change. In this type of flow, the average data rate and the peak data rate are the same. The maximum burst size is not applicable. This type of traffic is very easy for a network to handle since it is predictable. The network knows in advance how much bandwidth to allocate for this type of flow.

Variable Bit Rate

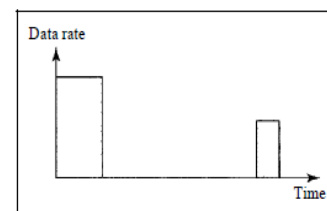
In the variable-bit-rate (VBR) category, the rate of the data flow changes in time, with the changes smooth instead of



a. Constant bit rate



b. Variable bit rate



c. Bursty

and sharp. In this type of flow, the average data

Three traffic profiles: rate and the peak data rate are different. The maximum burst size is usually a small value. This traffic is more difficult to handle than constant-bit-rate traffic, but it normally does not need to be reshaped.

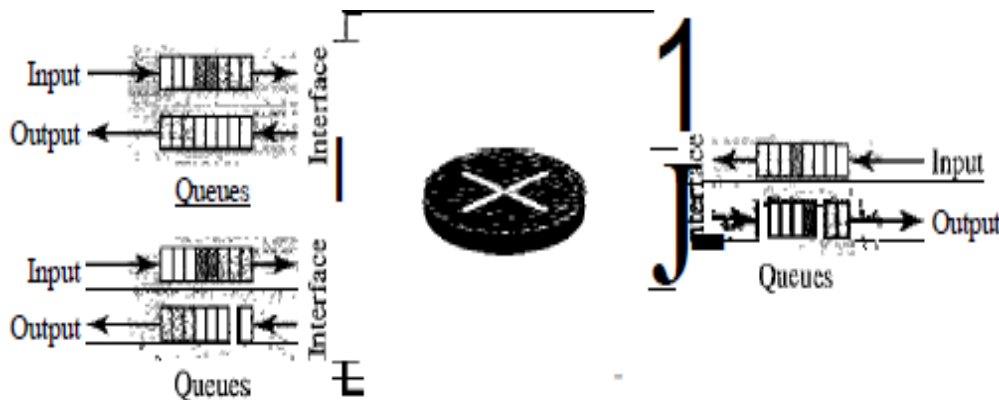
Bursty

In the **bursty data** category, the data rate changes suddenly in a very short time. It may jump from zero, for example, to 100 Mbps in a few microseconds and vice versa. It may also remain at this value for a while. The average bit rate and the peak rate are very different values in this type of flow. The maximum burst size is significant. This is the most difficult type of traffic for a network to handle because the profile is very unpredictable. To handle this type of traffic, the network needs to reshape it, using reshaping techniques. Bursty traffic is one of the main causes of congestion in a network.

➤ CONGESTION

An important issue in a packet-switched network is **congestion**. Congestion in a network may occur if the **load** on the network—the number of packets sent to the network—is greater than the **capacity** of the network—the number of packets the network can handle.

Congestion control refers to the mechanisms and techniques to control the congestion and keep the load below the capacity. We may ask why there is congestion on a network. Congestion happens in any system that involves waiting. For example, congestion happens on a freeway because any abnormality in the flow, such as an accident during rush hour, creates a bottleneck. Congestion in a network or internetwork occurs because routers and switches have queues—buffers that hold the packets before and after processing. A router, for example, has an input queue and an output queue for each interface. When a packet arrives at the incoming interface, it undergoes three steps before departing, as shown in the following figure.



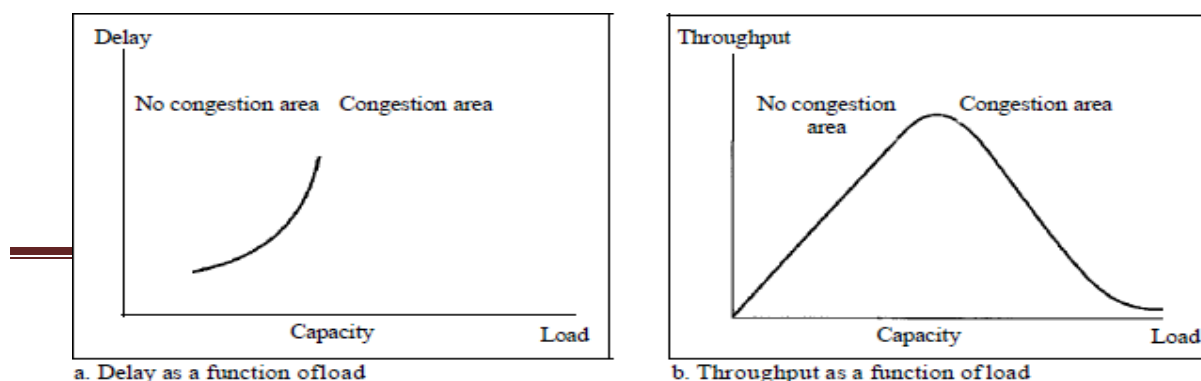
Queues in a router:

1. The packet is put at the end of the input queue while waiting to be checked.
2. The processing module of the router removes the packet from the input queue once it reaches the front of the queue and uses its routing table and the destination address to find the route.
3. The packet is put in the appropriate output queue and waits its turn to be sent. We need to be aware of two issues. First, if the rate of packet arrival is higher than the packet processing rate, the input queues become longer and longer. Second, if the packet departure rate is less than the packet processing rate, the output queues become longer and longer.

Network Performance

Congestion control involves two factors that measure the performance of a network: delay and throughput. The following figure shows these two performance measures as functions of load.

Packet delay and throughput as functions of load



Delay Versus Load

Note that when the load is much less than the capacity of the network, the delay is at a minimum. This minimum delay is composed of propagation delay and processing delay, both of which are negligible. However, when the load reaches network capacity, the delay increases sharply because we now need to add the waiting time in the queues (for all routers along the path) to the total delay. Note that the delay becomes infinite when the load is greater than the capacity. If this is obvious, consider the size of the queues when almost no packet reaches the destination, or reaches the destination with infinite delay; the queues become longer and longer. Delay has a negative effect on the load and consequently the congestion. When a packet is delayed, the source, not receiving the acknowledgment, retransmits the packet, which makes the delay, and the congestion, worse.

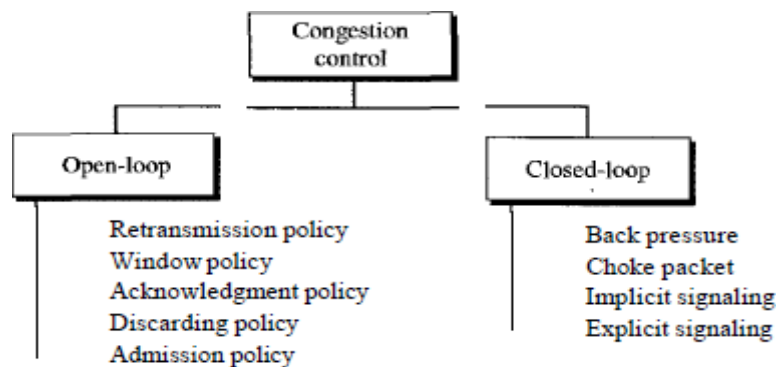
Throughput Versus Load

We defined throughput as the number of bits passing through a point in a second. We can extend that definition from bits to packets and from a point to a network. We can define throughput in a network as the number of packets passing through the network in a unit of time. Notice that when the load is below the capacity of the network, the throughput increases proportionally with the load. We expect the throughput to remain constant after the load reaches the capacity; instead the throughput declines sharply. The reason is the discarding of packets by the routers.

When the load exceeds the capacity, the queues become full and the routers have to discard some packets. Discarding a packet does not reduce the number of packets in the network because the sources retransmit the packets, using timeout mechanisms, when the packets do not reach the destinations.

➤ CONGESTION CONTROL

Congestion control refers to techniques and mechanisms that can either prevent congestion, before it happens, or remove congestion, after it has happened. In general, we can divide congestion control mechanisms into two broad



categories: open-loop congestion control (prevention) and closed-loop congestion control (removal) as shown in the following figure.

Congestion control categories:

Open-Loop Congestion Control

In open-loop congestion control, policies are applied to prevent congestion before it happens. In these mechanisms, congestion control is handled by either the source or the destination. We give a brief list of policies that can prevent congestion.

1. Retransmission Policy

Retransmission is sometimes unavoidable. If the sender feels that a sent packet is lost or corrupted, the packet needs to be retransmitted. Retransmission in general may increase congestion in the network. However, a good retransmission policy can prevent congestion. The retransmission policy and the retransmission timers must be designed to optimize efficiency and at the same time prevent congestion. For example, the retransmission policy used by TCP is designed to prevent or alleviate congestion.

2. Window Policy

The type of window at the sender may also affect congestion. The Selective Repeat window is better than the Go-Back-N window for congestion control. In the Go-Back-N window, when the timer for a packet times out, several packets may be resent, although some may have arrived safe and sound at the receiver. This duplication may make the congestion worse. The Selective Repeat window, on the other hand, tries to send the specific packets that have been lost or corrupted.

3. Acknowledgment Policy

The acknowledgment policy imposed by the receiver may also affect congestion. If the receiver does not acknowledge every packet it receives, it may slow down the sender and help prevent congestion. Several approaches are used in this case. A receiver may send an acknowledgment only if it has a packet to be sent or a special timer expires. A receiver may decide to acknowledge only N packets at a time. We need to know that the acknowledgments are also part of the load in a network. Sending fewer acknowledgments means imposing less load on the network.

4. Discarding Policy

A good discarding policy by the routers may prevent congestion and at the same time may not harm the integrity of the transmission. For example, in audio transmission, if the policy is to discard less sensitive packets when congestion is likely to happen, the quality of sound is still preserved and congestion is prevented or alleviated.

5. Admission Policy

An admission policy, which is a quality-of-service mechanism, can also prevent congestion in virtual-circuit networks. Switches in a flow first check the resource requirement of a flow before admitting it to the network. A router can deny establishing a virtual circuit connection if there is congestion in the network or if there is a possibility of future congestion.

Closed-Loop Congestion Control

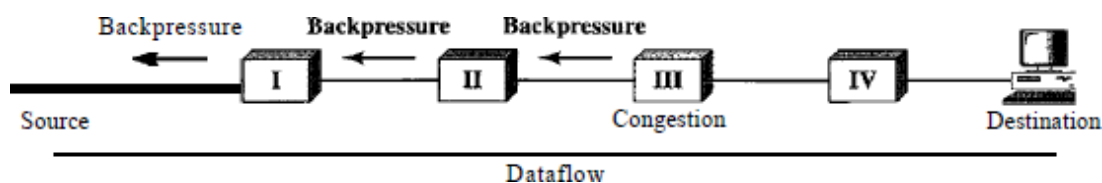
Closed-loop congestion control mechanisms try to alleviate congestion after it happens.

Several mechanisms have been used by different protocols. We describe a few of them here.

1. Backpressure

The technique of backpressure refers to a congestion control mechanism in which a congested node stops receiving data from the immediate upstream node or nodes. This may cause the upstream node or nodes to become congested, and they, in turn, reject data from their upstream nodes or nodes. And so on. Backpressure is a node-to-node congestion control that starts with a node and propagates, in the opposite direction of data flow, to the source. The backpressure technique can be applied only to virtual circuit networks, in which each node knows the upstream node from which a flow of data is coming. Following figure shows the idea of backpressure.

Backpressure method for alleviating congestion:



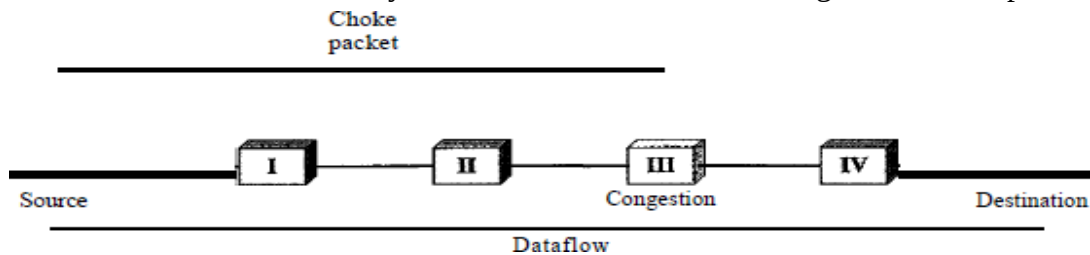
Node III in the figure has more input data than it can handle. It drops some packets in its input buffer and informs node II to slow down. Node II, in turn, may be congested because it is slowing down the output flow of data. If node II is congested, it informs node I to slow down, which in turn may create congestion. If so, node I inform the Source of data to slow down. This, in time, alleviates the congestion. Note that the pressure on node III is moved backward to the source to remove the congestion.

None of the virtual-circuit networks we studied in this book use backpressure. It was, however, implemented in the first virtual-circuit network, X.25. The technique cannot be implemented in a datagram network because in this type of network, a node(router) does not have the slightest knowledge of the upstream router.

2. Choke Packet

A choke packet is a packet sent by a node to the source to inform it of congestion.

Note the difference between the backpressure and choke packet methods. In backpressure, the warning is from one node to its upstream node, although the warning may eventually reach the source station. In the choke packet method, the warning is from the router, which has encountered congestion, to the source station directly. The intermediate nodes through which the packet has travelled are not



warned. We have seen an example of this type of control in ICMP. When a router in the Internet is overwhelmed with IP datagrams, it may discard some of them; but it informs the source host, using a source quench ICMP message. The warning message goes directly to the source station; the intermediate routers, and does not take any action. Following figure shows the idea of a choke packet.

Choke packet:

3.Implicit Signaling

In implicit signaling, there is no communication between the congested node or nodes and the source. The source guesses that there is congestion somewhere in the network from other symptoms. For example, when a source sends several packets and there is no acknowledgment for a while, one assumption is that the network is congested. The delay in receiving an acknowledgment is interpreted as congestion in the network; the source should slow down. We will see this type of signaling when we discuss TCP congestion control later in the chapter.

4.Explicit Signaling

The node that experiences congestion can explicitly send a signal to the source or destination.

The explicit signaling method, however, is different from the choke packet method. In the choke packet method, a separate packet is used for this purpose; in the explicit signaling method, the signal is included in the packets that carry data. Explicit signaling, as we will see in Frame Relay congestion control, can occur in either the forward or the backward direction.

5.Backward Signaling

A bit can be set in a packet moving in the direction opposite to the congestion. This bit can warn the source that there is congestion and that it needs to slow down to avoid the discarding of packets.

6.Forward Signaling

A bit can be set in a packet moving in the direction of the congestion. This bit can warn the destination that there is congestion. The receiver in this case can use policies, such as slowing down the acknowledgments, to alleviate the congestion.

➤ Quality Of Service

Quality of Service (QoS) is a set of technologies that work on a network to guarantee its ability to dependably run high-priority applications and traffic under limited network capacity. QoS technologies accomplish this by providing differentiated handling and capacity allocation to specific flows in network traffic. This enables the network administrator to assign the order in which packets are handled, and the amount of bandwidth afforded to that application or traffic flow

Important flow characteristics of the QoS are given below

1. **Reliability** If a packet gets lost or acknowledgement is not received (at sender), the re-transmission of data will be needed. This decreases the reliability. The importance of the reliability can differ according to the application **For example** E- mail and file transfer need to have a reliable transmission as compared to that of an audio conferencing

2. Delay

Delay of a message from source to destination is a very important characteristic. However, delay can be tolerated differently by the different applications

For example: The time delay cannot be tolerated in audio conferencing (needs a minimum time delay), while the time delay in the e-mail or file transfer has less importance

3. Jitter

The jitter is the variation in the packet delay. If the difference between delays is large, then it is called as **high jitter**. On the contrary, if the difference between delays is small, it is known as **low jitter**.

Example:

Case1: If 3 packets are sent at times 0, 1, and 2 and received at 10, 11, and 12. Here, the delay is same for all packets and it is acceptable for the telephonic conversation

Case2: If 3 packets 0, 1, 2 are sent and received at 31, 34, 39, so the delay is different for all packets. In this case, the time delay is not acceptable for the telephonic conversation

4. Bandwidth

Different applications need the different bandwidth

For example: Video conferencing needs more bandwidth in comparison to that of sending an e-mail.

➤ Techniques to Improve QOS:

Quality of service (QoS) in the case of networking implies the ability of a network to provide reliable service to the traffic over various technologies including Ethernet, wireless, IP, Asynchronous Mode etc.

QOS in case of network congestion must keep in record various elements causing this congestion. It may be due to the reason of low bandwidth or high traffic on a single route. So, routing protocol being used heavily impacts the Quality of service of networking. It depends on how efficient a routing algorithm is to detect the traffic on a particular route and to choose a route accordingly to prevent network congestion on a particular route.

Techniques to improve QOS

Generally, there are four techniques to improve quality of service –

- Scheduling
- Traffic shaping
- Resource Reservation
- Admission Control

Scheduling :

Packets from different flows arrive at a switch or router for processing. A good scheduling technique treats the different flows in a fair and appropriate manner. Several scheduling techniques are designed to improve the quality of service. Three of them here: FIFO queuing, priority queuing, and weighted fair queuing.

- 1) **FIFO Queuing: In first-in, first-out (FIFO)** queuing, packets wait in a buffer (queue) until the node (router or switch) is ready to process them. If the average arrival rate is higher than the average processing rate, the queue will fill up and new packets will be discarded.

- 2) **Priority Queuing:** In priority queuing, packets are first assigned to a priority class. Each priority class has its own queue. The packets in the highest-priority queue are processed first. Packets in the lowest-priority queue are processed last. Note that the system does not stop serving a queue until it is empty.

- 3) **Weighted Fair Queuing:** A better scheduling method is weighted fair queuing. In this technique, the packets are still assigned to different classes and admitted to different queues. The queues, however, are weighted based on the priority of the queues; higher priority means a higher weight. The system processes packets in each queue in a round-robin fashion with the number of packets selected from each queue based on the corresponding weight.

- **Traffic Shaping :**

Traffic shaping is a mechanism to control the amount and the rate of the traffic sent to the network. Two techniques can shape traffic: leaky bucket and token bucket.

1) Leaky Bucket: A technique called leaky bucket can smooth out bursty traffic. Bursty chunks are stored in the bucket and sent out at an average rate.

A FIFO queue holds the packets. If the traffic consists of fixed-size packets, the process removes a fixed number of packets from the queue at each tick of the clock. If the traffic consists of variable-length packets, the fixed output rate must be based on the number of bytes or bits. The following is an algorithm for variable-length packets:

1. Initialize a counter to n at the tick of the clock.
2. If n is greater than the size of the packet, send the packet and decrement the counter by the packet size. Repeat this step until n is smaller than the packet size.
3. Reset the counter and go to step 1.

2) Token Bucket:

The token bucket algorithm allows idle hosts to accumulate credit for the future in the form of tokens. For each tick of the clock, the system sends n tokens to the bucket. The system removes one token for every cell (or byte) of data sent. For example, if n is 100 and the host is idle for 100 ticks, the bucket collects 10,000 tokens. Now the host can consume all these tokens in one tick with 10,000 cells, or the host takes 1,000 ticks with 10 cells per tick. In other words, the host can send bursty data as long as the bucket is not empty

- **Resource Reservation :**

A flow of data needs resources such as a buffer, bandwidth, CPU time, and so on. The quality of service is improved if these resources are reserved beforehand.

- **Admission Control :**

Admission control refers to the mechanism used by a router, or a switch, to accept or reject a flow based on predefined parameters called flow specifications. Before a router accepts a flow for processing, it checks the flow specifications to see if its capacity (in terms of bandwidth, buffer size, CPU speed, etc.) and its previous commitments to other flows can handle the new flow.

➤ Integrated services

In computer networking, **integrated services** or **IntServ** is an architecture that specifies the elements to guarantee quality of service (QoS) on networks. IntServ can for example be used to allow video and sound to reach the receiver without interruption.

IntServ, every router in the system implements IntServ, and every application that requires some kind of QoS guarantee has to make an individual reservation. Flow specs describe what the reservation is for, while RSVP is the underlying mechanism to signal it across the network.

Flow Specification

There are two parts to a flow spec:

- What does the traffic look like? Done in the Traffic Specification part, also known as TSPEC.
- What guarantees does it need? Done in the service Request Specification part, also known as RSPEC.

TSPECs include token bucket algorithm parameters. The idea is that there is a token bucket which slowly fills up with tokens, arriving at a constant rate. Every packet which is sent requires a token, and if there are no tokens, then it cannot be sent. Thus, the rate at which tokens arrive dictates the average rate of traffic flow, while the depth of the bucket dictates how 'bursty' the traffic is allowed to be.

TSPECs typically just specify the token rate and the bucket depth. For example, a video with a refresh rate of 75 frames per second, with each frame taking 10 packets, might specify a token rate of 750 Hz, and a bucket depth of only 10. The bucket depth would be sufficient to accommodate the 'burst' associated with sending an entire frame all at once. On the other hand, a conversation would need a lower token rate, but a much higher bucket depth. This is because there are often pauses in conversations, so they can make do with less tokens by not sending the gaps between words and sentences. However, this means the bucket depth needs to be increased to compensate for the traffic being burstier.

RSPECs specify what requirements there are for the flow: it can be normal internet 'best effort', in which case no reservation is needed. This setting is likely to be used for webpages, FTP, and similar applications. The 'Controlled Load' setting mirrors the performance of a lightly loaded network: there may be occasional glitches when two people access the same resource by chance, but generally both delay and drop rate are fairly constant at the desired rate. This setting is likely to be used by soft QoS applications. The 'Guaranteed' setting gives an absolutely bounded service, where the delay is promised to never go above a desired amount, and packets never dropped, provided the traffic stays within spec.

RSVP

The Resource Reservation Protocol (RSVP): All machines on the network capable of sending QoS data send a PATH message every 30 seconds, which spreads out through the networks. Those who want to listen to them send a corresponding RESV (short for "Reserve") message which then traces the path backwards to the sender. The RESV message contains the flow specs.

The routers between the sender and listener have to decide if they can support the reservation being requested, and, if they cannot, they send a reject message to let the listener know about it. Otherwise, once they accept the reservation they have to carry the traffic.

The routers then store the nature of the flow, and also police it. This is all done in soft state, so if nothing is heard for a certain length of time, then the reader will time out and the reservation will be cancelled. This solves the problem if either the sender or the receiver crash or are shut down incorrectly without first cancelling the reservation. The individual routers may, at their option, police the traffic to check that it conforms to the flow specs.

Problems

In order for IntServ to work, all routers along the traffic path must support it. Furthermore, many states must be stored in each router. As a result, IntServ works on a small-scale, but as the system scales up to larger networks or the Internet, it becomes resource intensive to track all of the reservations.

One way to solve the **scalability** problem is by using a multi-level approach, where per-microflow resource reservation (such as resource reservation for individual users) is done in the edge network, while in the core network resources are reserved for aggregate flows only. The routers that lie between these different levels must adjust the amount of aggregate bandwidth reserved from the core network so that the reservation requests for individual flows from the edge network can be better satisfied

➤ QoS In Switched Networks:

QoS as used in two switched networks: Frame Relay and ATM. These two networks are virtual circuit networks that need a signaling protocol such as RSVP.

QoS in Frame Relay

Four different attributes to control traffic have been devised in frame Relay: access rate, committed burst size B_c , committed information rate (CIR), and excess burst size B_e , these are set during the negotiation between the user and the network.

Access Rate The access rate actually depends on the bandwidth of the channel connecting the user to the network. For example, if the user is connected to a frame relay network by a T-1 line, the access rate is 1.544 Mbps and can never be exceeded.

Committed Burst Size This is the maximum number of bits in a predefined time that the network is committed to transfer without discarding any frame or setting the DE bit. For example, if a B_c of 400 kbits for a period of 4s is granted, the user can send up to 400 kbits during a 4-s interval without worrying about any frame loss.

Committed Information Rate The committed Information rate (CIR) is similar to Committed Burst Size. Except that it defines an average rate in bits per second. The cumulative number of bits sent during the predefined period cannot exceed B_c . It can be calculated by using the following formula: $CIR = B_c / T$ bps. For example, if the B_c is 5 kbits in a period of 5s, the CIR is $5000/5$, or 1 kbps.

Excess Burst Size

Frame relay defines an excess burst size B_e . This is the maximum number of bits in excess of B_c that a user can send during a predefined time. The network is committed to transfer these bits there is no congestion.

User Rate

User who needs to send data faster may exceed the B_c level. This is the level $B_c + B_e$, there is a chance that the frames will reach the destination without being discarded.

QoS in ATM

ATM switches can use traffic policing to enforce the contract. The switch can measure the actual traffic flow and compare it against the agreed-upon traffic envelope. If the switch finds that traffic is outside of the agreed-upon parameters, it can set the cell loss priority (CLP) bit of the offending cells. Setting the CLP bit makes the cell discard-eligible, which means that any switch handling the cell is allowed to drop the cell during periods of congestion. Cell loss and cell delay are ATM QoS parameters; peak cell rate is one of its traffic parameters. QoS and traffic parameters together determine the ATM service category.

The QoS in ATM is based on the class, user-related attributes, and network attributes.

Classes The ATM Forum defines four classes: CBR, VBR, ABR and UBR

constant bit rate (CBR): CBR traffic is characterized by a continuous stream of bits at a steady rate, such as TDM traffic. Class A traffic is low-bandwidth traffic that is highly sensitive to delay and intolerant of cell loss. Carriers use the CBR class of service to provide Circuit Emulation Services (CESs) that emulate TDM like leased-line circuits.

variable bit rate, real time (VBR-RT): VBR-RT traffic has a bursty nature where end-to-end delay is critical. It can be characterized by voice or video applications that use compression, such as interactive videoconferencing.

variable bit rate, non-real time (VBR-NRT): VBR-NRT traffic has a bursty nature in which delay is not so critical, such as video playback, training tapes, and video mail messages.

available bit rate (ABR): ABR traffic can be characterized as bursty LAN traffic and data that is more tolerant of delays and cell loss. ABR is a best-effort service that is a managed service based on minimum cell rate (MCR) and with low cell loss.

unspecified bit rate (UBR): UBR is a best-effort service that does not specify bit rate or traffic parameters and has no QoS guarantees. Originally devised as a way to make use of excess bandwidth, UBR is subject to increased cell loss and the discard of whole packets.

User-Related Attributes

ATM defines two sets of attributes. User-related attributes are those attributes that define how fast the user wants to send data. The following are some user-related attributes.

SCR The sustained cell rate (SCR) is the average cell rate over a long time interval. The actual cell rate may be lower or higher than this value, but the average should be equal to or less than the SCR.

PCR The peak cell rate (PCR) defines the sender's maximum cell rate. The user's cell rate can sometimes reach this peak, as long as the SCR is maintained.

MCR The minimum cell rate (MCR) defines the minimum cell rate acceptable to the sender.

CVDT The cell variation delay tolerance (CVDT) is a measure of the variation in cell transmission times.

Network-Related Attributes

The network-related attributes are those that define characteristics of the network. The following are some network-related attributes:

CLR The cell loss ratio (CLR) defines the fraction of cells lost (or delivered so late that they are considered lost) during transmission.

CTD The cell Transfer delay (CTD) is the average time needed for a cell to travel from source to destination. The maximum CTD and the minimum CTD are also considered attributes.

CDV The cell delay variation (CDV) is the difference between the CTD maximum and the CTD minimum.

CER The cell error ratio (CER) defines the fraction of the cells delivered in error.

- **Security:** **Security** in networking is based on cryptography, the science and art of transforming messages to make them secure and immune to attack. Cryptography can provide confidentiality, integrity, authentication, and nonrepudiation of messages, Cryptography can also provide entity authentication

➤ Introduction

Cryptography refers to the science and art of transforming messages to make them secure and immune to attacks. It is a method of storing and transmitting data in a particular form so that only those for whom it is intended can read and process it. Cryptography not only protects data from theft or alteration but can also be used for user authentication.

Components

There are various components of cryptography which are as follows –

Plaintext and Ciphertext

The original message, before being transformed, is called plaintext. After the message is transformed, it is called ciphertext. An encryption algorithm transforms the plaintext into ciphertext; a decryption algorithm transforms the cipher text back into plaintext. The sender uses an encryption algorithm, and the receiver uses a decryption algorithm.

Cipher

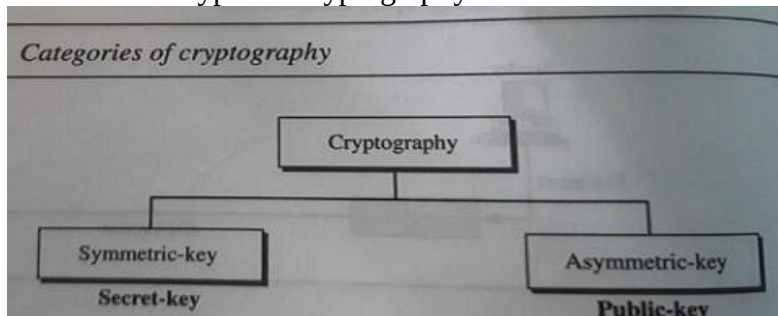
We refer to encryption and decryption algorithms as ciphers. The term cipher is also used to refer to different categories of algorithms in cryptography. This is not to say that every sender-receiver pair needs their very own unique cipher for secure communication. On the contrary, one cipher can serve millions of communicating pairs.

Key

A key is a number (or a set of numbers) that the cipher, as an algorithm, operates on. To encrypt a message, we need an encryption algorithm, an encryption key, and plaintext. These create the ciphertext. To decrypt a message, we need a decryption algorithm, a decryption key, and the ciphertext. These reveal the original plaintext.

Types

There are two types of cryptography which are as follows



Symmetric Key Cryptography

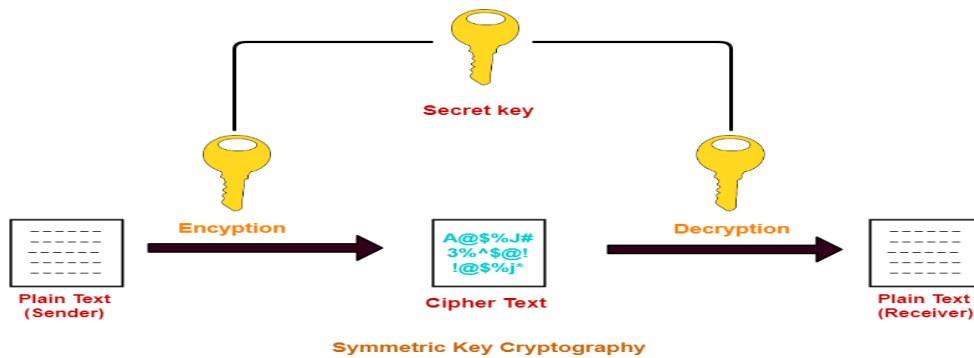
In symmetric-key cryptography, the same key is used by both parties. The sender uses this key and an encryption algorithm to encrypt data; the receiver uses the same key and the corresponding decryption algorithm to decrypt the data.

In this technique,

- Both sender and receiver uses a common key to encrypt and decrypt the message.
- This secret key is known only to the sender and to the receiver.
- It is also called as **secret key cryptography**.

Working-

The message exchange using symmetric key cryptography involves the following steps-



- Before starting the communication, sender and receiver shares the secret key.
- This secret key is shared through some external means.
- At sender side, sender encrypts the message using his copy of the key.
- The cipher text is then sent to the receiver over the communication channel.
- At receiver side, receiver decrypts the cipher text using his copy of the key.
- After decryption, the message converts back into readable format.

The advantages of symmetric key algorithms are-

- They are efficient.
- They take less time to encrypt and decrypt the message.

Disadvantages-

In symmetric key cryptography, The number of keys required is very large.

- Each pair of users require a unique secret key.
- If N people in the world wants to use this technique, then there needs to be $N(N-1) / 2$ secret keys.
- For 1 million people to communicate, a half billion secret keys would be needed.
- Sharing the secret key between the sender and receiver is an important issue.
- While sharing the key, attackers might intrude.

Asymmetric-Key Cryptography

In asymmetric or public-key cryptography, there are two keys: a private key and a public key. The private key is kept by the receiver. The public key is announced to the public.

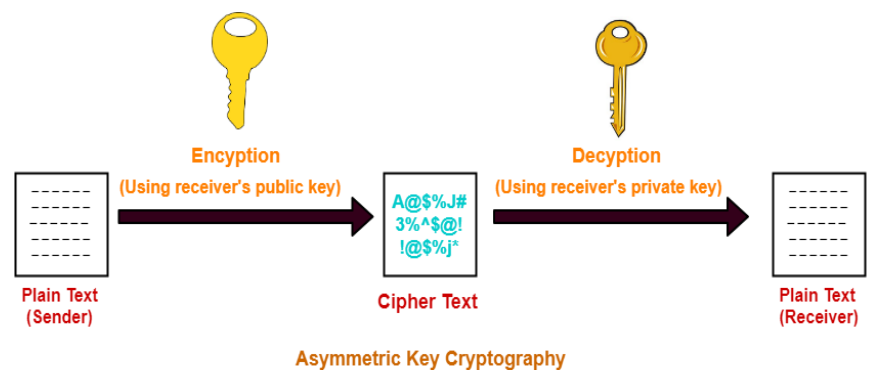
In public-key encryption/decryption, the public key that is used for encryption is different from the private key that is used for decryption. The public key is available to the public, and the private key is available only to an individual.

In this technique,

- Sender and receiver use different keys to encrypt and decrypt the message.
- It is called so because sender and receiver use different keys.
- It is also called as **public key cryptography**.

Working-

The message exchange using public key cryptography involves the following steps-



Asymmetric Encryption Algorithms



Advantages-

The advantages of public key cryptography are-

- It is more robust.
- It is less susceptible to third-party security breach attempts.

Disadvantages-

The disadvantages of public key cryptography are-

- It involves high computational requirements.
- It is slower than symmetric key cryptography.

Number of Keys Required-

To use public key cryptography,

- Each individual requires two keys- one public key and one private key.
- For n individuals to communicate, number of keys required = $2 \times n = 2n$ keys.

Asymmetric Encryption Algorithms-

The famous asymmetric encryption algorithms are-

RSA Algorithm

1. Diffie-Hellman Key Exchange

RSA Algorithm-

Let-Public key of the receiver = (e, n)

- Private key of the receiver = (d, n)

Then, RSA Algorithm works in the following steps-

At sender side,

- Sender represents the message to be sent as an integer between 0 and $n-1$.
- Sender encrypts the message using the public key of receiver.
- It raises the plain text message 'P' to the e^{th} power modulo n .
- This converts the message into cipher text 'C'. $C = P^e \bmod n$

The cipher text 'C' is sent to the receiver over the communication channel.

At receiver side,

- Receiver decrypts the cipher text using his private key.
- It raises the cipher text 'C' to the d^{th} power modulo n .
- This converts the cipher text back into the plain text 'P'. $P = C^d \bmod n$

Diffie Hellman Key Exchange-

As the name suggests,

- This algorithm is used to exchange the secret key between the sender and the receiver.
- This algorithm facilitates the exchange of secret key without actually transmitting it.

Diffie Hellman Key Exchange Algorithm-Let-

- Private key of the sender = X_s
- Public key of the sender = Y_s
- Private key of the receiver = X_r
- Public key of the receiver = Y_r

Using Diffie Hellman Algorithm, the key is exchanged in the following steps-

- One of the parties choose two numbers 'a' and 'n' and exchange with the other party.
 - 'a' is the primitive root of prime number 'n'.
 - After this exchange, both the parties know the value of 'a' and 'n'.
-

Both the parties already know their own private key.

- Both the parties calculate the value of their public key and exchange with each other.

Sender calculate its public key as-

$$Y_s = a^{X_s} \bmod n$$

Receiver calculate its public key as- $Y_r = a^{X_r} \bmod n$

Both the parties receive public key of each other.

- Now, both the parties calculate the value of secret key.

Sender calculates secret key as-

$$\text{Secret key} = (Y_r)^{X_s} \bmod n$$

Receiver calculates secret key as- $\text{Secret key} = (Y_s)^{X_r} \bmod n$

Finally, both the parties obtain the same value of secret key.

➤ Digital Signature

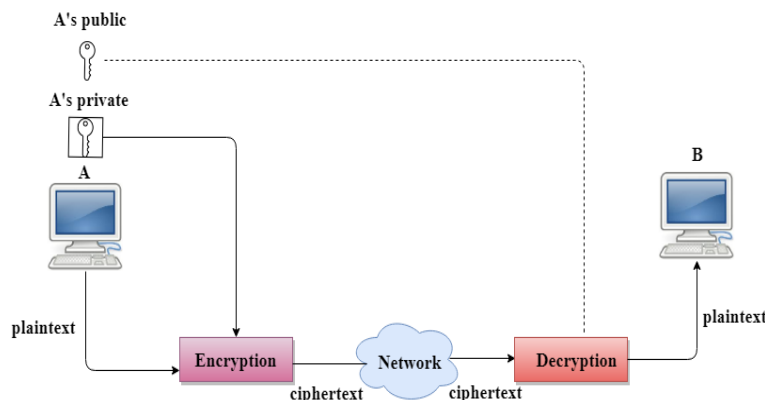
The Digital Signature is a technique which is used to validate the authenticity and integrity of the message. We know that there are four aspects of security: privacy, authentication, integrity, and non-repudiation. We have already discussed the first aspect of security and other three aspects can be achieved by using a digital signature.

The basic idea behind the Digital Signature is to sign a document. When we send a document electronically, we can also sign it. We can sign a document in two ways: to sign a whole document and to sign a digest.

Signing the Whole Document

- In Digital Signature, a public key encryption technique is used to sign a document. However, the roles of a public key and private key are different here. The sender uses a private key to encrypt the message while the receiver uses the public key of the sender to decrypt the message.
- In Digital Signature, the private key is used for encryption while the public key is used for decryption.
- Digital Signature cannot be achieved by using secret key encryption.

Digital Signature is used to achieve the following three aspects:



- Integrity:** The Digital Signature preserves the integrity of a message because, if any malicious attack intercepts a message and partially or totally changes it, then the decrypted message would be impossible.

- Authentication:** We can use the following reasoning to show how the message is authenticated. If an intruder (user X) sends a message pretending that it is coming from someone else (user A), user X uses her own private key to encrypt the message. The

message is decrypted by using the public key of user A. Therefore this makes the message unreadable. Encryption with X's private key and decryption with A's public key results in garbage value.

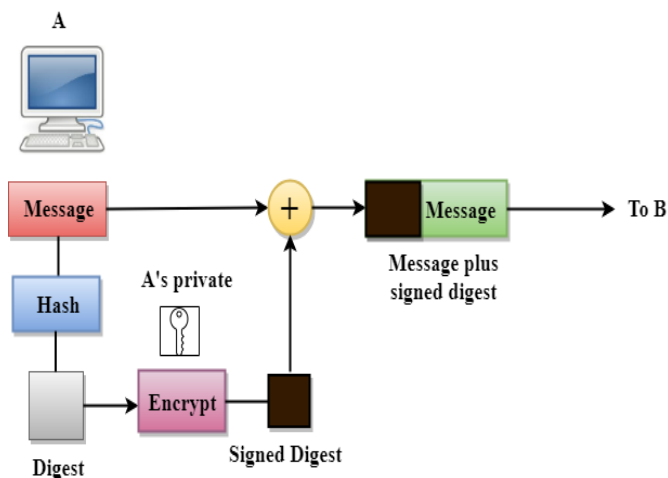
- Non-Repudiation:** Digital Signature also provides non-repudiation. If the sender denies sending the message, then her private key corresponding to her public key is tested on the plaintext. If the decrypted message is the same as the original message, then we know that the sender has sent the message.

- Public key encryption is efficient if the message is short. If the message is long, a public key encryption is inefficient to use. The solution to this problem is to let the sender sign a digest of the document instead of the whole document.
- The sender creates a miniature version (digest) of the document and then signs it, the receiver checks the signature of the miniature version.
- The hash function is used to create a digest of the message. The hash function creates a fixed-size digest from the variable-length message.
- The two most common hash functions used: MD5 (Message Digest 5) and SHA-1 (Secure Hash Algorithm 1). The first one produces 120-bit digest while the second one produces a 160-bit digest.
- A hash function must have two properties to ensure the success:
 - First, the digest must be one way, i.e., the digest can only be created from the message but not vice versa.
 - Second, hashing is a one-to-one function, i.e., two messages should not create the same digest.

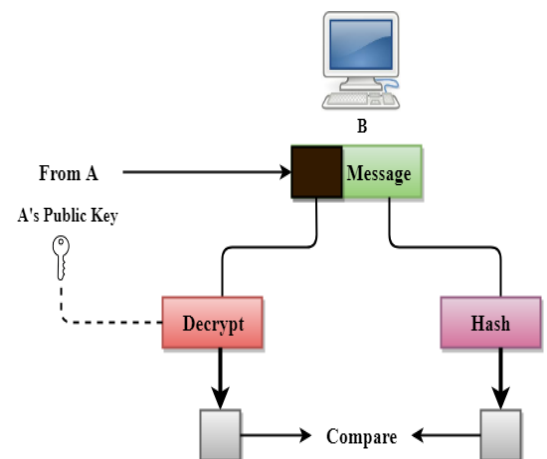
Following are the steps taken to ensure security:

- The miniature version (digest) of the message is created by using a hash function.
- The digest is encrypted by using the sender's private key.
- After the digest is encrypted, then the encrypted digest is attached to the original message and sent to the receiver.
- The receiver receives the original message and encrypted digest and separates the two. The receiver implements the hash function on the original message to create the second digest, and it also decrypts the received digest by using the public key of the sender. If both the digests are same, then all the aspects of security are preserved.

At the Sender site



At the Receiver site



➤ Application Layer

The application layer in the OSI model is the closest layer to the end user which means that the application layer and end user can interact directly with the software application. The application layer programs are based on client and servers.

The Application layer includes the following functions:

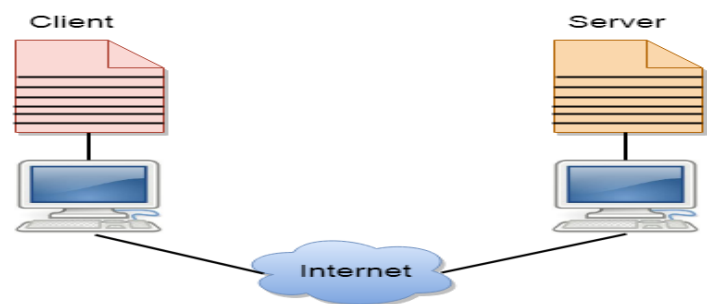
- **Identifying communication partners:** The application layer identifies the availability of communication partners for an application with data to transmit.
- **Determining resource availability:** The application layer determines whether sufficient network resources are available for the requested communication.
- **Synchronizing communication:** All the communications occur between the applications requires cooperation which is managed by an application layer.

Services of Application Layers

- **Network Virtual terminal:** An application layer allows a user to log on to a remote host. To do so, the application creates a software emulation of a terminal at the remote host. The user's computer talks to the software terminal, which in turn, talks to the host. The remote host thinks that it is communicating with one of its own terminals, so it allows the user to log on.
- **File Transfer, Access, and Management (FTAM):** An application allows a user to access files in a remote computer, to retrieve files from a computer and to manage files in a remote computer. FTAM defines a hierarchical virtual file in terms of file structure, file attributes and the kind of operations performed on the files and their attributes.
- **Addressing:** To obtain communication between client and server, there is a need for addressing. When a client made a request to the server, the request contains the server address and its own address. The server response to the client request, the request contains the destination address, i.e., client address. To achieve this kind of addressing, DNS is used.
- **Mail Services:** An application layer provides Email forwarding and storage.
- **Directory Services:** An application contains a distributed database that provides access for global information about various objects and services.
Authentication: It authenticates the sender or receiver's message or both.

➤ Client and Server model

- A client and server networking model is a model in which computers such as servers provide the network services to the other computers such as clients to perform a user based tasks. This model is known as client-server networking model.
- The application programs using the client-server model should follow the given below strategies:



An application program is known as a client program, running on the local machine that requests for a service from an application program known as a server program, running on the remote machine.

- A client program runs only when it requests for a service from the server while the server program runs all time as it does not know when its service is required.
- A server provides a service for many clients not just for a single client. Therefore, we can say that client-server follows the many-to-one relationship. Many clients can use the service of one server.

- Services are required frequently, and many users have a specific client-server application program. For example, the client-server application program allows the user to access the files, send e-mail, and so on. If the services are more customized, then we should have one generic application program that allows the user to access the services available on the remote computer.

Client

A client is a program that runs on the local machine requesting service from the server. A client program is a finite program means that the service started by the user and terminates when the service is completed.

Server

A server is a program that runs on the remote machine providing services to the clients. When the client requests for a service, then the server opens the door for the incoming requests, but it never initiates the service.

A server program is an infinite program means that when it starts, it runs infinitely unless the problem arises. The server waits for the incoming requests from the clients. When the request arrives at the server, then it responds to the request.

Advantages of Client-server networks:

- **Centralized:** Centralized back-up is possible in client-server networks, i.e., all the data is stored in a server.
- **Security:** These networks are more secure as all the shared resources are centrally administered.
- **Performance:** The use of the dedicated server increases the speed of sharing resources. This increases the performance of the overall system.
- **Scalability:** We can increase the number of clients and servers separately, i.e., the new element can be added, or we can add a new node in a network at any time.

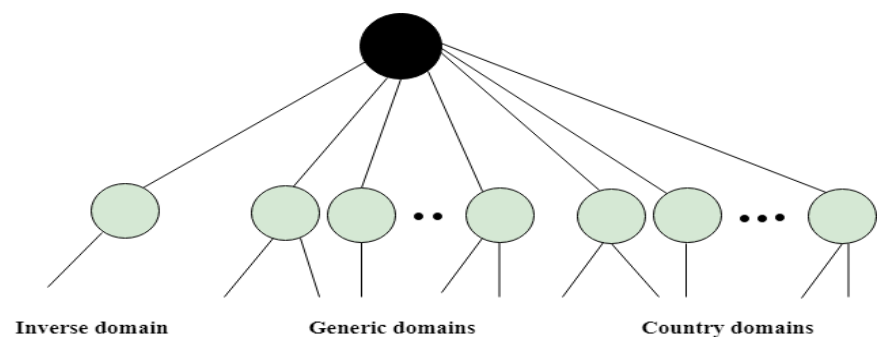
Disadvantages of Client-Server network:

- **Traffic Congestion** is a big problem in Client/Server networks. When a large number of clients send requests to the same server may cause the problem of Traffic congestion.
- It does not have a robustness of a network, i.e., when the server is down, then the client requests cannot be met.
- A client/server network is very decisive. Sometimes, regular computer hardware does not serve a certain number of clients. In such situations, specific hardware is required at the server side to complete the work.
- Sometimes the resources exist in the server but may not exist in the client. For example, If the application is web, then we cannot take the print out directly on printers without taking out the print view window on the web

➤ DNS,DNS in the Internet

An application layer protocol defines how the application processes running on different systems, pass the messages to each other.

- DNS stands for Domain Name System.
- DNS is a directory service that provides a mapping between the name of a host on the network and its numerical address.



- DNS is required for the functioning of the internet.
- Each node in a tree has a domain name, and a full domain name is a sequence of symbols specified by dots.
- DNS is a service that translates the domain name into IP addresses. This allows the users of networks to utilize user-friendly names when looking for other hosts instead of remembering the IP addresses.
- For example, suppose the FTP site at EduSoft had an IP address of 132.147.165.50, most people would reach this site by specifying ftp.EduSoft.com. Therefore, the domain name is more reliable than IP address.

DNS is a TCP/IP protocol used on different platforms. The domain name space is divided into three different sections: generic domains, country domains, and inverse domain.

Generic Domains

- It defines the registered hosts according to their generic behavior.
- Each node in a tree defines the domain name, which is an index to the DNS database.
- It uses three-character labels, and these labels describe the organization type.

Label	Description
aero	Airlines and aerospace companies
biz	Businesses or firms
com	Commercial Organizations
coop	Cooperative business Organizations
edu	Educational institutions
gov	Government institutions
info	Information service providers
int	International Organizations
mil	Military groups
museum	Museum & other nonprofit organizations
name	Personal names
net	Network Support centers
org	Nonprofit Organizations
pro	Professional individual Organizations

Country Domain

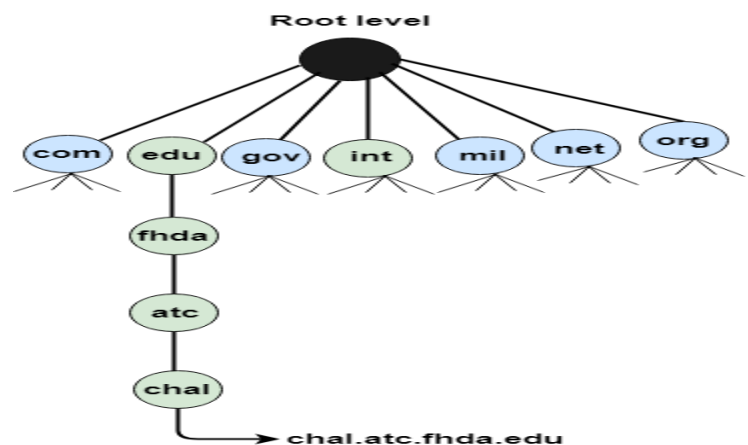
The format of country domain is same as a generic domain, but it uses two-character country abbreviations (e.g., us for the United States) in place of three character organizational abbreviations.

Inverse Domain

The inverse domain is used for mapping an address to a name. When the server has received a request from the client, and the server contains the files of only authorized clients. To determine whether the client is on the authorized list or not, it sends a query to the DNS server and ask for mapping an address to the name.

Working of DNS

- DNS is a client/server network communication protocol. DNS clients send requests to the server while DNS servers send responses to the client.



- Client requests contain a name which is converted into an IP address known as a forward DNS lookups while requests containing an IP address which is converted into a name known as reverse DNS lookups.
- DNS implements a distributed database to store the name of all the hosts available on the internet.
- If a client like a web browser sends a request containing a hostname, then a piece of software such as **DNS resolver** sends a request to the DNS server to obtain the IP address of a hostname. If DNS server does not contain the IP address associated with a hostname, then it forwards the request to another DNS server. If IP address has arrived at the resolver, which in turn completes the request over the internet protocol

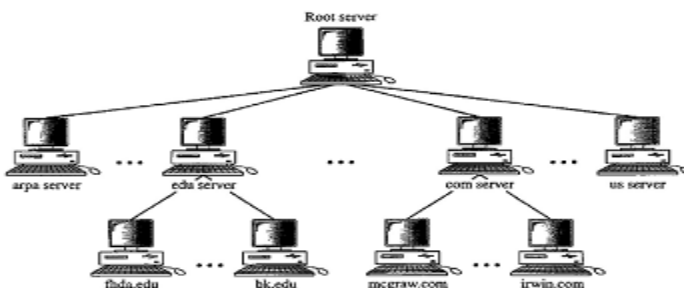
➤ **Distribution of NAMESPACE**

- The information contained in the domain name space must be stored. However, it is very inefficient and also unreliable to have just one computer store such a huge amount of information. **It is inefficient because responding to requests from allover the world places a heavy load on the system. It is not unreliable because any failure makes the data inaccessible.**

The solution to these problems is to distribute the information among many computers called **DNS servers**. One way to do this is to divide the whole space into many domains based on the first level.

- In other words, we let the root stand alone and create as many domains (subtrees) as there are first-level nodes. Because a domain created in this way could be very large, **DNS allows domains to be divided further into smaller domains (subdomains).**
- Each server can be responsible (authoritative) for either a large or a small domain. **In other words, we have a hierarchy of servers in the same way that we have a hierarchy of names.**

Hierarchy of name servers



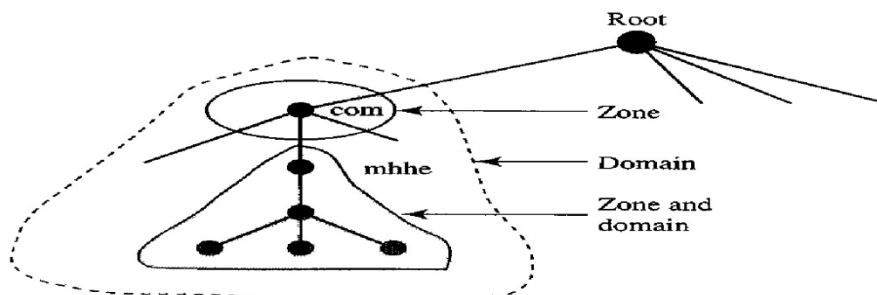
- Since the complete domain name hierarchy cannot be stored on a single server, it is divided among many servers. What a server is responsible for or **has authority over is called a zone.**
- We can define a zone as a contiguous part of the entire tree. If a server accepts responsibility for a domain and does not divide the domain into smaller domains, **the domain and the zone refer to the same thing.**
- The server makes a database called a **zone file and keeps all the information for every node under that domain.**

However, if a server divides its domain into subdomains and delegates part of its authority to other servers, **domain and zone refer to different things.**

- The information about the nodes in the subdomains is stored in the servers at the lower levels, with the original server keeping **some sort of reference to these lower-level servers.**
- Of course the original server does not free itself from responsibility totally: It still has a zone, but the detailed information is kept by the lower-level servers (see Figure below).

- A server can also divide part of its domain and delegate responsibility but still keep part of the domain for itself. In this case, its zone is made of detailed information for the part of the domain that is not delegated and references to those parts that are delegated.
- **A root server is a server whose zone consists of the whole tree.** A root server usually does not store any information about domains but delegates its authority to other servers, keeping references to those servers. There are several root servers, each covering the whole domain name space. **The servers are distributed all around the world.**
- DNS defines two types of servers: **primary and secondary**. A primary server is a server that stores a file about the zone for which it is an authority. It is responsible for creating, maintaining, and updating the zone file. **It stores the zone file on a local disk.**
- A secondary server is a server that transfers the complete information about a zone from another server (primary or secondary) and stores the file on its local disk. The secondary server **neither creates nor updates the zone files.**
If updating is required, it must be done by the primary server, **which sends the updated version to the secondary.**
- The primary and secondary servers are both- authoritative for the zones they serve. **The idea is not to put the secondary server at a lower level of authority but to create redundancy for the data so that if one server fails, the other can continue serving clients.** Note also that a server can be a primary server for a specific zone and a secondary server for another zone. Therefore, when we refer to a server as a primary or secondary server, we should be careful to which zone we refer.

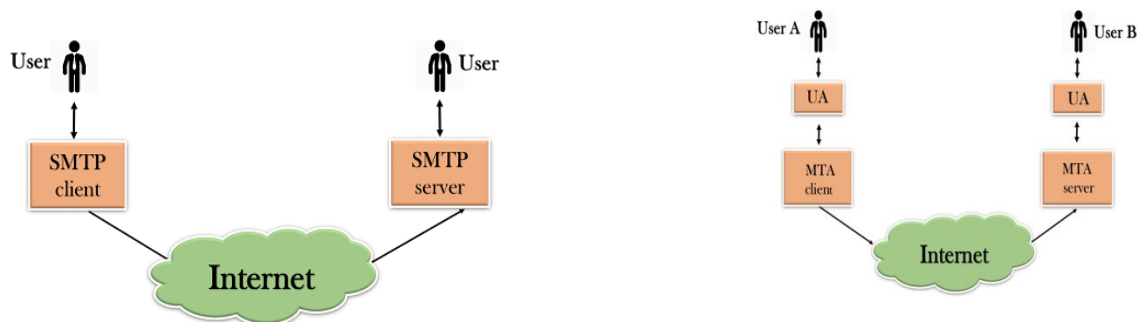
Zones and domains



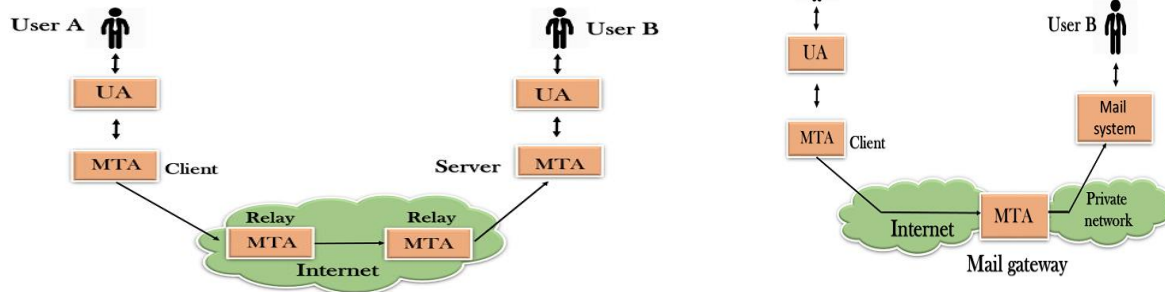
➤ **SMTP**

- SMTP stands for Simple Mail Transfer Protocol.
- SMTP is a set of communication guidelines that allow software to transmit an electronic mail over the internet is called **Simple Mail Transfer Protocol**.
- It is a program used for sending messages to other computer users based on e-mail addresses.
- It provides a mail exchange between users on the same or different computers, and it also supports:
 - It can send a single message to one or more recipients.
 - Sending message can include text, voice, video or graphics.
 - It can also send the messages on networks outside the internet.
- The main purpose of SMTP is used to set up communication rules between servers. The servers have a way of identifying themselves and announcing what kind of communication they are trying to perform. They also have a way of handling the errors such as incorrect email address. For example, if the recipient address is wrong, then receiving server reply with an error message of some kind.

Components of SMTP



- First, we will break the SMTP client and SMTP server into two components such as user agent (UA) and mail transfer agent (MTA). The user agent (UA) prepares the message, creates the envelope and then puts the message in the envelope. The mail transfer agent (MTA) transfers this mail across the internet.
- SMTP allows a more complex system by adding a relaying system. Instead of just having one MTA at sending side and one at receiving side, more MTAs can be added, acting either as a client or server to relay



- The relaying system without TCP/IP protocol can also be used to send the emails to users, and this is achieved by the use of the mail gateway. The mail gateway is a relay MTA that can be used to receive an email.

Working of SMTP

- Composition of Mail:** A user sends an e-mail by composing an electronic mail message using a Mail User Agent (MUA). Mail User Agent is a program which is used to send and receive mail. The message contains two parts: body and header. The body is the main part of the message while the header includes information such as the sender and recipient address. The header also includes descriptive information such as the subject of the message. In this case, the message body is like a letter and header is like an envelope that contains the recipient's address.
- Submission of Mail:** After composing an email, the mail client then submits the completed e-mail to the SMTP server by using SMTP on TCP port 25.
- Delivery of Mail:** E-mail addresses contain two parts: username of the recipient and domain name. For example, vivek@gmail.com, where "vivek" is the username of the recipient and "gmail.com" is the domain name. If the domain name of the recipient's email address is different from the sender's domain name, then MSA will send the mail to the Mail Transfer Agent (MTA). To relay the email, the MTA will find the target domain. It checks the MX record from Domain Name System to obtain the target domain. The MX record contains the domain name and IP address of the recipient's domain. Once the record is located, MTA connects to the exchange server to relay the message.

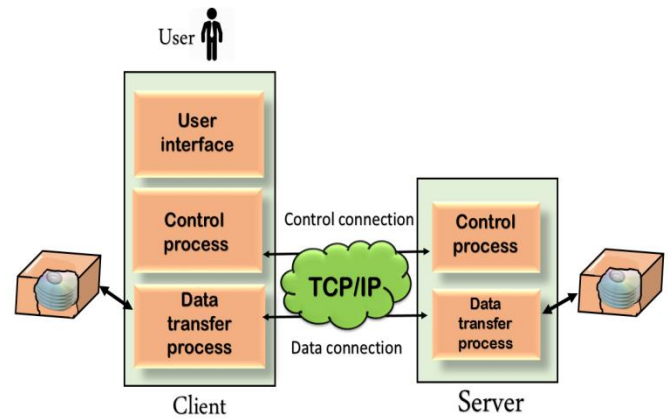
4. **Receipt and Processing of Mail:** Once the incoming message is received, the exchange server delivers it to the incoming server (Mail Delivery Agent) which stores the e-mail where it waits for the user to retrieve it.
5. **Access and Retrieval of Mail:** The stored email in MDA can be retrieved by using MUA (Mail User Agent). MUA can be accessed by using login and password

➤ File Transfer

Transferring files from one computer to another are one of the most common tasks expected from a networking or internetworking environment. The greatest volume of data exchange in the Internet today is due to file transfer, using one popular protocol involved in transferring files: File Transfer Protocol (FTP).

➤ FTP

- FTP stands for File transfer protocol.
- FTP is a standard internet protocol provided by TCP/IP used for transmitting the files from one host to another.
- It is mainly used for transferring the web page files from their creator to the computer that acts as a server for other computers on the internet.
- It is also used for downloading the files to computer from other servers.



Objectives of FTP

- It provides the sharing of files.
- It is used to encourage the use of remote computers.
- It transfers the data more reliably and efficiently.

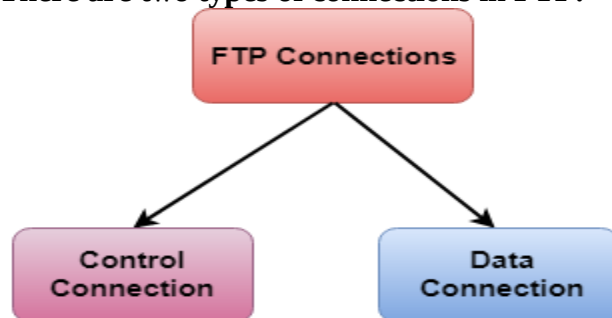
Why FTP?

Although transferring files from one system to another is very simple and straightforward, but sometimes it can cause problems. For example, two systems may have different file conventions. Two systems may have different ways to represent text and data. Two systems may have different directory structures. FTP protocol overcomes these problems by establishing two connections between hosts. One connection is used for data transfer, and another connection is used for the control connection.

Mechanism of FTP

The above figure shows the basic model of the FTP. The FTP client has three components: the user interface, control process, and data transfer process. The server has two components: the server control process and the server data transfer process.

There are two types of connections in FTP:



- **Control Connection:** The control connection uses very simple rules for communication. Through control connection, we can transfer a line of command or line of response at a time. The control connection is made between the control processes. The control connection remains connected during the entire interactive FTP session.

- **Data Connection:** The Data Connection uses very complex rules as data types may vary. The data

connection is made between data transfer processes. The data connection opens when a command comes for transferring the files and closes when the file is transferred.

FTP Clients

- FTP client is a program that implements a file transfer protocol which allows you to transfer files between two hosts on the internet.
- It allows a user to connect to a remote host and upload or download the files.
- It has a set of commands that we can use to connect to a host, transfer the files between you and your host and close the connection.
- The FTP program is also available as a built-in component in a Web browser. This GUI based FTP client makes the file transfer very easy and also does not require to remember the FTP commands.

Advantages of FTP:

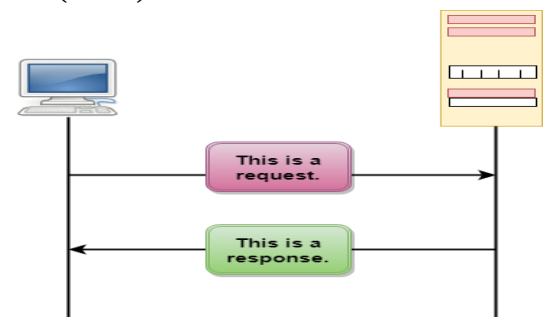
- **Speed:** One of the biggest advantages of FTP is speed. The FTP is one of the fastest way to transfer the files from one computer to another computer.
- **Efficient:** It is more efficient as we do not need to complete all the operations to get the entire file.
- **Security:** To access the FTP server, we need to login with the username and password. Therefore, we can say that FTP is more secure.
- **Back & forth movement:** FTP allows us to transfer the files back and forth. Suppose you are a manager of the company, you send some information to all the employees, and they all send information back on the same server.

Disadvantages of FTP:

- The standard requirement of the industry is that all the FTP transmissions should be encrypted. However, not all the FTP providers are equal and not all the providers offer encryption. So, we will have to look out for the FTP providers that provides encryption.
- FTP serves two operations, i.e., to send and receive large files on a network. However, the size limit of the file is 2GB that can be sent. It also doesn't allow you to run simultaneous transfers to multiple receivers.
- Passwords and file contents are sent in clear text that allows unwanted eavesdropping. So, it is quite possible that attackers can carry out the brute force attack by trying to guess the FTP password.
- It is not compatible with every system.

➤ HTTP

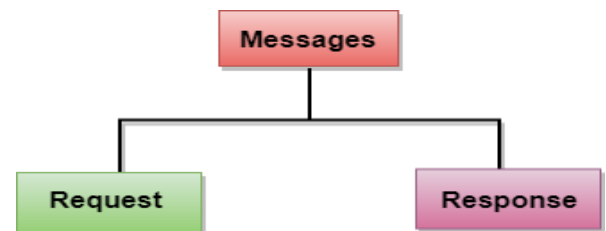
- HTTP stands for **HyperText Transfer Protocol**.
- It is a protocol used to access the data on the World Wide Web (www).
- The HTTP protocol can be used to transfer the data in the form of plain text, hypertext, audio, video, and so on.
- This protocol is known as HyperText Transfer Protocol because of its efficiency that allows us to use in a hypertext environment where there are rapid jumps from one document to another document.
- HTTP is similar to the FTP as it also transfers the files from one host to another host. But, HTTP is simpler than FTP as HTTP uses only one connection, i.e., no control connection to transfer the files.
- HTTP is used to carry the data in the form of MIME-like format.



- HTTP is similar to SMTP as the data is transferred between client and server. The HTTP differs from the SMTP in the way the messages are sent from the client to the server and from server to the client. SMTP messages are stored and forwarded while HTTP messages are delivered immediately.

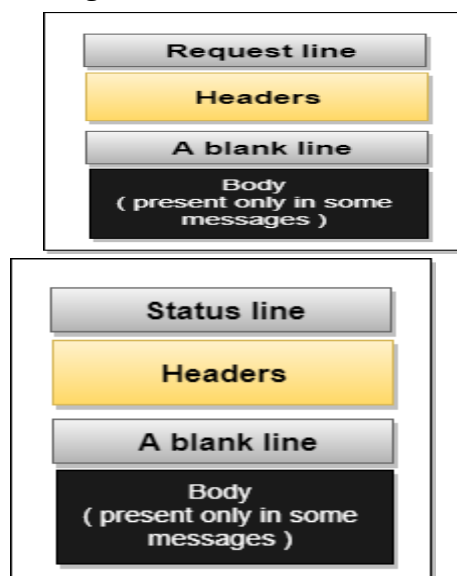
Features of HTTP:

- **Connectionless protocol:** HTTP is a connectionless protocol. HTTP client initiates a request and waits for a response from the server. When the server receives the request, the server processes the request and sends back the response to the HTTP client after which the client disconnects the connection. The connection between client and server exist only during the current request and response time only.
- **Media independent:** HTTP protocol is a media independent as data can be sent as long as both the client and server know how to handle the data content. It is required for both the client and server to specify the content type in MIME-type header.
- **Stateless:** HTTP is a stateless protocol as both the client and server know each other only during the current request. Due to this nature of the protocol, both the client and server do not retain the information between various requests of the web pages.



HTTP Transactions

The above figure shows the HTTP transaction between client and server. The client initiates a transaction by sending a request message to the server. The server replies to the request message by sending a response message.



Messages

HTTP messages are of two types: request and response. Both the message types follow the same message format.

Request Message: The request message is sent by the client that consists of a request line, headers, and sometimes a body.

Response Message: The response message is sent by the server to the client that consists of a status line, headers, and sometimes a body.

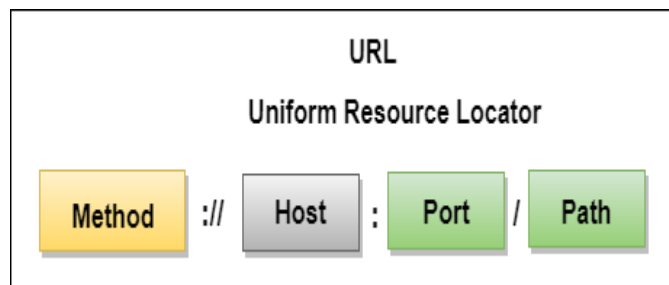
Uniform Resource Locator (URL)

- A client that wants to access the document in an internet needs an address and to facilitate the access of documents, the HTTP uses the concept of Uniform Resource Locator (URL).

- The Uniform Resource Locator (URL) is a standard way of specifying any kind of information on the internet.

- The URL defines four parts: method, host computer, port, and path. **Method:** The method is the protocol used to retrieve the document

from a server. For example, HTTP.

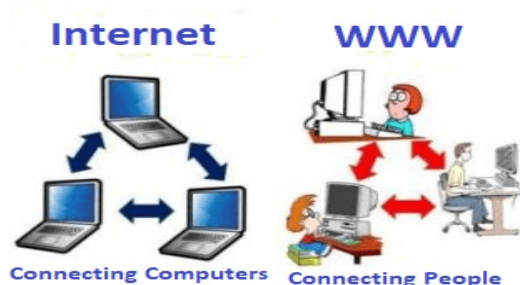


- **Host:** The host is the computer where the information is stored, and the computer is given an alias name. Web pages are mainly stored in the computers and the computers are given an alias name that begins with the characters "www". This field is not mandatory.

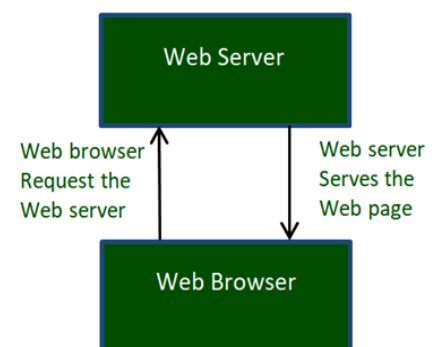
- **Port:** The URL can also contain the port number of the server, but it's an optional field. If the port number is included, then it must come between the host and path and it should be separated from the host by a colon.
- **Path:** Path is the pathname of the file where the information is stored. The path itself contains slashes that separate the directories from the subdirectories and files.
- What is World Wide Web?
- World Wide Web, which is also known as a Web, is a collection of websites or web pages stored in web servers and connected to local computers through the internet. These websites contain text pages, digital images, audios, videos, etc. Users can access the content of these sites from any part of the world over the internet using their devices such as computers, laptops, cell phones, etc. The WWW, along with internet, enables the retrieval and display of text and media to your device.
- The building blocks of the Web are web pages which are formatted in HTML and connected by links called "hypertext" or hyperlinks and accessed by HTTP. These links are electronic connections that link related pieces of information so that users can access the desired information quickly. Hypertext offers the advantage to select a word or phrase from text and thus to access other pages that provide additional information related to that word or phrase.
- A web page is given an online address called a Uniform Resource Locator (URL). A particular collection of web pages that belong to a specific URL is called a website,
- , e.g., www.facebook.com, www.google.com, etc. So, the World Wide Web is like a huge electronic book whose pages are stored on multiple servers across the world.
- Small websites store all of their WebPages on a single server, but big websites or organizations place their WebPages on different servers in different countries so that when users of a country search their site they could get the information quickly from the nearest server.
- So, the web provides a communication platform for users to retrieve and exchange information over the internet. Unlike a book, where we move from one page to another in a sequence, on World Wide Web we follow a web of hypertext links to visit a web page and from that web page to move to other web pages. You need a browser, which is installed on your computer, to access the Web.
- Difference between World Wide Web and Internet:
- Some people use the terms 'internet' and 'World Wide Web' interchangeably. They think they are the same thing, but it is not so. Internet is entirely different from WWW. It is a worldwide network of devices like computers, laptops, tablets, etc. It enables users to send emails to other users and chat with them online. For example, when you send an email or

chatting with someone online, you are using the internet.

- But, when you have opened a website like google.com for information, you are

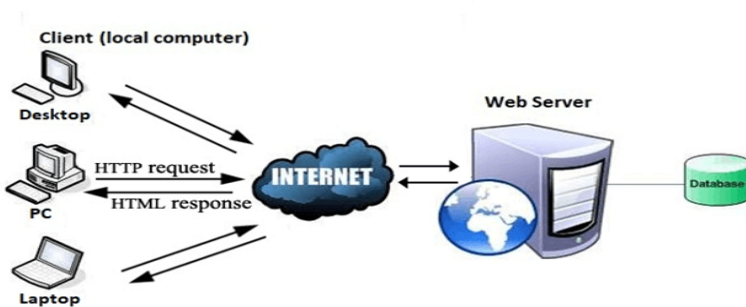


using the World Wide Web; a network of servers over the internet. You request a webpage from your computer using a browser, and the server renders that page to your browser. Your computer is called a client who runs a program (web browser), and asks the other computer (server) for the information it needs.



- **How the World Wide Web Works?**

- Now, we have understood that WWW is a collection of websites connected to the internet so that people can search and share information. Now, let us understand how it works!
- The Web works as per the internet's basic client-server format as shown in the following image. The servers store and transfer web pages or information to user's computers on the network when requested by the users. A web server is a software program which serves the web pages requested by web users using a browser. The computer of a user who requests documents from a server is known as a client. Browser, which is installed on the user's computer, allows users to view the retrieved documents.
- All the websites are stored in web servers. Just as someone lives on rent in a house, a website occupies a space in a server and remains stored in it. The server hosts the website whenever a user requests its WebPages, and the website owner has to pay the hosting price for the same.



clients (computers of users). These technologies include Hypertext Markup Language (HTML), Hypertext Transfer Protocol (HTTP) and Web browsers.

➤ **Electronic Mail**

One of the most popular Internet services is electronic mail(e-mail). Email is a service which allows us to send the message in electronic mode over the internet. It offers an efficient, inexpensive and real time mean of distributing information among people.

E-Mail Address

Each user of email is assigned a unique name for his email account. This name is known as E-mail address. Different users can send and receive messages according to the e-mail address.

E-mail is generally of the form username@domainname. For example, webmaster@tutorialspoint.com is an e-mail address where webmaster is username and tutorialspoint.com is domain name.

- The username and the domain name are separated by @ (**at**) symbol.
- E-mail addresses are not case sensitive.
- Spaces are not allowed in e-mail address.

E-mail Message Components

E-mail message comprises of different components: E-mail Header, Greeting, Text, and Signature. These components are described in the following diagram:

E-mail Header

The first five lines of an E-mail message is called E-mail header. The header part comprises of following fields:

From

The **From** field indicates the sender's address i.e. who sent the e-mail.

Date

The **Date** field indicates the date when the e-mail was sent.

To The To field indicates the recipient's address i.e. to whom the e-mail is sent.

Subject The Subject field indicates the purpose of e-mail. It should be precise and to the point.

CC CC stands for Carbon copy. It includes those recipient addresses whom we want to keep informed but not exactly the intended recipient.

BCC BCC stands for Black Carbon Copy. It is used when we do not want one or more of the recipients to know that someone else was copied on the message.

Greeting Greeting is the opening of the actual message. Eg. Hi Sir or Hi Guys etc.

Text It represents the actual content of the message.

Signature This is the final part of an e-mail message. It includes Name of Sender, Address, and Contact Number.

Advantages

E-mail has proved to be powerful and reliable medium of communication. Here are the benefits of **E-mail**:

Reliable

Many of the mail systems notify the sender if e-mail message was undeliverable.

Convenience

There is no requirement of stationary and stamps. One does not have to go to post office. But all these things are not required for sending or receiving an mail.

Speed

E-mail is very fast. However, the speed also depends upon the underlying network.

Inexpensive

The cost of sending e-mail is very low.

Printable

It is easy to obtain a hardcopy of an e-mail. Also an electronic copy of an e-mail can also be saved for records.

Global

E-mail can be sent and received by a person sitting across the globe.

Generality

It is also possible to send graphics, programs and sounds with an e-mail.

Disadvantages

Apart from several benefits of E-mail, there also exists some disadvantages as discussed below:

Forgery

E-mail doesn't prevent from forgery, that is, someone impersonating the sender, since sender is usually not authenticated in any way.

Overload

Convenience of E-mail may result in a flood of mail.

Misdirection

It is possible that you may send e-mail to an unintended recipient.

Junk

Junk emails are undesirable and inappropriate emails. Junk emails are sometimes referred to as spam.

No Response

It may be frustrating when the recipient does not read the e-mail and respond on a regular basis.

E-mail Protocols are set of rules that help the client to properly transmit the information to or from the mail server. Here in this tutorial, we will discuss various protocols such as **SMTP**, **POP**, and **IMAP**.

SMTP

SMTP stands for **Simple Mail Transfer Protocol**. It was first proposed in 1982. It is a standard protocol used for sending e-mail efficiently and reliably over the internet.

Key Points:

- SMTP is application level protocol.
- SMTP is connection oriented protocol.
- SMTP is text based protocol.
- It handles exchange of messages between e-mail servers over TCP/IP network.
- Apart from transferring e-mail, SMTP also provides notification regarding incoming mail.
- When you send e-mail, your e-mail client sends it to your e-mail server which further contacts the recipient mail server using SMTP client.
- These SMTP commands specify the sender's and receiver's e-mail address, along with the message to be send.
- The exchange of commands between servers is carried out without intervention of any user.
- In case, message cannot be delivered, an error report is sent to the sender which makes SMTP a reliable protocol The following table describes some of the SMTP commands:.

S.N.	Command Description
1	HELLO This command initiates the SMTP conversation.
2	EHELO This is an alternative command to initiate the conversation. ESMTP indicates that the sender server wants to use extended SMTP protocol.
3	MAIL FROM This indicates the sender's address.
4	RCPT TO It identifies the recipient of the mail. In order to deliver similar message to multiple users this command can be repeated multiple times.
5	SIZE This command let the server know the size of attached message in bytes.
6	DATA The DATA command signifies that a stream of data will follow. Here stream of data refers to the body of the message.
7	QUIT This commands is used to terminate the SMTP connection.
8	VERFY This command is used by the receiving server in order to verify whether the given username is valid or not.
9	EXPN It is same as VRFY, except it will list all the users name when it used with a distribution list.

➤ IMAP

IMAP stands for **Internet Message Access Protocol**. It was first proposed in 1986. There exist five versions of IMAP as follows:

1. Original IMAP
2. IMAP2
3. IMAP3
4. IMAP2bis
5. IMAP4

Key Points:

- IMAP allows the client program to manipulate the e-mail message on the server without downloading them on the local computer.
- The e-mail is hold and maintained by the remote server.
- It enables us to take any action such as downloading, delete the mail without reading the mail.It enables us to create, manipulate and delete remote message folders called mail boxes.
- IMAP enables the users to search the e-mails.
- It allows concurrent access to multiple mailboxes on multiple mail servers. The following table describes some of the IMAP commands:

S.N.	Command Description
1	IMAP_LOGIN This command opens the connection.
2	CAPABILITY This command requests for listing the capabilities that the server supports.
3	NOOP This command is used as a periodic poll for new messages or message status updates during a period of inactivity.
4	SELECT This command helps to select a mailbox to access the messages.
5	EXAMINE It is same as SELECT command except no change to the mailbox is permitted.
6	CREATE It is used to create mailbox with a specified name.
7	DELETE It is used to permanently delete a mailbox with a given name.
8	RENAME It is used to change the name of a mailbox.
9	LOGOUT This command informs the server that client is done with the session. The server must send BYE untagged response before the OK response and then close the network connection.

➤ POP

POP stands for Post Office Protocol. It is generally used to support a single client. There are several versions of POP but the POP 3 is the current standard.

Key Points

- POP is an application layer internet standard protocol.
- Since POP supports offline access to the messages, thus requires less internet usage time.
- POP does not allow search facility.
- In order to access the messages, it is necessary to download them.
- It allows only one mailbox to be created on server.
- It is not suitable for accessing non mail data.
- POP commands are generally abbreviated into codes of three or four letters. Eg. STAT. The following table describes some of the POP commands:

S.N.	Command Description
1	LOGIN This command opens the connection.
2	STAT It is used to display number of messages currently in the mailbox.
3	LIST It is used to get the summary of messages where each message summary is shown.
4	RETR This command helps to select a mailbox to access the messages.
5	DELE It is used to delete a message.
6	RSET It is used to reset the session to its initial state.
7	QUIT It is used to log off the session.

MIME stands for (**Multipurpose Internet Mail Extensions**). It is widely used internet standard for coding binary files to send them as e-mail attachments over the internet. MIME allows an E-mail message to contain a non-ASCII file such as a video image or a sound and it provides a mechanism to transfer a non text characters to text characters.

• MIME was invented to overcome the following limitations of SMTP:

1. SMTP cannot transfer executable files and binary objects.
2. SMTP cannot transmit text data of other language, e.g. French, Japanese, Chinese etc, as these are

represented in 8-bit codes.

3. SMTP services may reject mails having size greater than a certain size.
4. SMTP cannot handle non-textual data such as pictures, images, and video/audio content.

- **The MIME specification includes the following elements**

1. Message header fields. Five message header fields are defined. These fields provide information about the body of the message.
2. Content formats. A number of content formats are defined, thus standardizing representations that support multimedia electronic mail
3. **Transfer encoding.** Transfer encoding are defined that enable the conversion of any content format into a form that is protected from alteration by the mail system.

Traditional e-mail sent over the Internet using Simple Mail Transfer Protocol (SMTP) as specified by Request for Comments (The IETF standards documents are called RFC) 822 defines messages as consisting of a header and a body part, both of which are encoded using 7-bit ASCII text encoding. The header of an SMTP message consists of a series of field/value pairs that are structured so that the message can be delivered to its intended recipient. The body is unstructured text and contains the actual message.

- **Multipurpose Internet Mail Extensions (MIME)** five additional extensions to SMTP messages supports multipart messages with more two parts, and allows the encoding of 8-bit binary data such as image files so that they can be using SMTP. The encoding method for translation binary information used by MIME, Base64 Encoding, essentially provides a mechanism for translating non text information into text characters. The MIME extensions are implemented as fields in the e-mail message header.
- These fields are the following: Content type, Content transfer encoding method, MME version number Content ID (optional), Content description (optional).

MIME Header

The five header fields defined in MIME are as follows:

1. **MIME-version.** It indicates the MIME version being used. The current version is 1.1. It is represented as : MIME-version: 1.1.

2. **Content-type.** It describes the type and subtype of the data in the body of the message. The content type and content subtype are separated by slash.

This field describes how the object in the body is to be interpreted. The default value is plaintext in US ASCII. Content type field is represented as:

Context-type: <type/subtype; parameters>

There are seven different types and fourteen sub-types of content. The various content type are listed in the table below: Content-transfer encoding. It describes how the object within the body has been encoded to US ASCII to make it acceptable for mail transfer. Thus it specifies the method used to encode the message into 0s and 1s for transport.

The content transfer encoding field is represented as :

Type	Sub type	Description
Text	Plain	Unformatted text in US ASCII ISO 8859.
Image	jpeg	Image in JPEG Format.
	gif	Image in GIF format.
Video	mpeg	MPEG format.
Audio	Basic	Single- channel encoding of voice at 8kHz.
Message	rfc 822	The body is an encapsulated message that conforms to RFC 822.
	partial	Large mail is fragmented.
	External Body	contains pointer to an object that exists elsewhere and is accessible via FTP. TFTP etc.
Multipart	Mixed	The different parts are independent but are to be transmitted together. They should be presented to receiver in the appear in mail message.
	Parallel	same as mixed but order not defined.
	Alternate	The different parts are alternate versions of the same information
	Digest	similar to mixed, but the default type/subtype of each part is message/rfc 822.
Application	Postscript	Adobe postscript .
	Octet-stream	General binary data consisting of 8-bit bytes (Octets).

3. Content-transfer-encoding : <type>

The various encoding methods used are given in the table below:

Type	Description
7-bit	The body contains The 7-bit ASCII Characters With maximum length of 1000 characters
8-bit	There can be non-ASCII 8-bit characters but the maximum length of the body is limited to 1000 characters.
Binary	Binary 8-bit characters without limitation of 1000 characters in the body.
Quoted-printable	This is useful when data consists of largely printable characters. Characters in the range decimal equivalent 33 to 61 in ASCII are represented in ASCII. Others are represented as two-digit hex representation preceded by '=' sign, Non-text characters are replaced with six-digit hex sequence
Base 64	6-bit block of input data is encoded into 8-bit block of output.

This field is represented as:

Content-description: <description>

The various fields in the MIME header are

E-Mail Header	
MIME-VERSION :1.1 Content-type :type/subtype Content-transfer-encoding : encoding type Content-id : message id Content-description : textual explanation of non textual contents	MIME Header
E-Mail Body	
MIME Header	

➤ Video Conferencing

Video conference refers to a communication over a distance between three or more people where video and audio is transmitted in near real-time. Video conferencing is live, visual connection between two or more remote parties over the internet that simulates a face-to-face meeting. Video conferencing is important because it joins people who would not normally be able to form a face-to-face connection.

At its simplest, video conferencing provides transmission of static images and text between two locations. At its most sophisticated, it provides transmission of full-motion video images and high-quality audio between multiple locations.

In the business world, desktop video conferencing is a core component of unified communications platforms that also include calling and messaging capabilities. Standalone on-premises and cloud-based video conferencing platforms are also available from numerous vendors who support desktop- and room-based video, as well as the ability to embed video conferencing into business applications, such as telehealth, customer service and distance learning.

The widespread availability of cloud-based services enables organizations to implement video conferencing with minimal upfront investment and to take advantage of rapidly emerging AI-powered features to improve audio and video performance.

How video conferencing works

The video conferencing process can be split into two steps: compression and transfer.

During compression, the camera and microphone capture analog audiovisual (AV) input. The data collected is in the form of continuous waves of frequencies and amplitudes. These represent the captured sounds,

4. Content-Id. It is used to uniquely identify the MIME entities in multiple contexts i.e. it uniquely identifies the whole message in a multiple message environment. This field is represented as:

Content-id : id = <content-id>

5. Content-description. It is a plaintext description of the object within the body; It specifies whether the body is image, audio or video.

colors, brightness, depth and shades. Once captured, codecs convert data into digital packets, typically with compression to minimize bandwidth usage.

During the transfer phase, packets are sent over the network, typically to the cloud service provider, which then transmits them to other conference participants (and combines voice and video from multiple participants).

Once packets reach the endpoint, the codecs decompress the data. The codecs convert it back into analog audio and video. This enables the receiving screen and speakers to correctly view and hear the AV data.

Components of video conferencing systems

The components of a video conferencing system include the following:

- A network for data transfer, such as wired/wireless local area network, wide area network, cellular wireless and residential broadband.
- Two or more video cameras or webcams that provide video input.
- Two or more microphones -- either an external microphone or one built into the accessing device.
- A computer screen, monitor, TV or projector that can broadcast video output.
- Headphones, laptop speakers or external speakers that can be used for audio output.
- Codecs, which can be hardware- or software-based, to reduce bandwidth by compressing and decompressing AV data. They typically include acoustic echo cancellation capabilities, which reduce audio delays to support real-time communication. Codecs may also include features like noise cancellation and acoustic fencing to minimize background noise during conferences.

Benefits of video conferencing

Video conferencing services carry many benefits. In businesses, they can increase productivity among employees, as well as provide an improved way of communicating and interacting with colleagues, partners and customers.

For businesses, the tangible benefits of video conferencing include lower travel costs -- especially, for employee training -- and shortened meeting and project times as a result of improved communications among team members. Businesses can also increase revenue through higher quality virtual sales meetings.

The intangible benefits of video conferencing include more efficient meetings with the exchange of nonverbal communications and a stronger sense of community among business contacts, both within and between companies, as well as with customers.

On a personal level, the face-to-face connection enables participants to develop a stronger sense of familiarity with individuals they may never actually meet in person. Since the start of the COVID-19 pandemic, many companies now use video conferencing for community-building activities and social gatherings, including lunch and learns, health and wellness activities, happy hours and games.

Disadvantages of video conferencing

While video conferencing provides numerous benefits for businesses and individuals, it also has several disadvantages. For example, high-quality video calling and conferencing demands a consistently reliable high-speed internet connection with minimal latency and jitter. Only a strong internet connection can guarantee the voice audio and visuals will be reliably and smoothly communicated. Any issues with bandwidth or internet connectivity could cause the audio and/or video displays to be interrupted or lost, so quality of service measures may be required for important calls.

Another disadvantage is the cost of high-quality video conferencing systems. While many companies adopt video conferencing services to reduce business travel costs, they will still end up spending large amounts of money on a video conferencing system, especially for larger offices. In addition to all the

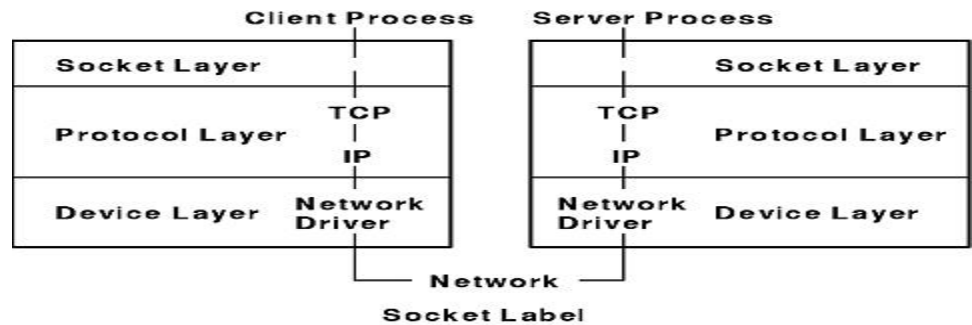
costly equipment and technology, companies will often also need to pay for the installation, deployment and maintenance of endpoints and on-premises servers if using an on-premises platform. made video conferencing ubiquitous on desktops and mobile devices that have an embedded camera. Facebook also expanded its reach from consumer video into enterprise video with Workplace Rooms. In the enterprise space, video conferencing platforms include those offering conferencing applications, personal or room-based conferencing endpoints, or both applications and endpoints.

➤ Socket Interfaces

The kernel structure consists of three layers: the socket layer, the protocol layer, and the device layer.

The *socket layer* supplies the interface between the subroutines and lower layers, the *protocol layer* contains the

protocol modules used for communication, and the *device layer* contains the device drivers that control the network devices. Protocols and drivers are dynamically loadable. The Socket Label figure illustrates the relationship between the layers.



Processes communicate using the client and server model. In this model, a server process, one end point of a two-way communication path, listens to a socket. The client process, the other end of the communication path, communicates to the server process over another socket. The client process can be on another machine. The kernel maintains internal connections and routes data from client to server.

Within the socket layer, the socket data structure is the focus of activity. The system-call interface subroutines manage the activities related to a subroutine, collecting the subroutine parameters and converting program data into the format expected by second-level subroutines.

Most of the socket facilities are implemented within second-level subroutines. These second-level subroutines directly manipulate socket data structures and manage the synchronization between asynchronous activities.

Socket Interface to Network Facilities

- This section explains the socket interprocess communication (IPC) facilities.
- This section explains the socket interprocess communication (IPC) facilities.
- The socket interprocess communication (IPC) facilities, illustrated by the Operating System Layer Examples figure (Figure 1), are layered on top of networking facilities. Data flows from an application program through the socket layer to the networking support. A protocol-related state is maintained in auxiliary data structures that are specific to the supporting protocols. The socket level passes responsibility for storage associated with transmitted data to the network level.
- Some of the communication domains supported by the socket IPC facility provide access to network protocols. These protocols are implemented as a separate software layer logically below the socket software in the kernel. The kernel provides ancillary services, such as buffer management, message routing, standardized interfaces to the protocols, and interfaces to the network interface drivers for the use of the various network protocols.
- User request and control output subroutines serve as the interface from the socket subroutines to the communication protocols.