

**KARPAGAM ACADEMY OF HIGHER EDUCATION**

Coimbatore-641 021

(For the candidates admitted from 2016 onwards)

DEPARTMENT OF COMPUTER SCIENCE, CA & IT**SUBJECT NAME: DATA COMMUNICATION NETWORKS****SEMESTER : V****SUBJECT CODE: 15CSU505C****CLASS: III B. Sc - CS****COURSE OBJECTIVE:**

This course is to master the fundamentals of data communications networks by gaining a working knowledge of data transmission concepts, understanding the operation of all seven layers of OSI Model and the protocols used in each layer.

COURSE OUTCOME:

1. Build an understanding of the fundamental concepts of computer networking.
2. Familiarize the student with the basic taxonomy and terminology of the computer networking.
3. Introduce the student to advanced networking concepts, preparing the student for entry Advanced courses in computer networking.
4. Allow the student to gain expertise in some specific areas of networking such as the design and maintenance of individual networks.

UNIT-I

Data Communication: An Overview – Protocols and Standards. Network Models: The OSI Model and Layers – TCP/IP Protocol Suite – Addressing. Physical Layer: Analog and Digital Signals – Transmission Impairments.

UNIT-II

Physical Layer: Multiplexing – Frequency Division Multiplexing-Wavelength Division Multiplexing– Synchronous Time-Division Multiplexing– Statistical Time Division Multiplexing. Transmission Media - Guided Media- Twisted pair and coaxial cable - Fiber optic cable-Unguided Transmission Media Switching – Circuit Switched Networks-Datagram Networks – Virtual Circuit networks.

UNIT-III

Data Link Layer: Error Correction and Detection – Framing – Flow and Error Control – Protocols – Noisy and Noiseless Channel – Multiple Access.

UNIT-IV

Network Layer: IPv4 addresses – Internetworking – IPv4 – Delivery and Forwarding – Unicast Routing Protocols.

Transport Layer: Process to Process Delivery – User Datagram Protocol – Transmission Control Protocol.

UNIT-V

Transport Layer: Data Traffic – Congestion Control.

Application Layer: Electronic Mail – File Traffic – WWW and HTTP – Symmetric Key and Asymmetric Key Cryptography – Security Services – Message Integrity – Message Authentication – Digital Signature.

TEXT BOOK

1. Behrouz A. Forouzan. 2006. Data Communication and Networking, 4th Edition, McGraw Hill, New Delhi.

REFERENCES

1. Andrews S. Tanenbaum. 2003. Computer Networks. 4th Edition, Prentice Hall of India, New Delhi.
2. Douglas E. Comer. 2000. Computer Networks and Internets, 2nd Edition. Pearson Education Asia, New Delhi.
3. Stanford H. Rowe and Marsha L. Schuh. 2005. Computer Networking, 1st Edition, Pearson Education,
4. William Stallings. 2007. Data and Communication Network, 8th Edition, Tata McGraw Hill, New Delhi.

WEB SITES

1. www.mhhe.com/engcs/compsci/forouzan/
2. www.amazon.com/Data-Communications-Networking-Behrouz-Forouzan/dp/0072923547
3. highered.mcgraw-hill.com/sites/0072515848/information_center_view0/ -

ESE MARKS ALLOCATION

1.	Section A 20 X 1 = 20 (Online Examination)	20
2.	Section B 5 X 8 = 40 (Either 'A' or 'B' Choice)	40
4.	Total	60

**KARPAGAM ACADEMY OF HIGHER EDUCATION**

Coimbatore-641 021

(For the candidates admitted from 2016 onwards)

DEPARTMENT OF COMPUTER SCIENCE, CA & IT**SUBJECT NAME: DATA COMMUNICATION NETWORKS****SUBJECT CODE: 15CSU502****SEMESTER: V****STAFF: Dr. T. GENISH****CLASS: III B. Sc (CS)****LECTURE PLAN**

Sl.No	Lecture Duration (Periods)	Topics to be covered	Support Materials
Unit- I			
1	1	Data Communication : Overview	T1:1-7,R1
2	1	Network Models and categories of networks	T1:8-16,W1
3	1	Protocols and Standards	T1:19-21
4	1	The OSI Model- Physical, data link	T1:27-29,W2
5	1	Network layer	W2
6	1	Transport, session, presentation layers	T1:36-42,W3
7	1	Presentation layer	W3
8	1	Application layer	W3
9	1	TCP/IP Protocol Suite	T1:42-45
10	1	Addressing	T1:45-50
11	1	Analog Signals	T1:59-80
12	1	Digital Signals	T1:71-80
13	1	Transmission Impairments	T1:80-84

14	1	Recapitulation and Possible Questions Discussion	
15	1	Recapitulation and Possible Questions Discussion	
		Total No. Of Hours Planned	15
TEXT BOOK:		T1: Behrouz A. Forouzan. 2006. Data Communication and Networking, 4 th Edition, McGraw Hill, New Delhi.	
REFERENCES		R1: Andrews S. Tanenbaum. 2003. Computer Networks. 4 th Edition, Prentice Hall of India, New Delhi.	
WEB SITES		W1: www.mhhe.com/engcs/compsci/forouzan/ W2: www.amazon.com/Data-Communications-Networking-Behrouz-Forouzan/dp/0072923547 W3: highered.mcgraw-hill.com/sites/0072515848/information_center_view0/	
Sl.No	Lecture Duration (Periods)	Topics to be covered	Support Materials
Unit- II			
1	1	Multiplexing – Frequency Division Multiplexing	T1:161-169, W2
2	1	Wavelength Division Multiplexing	
3	1	Multiplexing – Introduction	T1:169-179
4	1	Synchronous Time-Division Multiplexing	
5	1	Multiplexing – Statistical Time Division Multiplexing	T1:179-180
6	1	Transmission Media - Guided Media	T1:191-195
7	1	Twisted pair and coaxial cable	
8	1	Transmission Media- Fibre optic cable,	T1:196-208
9	1	Unguided Transmission Media	T1:203-208
10	1	Switching – Circuit Switched Networks	T1:214-218

11	1	Switching – Datagram Networks	T1:218-221,R2
12	1	Switching – Virtual Circuit networks	T1:221-227
13	1	Recapitulation and Possible Questions Discussion	
14	1	Recapitulation and Possible Questions Discussion	
		Total No. Of Hours Planned	14
TEXT BOOK:		T1: Behrouz A. Forouzan. 2006. Data Communication and Networking, 4 th Edition, McGraw Hill, New Delhi.	
REFERENCES		R2: Douglas E. Comer. 2000. Computer Networks and Internets, 2 nd Edition. Pearson Education Asia, New Delhi.	
WEB SITES		W2: www.amazon.com/Data-Communications-Networking-Behrouz-Forouzan/dp/0072923547	
Sl.No	Lecture Duration (Periods)	Topics to be covered	Support Materials
Unit- III			
1	1	Data link layer: Error Detection	T1:267-284
2	1	Error Correction	
3	1	Framing	T1:307-310
4	1	Flow and Error Control	T1:307-317
5	1	Noiseless Channel	
6	1	Noisy Channel Stop-and-wait	T1:318-324
7	1	Noisy Channel Go-Back-N	T1:324-332
8	1	Noisy Channel Selective Repeat	T1:332-340
9	1	Piggy backing	
10	1	Multiple Access - Random Access Protocols	T1:364-379
11	1	Multiple Access - Controlled Access Protocols	T1:379-383
12	1	Multiple Access - Channelization Protocols	T1:383-390

13	1	Recapitulation and Possible Questions Discussion	
14	1	Recapitulation and Possible Questions Discussion	
		Total No. Of Hours Planned	14
TEXT BOOK:		T1: Behrouz A. Forouzan. 2006. Data Communication and Networking, 4 th Edition, McGraw Hill, New Delhi.	
Sl.No	Lecture Duration (Periods)	Topics to be covered	Support Materials
Unit- IV			
1	1	Network Layer – IPv4 Addresses	T1:549-566
2	1	Internetworking: IPv4 – Datagram Fragmentation	T1:579-596
3	1	Checksum	
4	1	Delivery	T1:647-658
5	1	Forwarding	
6	1	Unicast Routing Protocols – Distance Vector Routing	T1:658-666
7	1	Unicast Routing Protocols – Link State Routing	T1:666-674
8	1	Unicast Routing Protocols – Path Vector Routing	T1:674-678
9	1	Transport Layer – Process to Process Delivery	T1:703-709
10	1	User Datagram Protocol	T1:709-715
11	1	TCP – Services	T1:715-728
12	1	Features	
13	1	TCP connection	
14	1	TCP –Flow control	T1:728-735
15	1	Error Control	

16	1	Congestion Control	T1:740-743
17	1	Recapitulation and Possible Questions Discussion	
		Total No. Of Hours Planned	17
TEXT BOOKS:		T1: Behrouz A. Forouzan. 2006. Data Communication and Networking, 4 th Edition, McGraw Hill, New Delhi.	
Sl.No	Lecture Duration (Periods)	Topics to be covered	Support Materials
Unit- V			
1	1	Data Traffic	T1:761-833
2	1	Congestion Control	
3	1	Application Layer - Email – Architecture, User Agent	
4	1	Email – File Transfer	T1:851-948
5	1	WWW and HTTP – Architecture Web documents	
6	1	Symmetric Key Cryptography	
7	1	Asymmetric Key Cryptography	T1:949-976
8	1	Security Services, Message Integrity Message Authentication	
9	1	Digital Signature	
10	1	Recapitulation and Possible Questions Discussion	
11	1	Previous year end-semester question paper discussion	
12	1	Previous year end-semester question paper discussion	
13	1	Previous year end-semester question paper discussion	

14	1	Previous year end-semester question paper discussion	
15	1	Previous year end-semester question paper discussion	
		Total No. Of Hours Planned	15
TEXT BOOKS:		T1: Behrouz A. Forouzan. 2006. Data Communication and Networking, 4 th Edition, McGraw Hill, New Delhi.	
Overall Total (All Units)			75

SUPPORT MATERIALS:**TEXT BOOK:**

T1: Behrouz A. Forouzan. 2006. Data Communication and Networking, 4th Edition, McGraw Hill, New Delhi.

REFERENCES

R1:Andrews S. Tanenbaum. 2003. Computer Networks. 4th Edition, Prentice Hall of India, New Delhi.

R2:Douglas E. Comer. 2000. Computer Networks and Internets, 2nd Edition. Pearson Education Asia, New Delhi.

WEBSITES:

W1: www.mhhe.com/engcs/compsci/forouzan/

W2:www.amazon.com/Data-Communications-Networking-Behrouz-Forouzan/dp/0072923547

W3:highered.mcgraw-hill.com/sites/0072515848/information_center_view0/

UNIT-I**UNIT-I- Syllabus**

Data Communication: An Overview – Protocols and Standards. Network Models: The OSI Model and Layers – TCP/IP Protocol Suite – Addressing. Physical Layer: Analog and Digital Signals – Transmission Impairments.

An Overview of data communication and networking

Data: Data refers to information presented in whatever form is agreed upon by the parties creating and using the data.

Data Communication

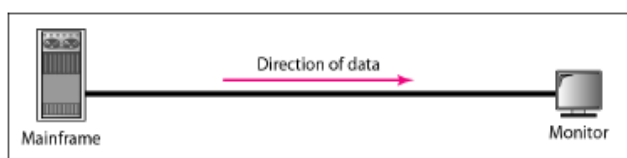
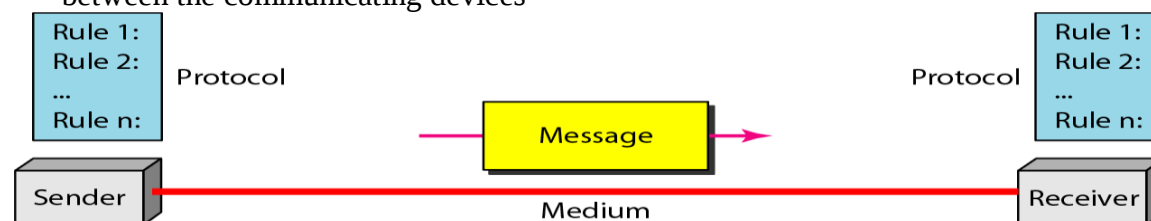
Networks exist so that data may be sent from one place to another. It is the exchange of data between two devices via some form of transmission medium such as a wire cable.

- For a communication to occur the communicating system must be made up of software and hardware
- Three fundamental characteristics for data communication system are
 1. Delivery- deliver data to correct destination
 2. Accuracy-must deliver the data accurately
 3. Timeliness- the system must deliver the data in a timely manner (eg: audio, video – real time transmission).
 4. Jitter- Jitter refers to the variation in the packet arrival time. It is the uneven delay in the delivery of audio or video packets.

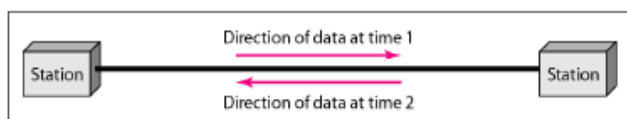
Components of data communication:

It has 5 components

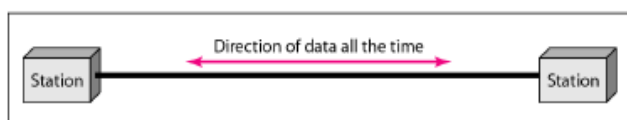
1. Message: It is the information to be communicated. It consists of text, numbers, pictures, sound, or any combination of these.
2. Sender: it is a device that sends the data message. It can be a computer, workstation, telephone handset, video camera etc.
3. Receiver: it is a device which receives the message. It can be a computer, workstation etc.
4. Medium: The transmission medium is the physical path by which a message travels from sender to receiver. It can be twisted pair wire, coaxial cable, fiber optic cable, or radio waves.
5. Protocols: It is a set of rules that governs data communication. It represents an agreement between the communicating devices



a. Simplex



b. Half-duplex



c. Full-duplex

Direction of data flow

Simplex: In this mode the communication is unidirectional eg.

Keyboards, monitors

- Half duplex: Each station can both transmit and receive but not at the same time, when one device is sending the other can only receive. (eg. Walkie talkies)
- Full duplex: both stations can transmit and receive simultaneously. (eg: telephone networks)

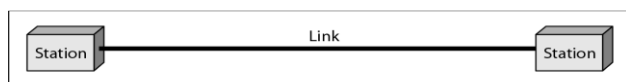
Networks:

-A network is a set of devices (often referred to as nodes) connected by communication links. A node can be a computer, printer, or any other device capable of sending and/or receiving data generated by other nodes on the network.

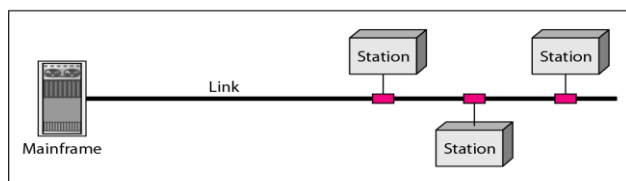
-Data communication between remote parties can be achieved through a process called networking, involving the connection of computers, media and networking devices.

Network criteria

- Performance—Can be measured in many ways
 - transit time: amount of time required for a message to travel from one device to another
 - response time: time elapsed between an inquiry and a response
 - Number of users
 - Type of transmission medium
 - Hardware capabilities and software efficiency
- Reliability—A measure of frequency of failure and the time needed to recover, network robustness
- Security—Protecting of data from unauthorized users



a. Point-to-point



b. Multipoint

Types of connections: point-to-point and multipoint

- a) Point-to-Point A point-to-point connection provides a dedicated link between two devices. The entire capacity of the link is reserved for transmission between those two devices. Most point-to-point connections use an actual length of wire or

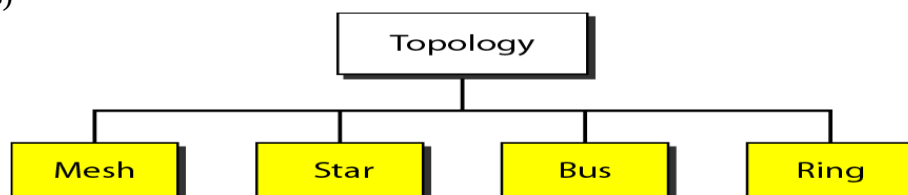
cable to connect the two ends, but other options, such as microwave or satellite links, are also possible.

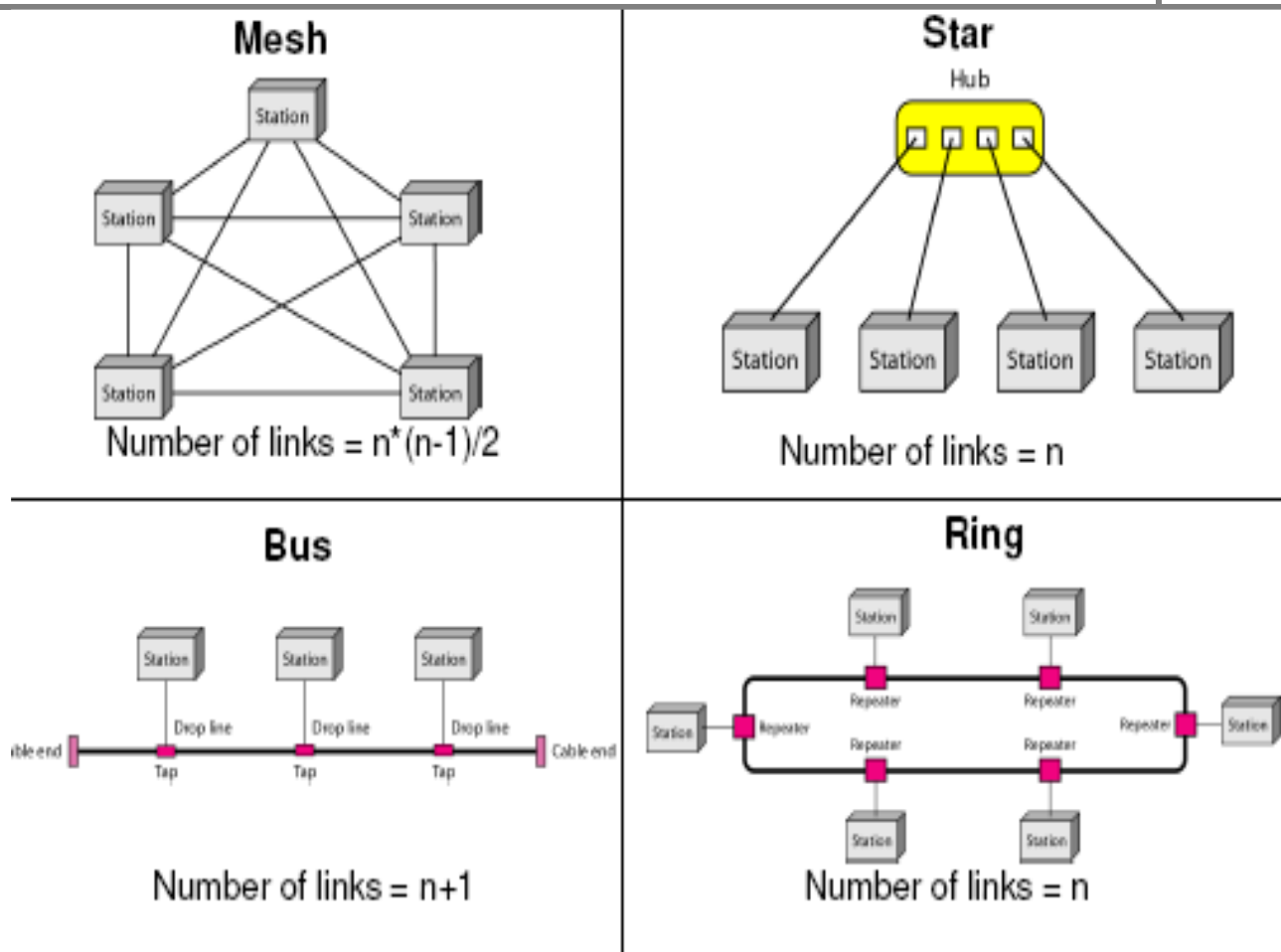
- b) Multipoint A multipoint (also called multi drop) connection is one in which more than two specific devices share a single link. In a multipoint environment, the capacity of the channel is shared, either spatially or temporally. If several devices can use the link simultaneously, it is a *spatially shared* connection. If users must take turns, it is a *timeshared* connection.

Physical Topology

Physical topology refers to the way in which a network is laid out physically.

Network topology is the geometric representation of the relationship of all the links and linking devices (nodes)





Mesh Topology: In a mesh topology, every device has a dedicated point-to-point link to every other device. The term dedicated means that the link carries traffic only between the two devices it connects. To find the number of physical links in a fully connected mesh network with n nodes, we first consider that each node must be connected to every other node. Node 1 must be connected to $n - 1$ nodes, node 2 must be connected to $n - 1$ nodes, and finally node n must be connected to $n - 1$ nodes. We need $n(n - 1)$ physical links. However, if each physical link allows communication in both directions (duplex mode), we can divide the number of links by 2. In other words, we can say that in a mesh topology, we need $n(n - 1) / 2$ duplex-mode links.

Advantages of Mesh Topology

- A dedicated links guarantees that each connection can carry its own data load, thus eliminating the traffic problems that can occur when links must be shared by multiple devices.
- A mesh topology is robust. If one link becomes unusable, it does not incapacitate the entire system.
- There is the advantage of privacy or security. When every message travels along a dedicated line, only the intended recipient sees it. Physical boundaries prevent other users from gaining access to messages.
- Finally, point-to-point links make fault identification and fault isolation easy. Traffic can be routed to avoid links with suspected problems. This facility enables the network manager to discover the precise location of the fault and aids in finding its cause and solution.

Disadvantages of Mesh Topology

The main disadvantages of a mesh are related to the amount of cabling and the number of I/O ports required.

- First, because every device must be connected to every other device, installation and reconnection are difficult.
- Second, the sheer bulk of the wiring can be greater than the available space (in walls, ceilings, or floors) can accommodate.
- Finally, the hardware required to connect each link (I/O ports and cable) can be prohibitively expensive.
- For these reasons a mesh topology is usually implemented in a limited fashion, for example, as a backbone connecting the main computers of a hybrid network that can include several other topologies.

One practical example of a mesh topology is the connection of telephone regional offices in which each regional office needs to be connected to every other regional office.

Star Topology

In a star topology, each device has a dedicated point-to-point link only to a central controller, usually called a hub. The devices are not directly linked to one another. Unlike a mesh topology, a star topology does not allow direct traffic between devices. The controller acts as an exchange: If one device wants to send data to another, it sends the data to the controller, which then relays the data to the other connected device.

Advantages of Star Topology

- A star topology is less expensive than a mesh topology. In a star, each device needs only one link and one I/O port to connect it to any number of others which also makes it easy to install and reconfigure.
- Other advantages include robustness. If one link fails, only that link is affected. All other links remain active. This factor also lends itself to easy fault identification and fault isolation.
- As long as the hub is working, it can be used to monitor link problems and bypass defective links.

Disadvantages of Star Topology

- One big disadvantage of a star topology is the dependency of the whole topology on one single point, the hub. If the hub goes down, the whole system is dead.
- Although a star requires far less cable than a mesh, each node must be linked to a central hub. For this reason, often more cabling is required in a star than in some other topologies (such as ring or bus).

The star topology is used in local-area networks (LANs), High-speed LANs often use a star topology with a central hub.

Bus Topology

The preceding examples all describe point-to-point connections. A **bus topology**, on the other hand, is multipoint. One long cable acts as a **backbone** to link all the devices in a network

Nodes are connected to the bus cable by drop lines and taps. A drop line is a connection running between the device and the main cable. A tap is a connector that either splices into the main cable or punctures the sheathing of a cable to create a contact with the metallic core. As a signal travels along the backbone, some of its energy is transformed into heat. Therefore, it becomes weaker and weaker as it travels farther and farther. For this reason there is a limit on the number of taps a bus can support and on the distance between those taps.

Advantages of Bus Topology

- Advantages of a bus topology include ease of installation. Backbone cable can be laid along the most efficient path, then connected to the nodes by drop lines of various lengths. In this way, a bus uses less cabling than mesh or star topologies.

- In a star, for example, four network devices in the same room require four lengths of cable reaching all the way to the hub. In a bus, this redundancy is eliminated. Only the backbone cable stretches through the entire facility. Each drop line has to reach only as far as the nearest point on the backbone.

Disadvantages

- Disadvantages include difficult reconnection and fault isolation.
- A bus is usually designed to be optimally efficient at installation. It can therefore be difficult to add new devices.
- Signal reflection at the taps can cause degradation in quality. This degradation can be controlled by limiting the number and spacing of devices connected to a given length of cable.
- Adding new devices may therefore require modification or replacement of the backbone.
- In addition, a fault or break in the bus cable stops all transmission, even between devices on the same side of the problem.

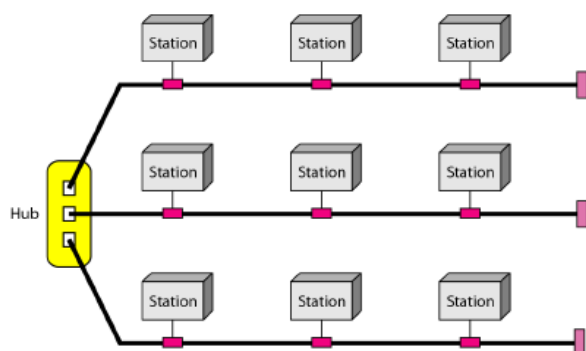
Bus topology was the one of the first topologies used in the design of early local area networks. Ethernet LANs can use a bus topology, but they are less popular now.

Ring Topology

Ring Topology In a ring topology, each device has a dedicated point-to-point connection with only the two devices on either side of it. A signal is passed along the ring in one direction, from device to device, until it reaches its destination. Each device in the ring incorporates a repeater. When a device receives a signal intended for another device, its repeater regenerates the bits and passes them along.

A ring is relatively easy to install and reconfigure. Each device is linked to only its immediate neighbors (either physically or logically). To add or delete a device requires changing only two connections. The only constraints are media and traffic considerations (maximum ring length and number of devices). In addition, fault isolation is simplified. Generally in a ring, a signal is circulating at all times. If one device does not receive a signal within a specified period, it can issue an alarm. The alarm alerts the network operator to the problem and its location. However, unidirectional traffic can be a disadvantage. In a simple ring, a break in the ring (such as a disabled station) can disable the entire network. This weakness can be solved by using a dual ring or a switch capable of closing off the break.

Ring topology was prevalent when IBM introduced its local-area network Token Ring. Today, the need for higher-speed LANs has made this topology less popular.



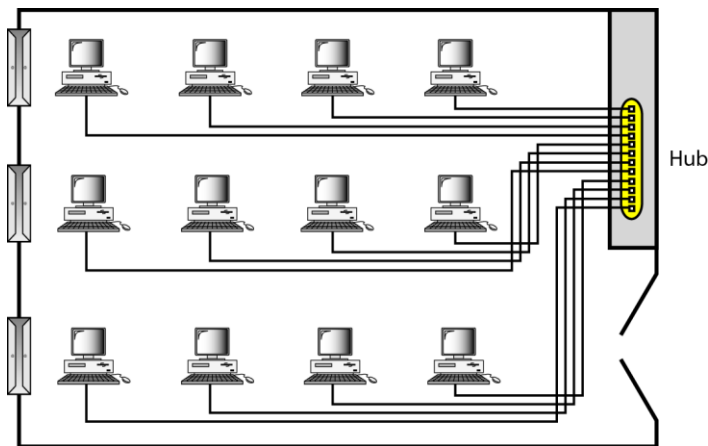
A hybrid topology

A network can be hybrid. For example, we can have a main star topology with each branch connecting several stations in a bus topology as shown in Figure below. A star backbone with three bus networks

Categories of networks

• Local Area Networks (LANs)

A local area network (LAN) is usually privately owned and links the devices in a single office, building, or campus. Depending on the needs of an organization and the type of technology used, a LAN can be as simple as two PCs and a printer in someone's home office; or it can extend throughout

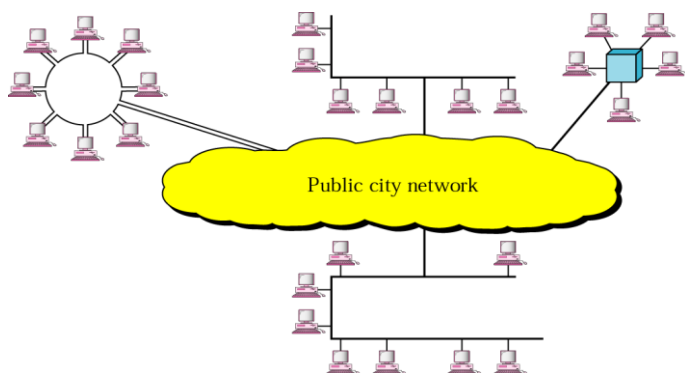


a company and include audio and video peripherals. Currently, LAN size is limited to a few kilometers.

LANs are designed to allow resources to be shared between personal computers or workstations. The resources to be shared can include hardware (e.g., a printer), software (e.g., an application program), or data. A common example of a LAN, found in many business environments, links a workgroup of task-related computers, for example, engineering workstations or accounting PCs.

One of the computers may be given a large capacity disk drive and may become a server to clients. Software can be stored on this central server and used as needed by the whole group. In this example, the size of the LAN may be determined by licensing restrictions on the number of users per copy of software, or by restrictions on the number of users licensed to access the operating system.

In addition to size, LANs are distinguished from other types of networks by their transmission media and topology. In general, a given LAN will use only one type of transmission medium. The most common LAN topologies are bus, ring, and star. Early LANs had data rates in the 4 to 16 megabits per second (Mbps) range. Today, however, speeds are normally 100 or 1000 Mbps

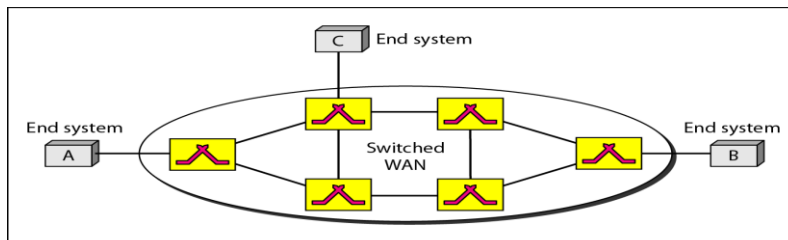


Metropolitan Area Networks (MANs)

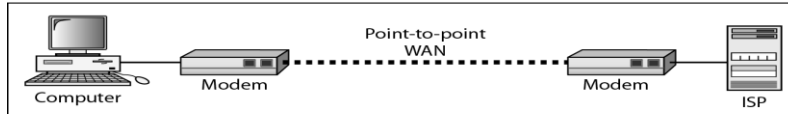
A metropolitan area network (MAN) is a network with a size between a LAN and a WAN. It normally covers the area inside a town or a city. It is designed for customers who need a high-speed connectivity, normally to the Internet, and have endpoints spread over a city or part of city. A good example of a MAN is the part of the telephone company network that can provide a high-speed DSL

line to the customer. Another example is the cable TV network that originally was designed for cable TV, but today can also be used for high-speed data connection to the Internet

• Wide Area Networks (WANs)



a. Switched WAN



b. Point-to-point WAN

A wide area network (WAN) provides long-distance transmission of data, image, audio, and video information over large geographic areas that may comprise a country, a continent, or even the whole world. A WAN can be as complex as the backbones that connect the Internet or as simple as a dial-up line that connects a home computer to the Internet. We normally refer to the first as a switched WAN and to the second as a point-to-point WAN.

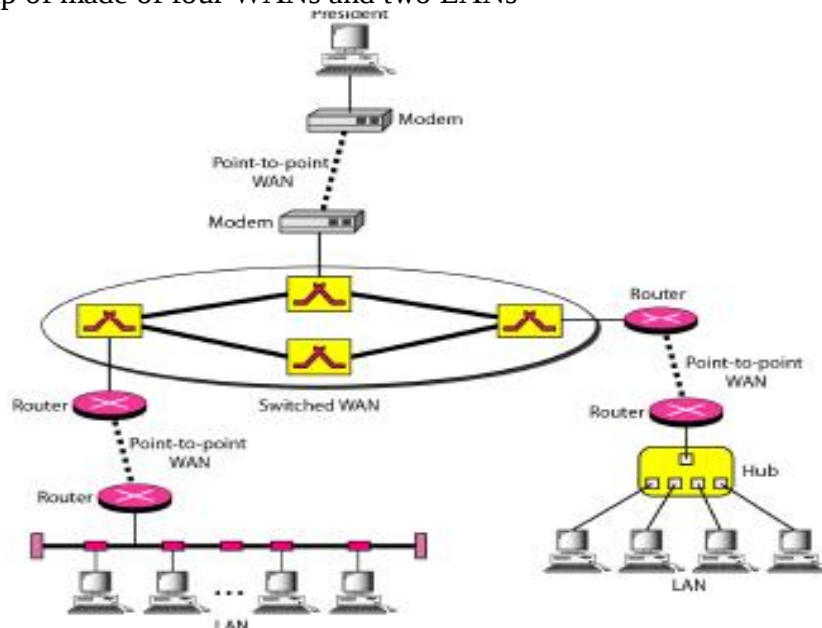
- The switched WAN connects the end systems, which usually comprise a router (internetworking connecting device) that connects to another LAN or WAN.
- The point-to-point WAN is normally a line leased from a telephone or cable TV provider that connects a home computer or a small LAN to an Internet service provider (ISP). This type of WAN is often used to provide Internet access.

• Internetworks

- Internetwork (internet) : two or more networks are connected by internetworking devices
- Internetworking devices: router, gateway, etc.
- The Internet: a specific worldwide network

An heterogeneous network

It is made up of made of four WANs and two LANs



Comparison of LANs, MANs, & WANs

Interprocessor distance	Processors located in same	Example
1 m	Square meter	Personal area network
10 m	Room	Local area network
100 m	Building	
1 km	Campus	
10 km	City	Metropolitan area network
100 km	Country	Wide area network
1000 km	Continent	
10,000 km	Planet	
		The Internet

LANs: 1 – 1000 Mbps

MANs: 10 – 40 Gbps

WANs: Tbps

THE INTERNET

The Internet has revolutionized many aspects off our daily lives. it has affected the way we do business as well as the way we spend our leisure time..

The Internet is a communication system that has brought a wealth off information to our finger tips and organized it for our use.

A brief history

- Mid-1960s
 - Standalone devices
 - ARPA (Advanced Research Projects Agency) was interested in finding a way to connect computers to share information
 - Backbones: None - Hosts: None
- 1967
 - ARPA presented its ideas for ARPANET
 - Backbones: None - Hosts: None
- 1969
 - The first physical network was constructed
 - Backbones: 50Kbps ARPANET - Hosts: 4
- 1972
 - The first e-mail program was created by Ray Tomlinson of BBN
- 1973
 - Development began on the protocol later to be called TCP/IP (by Vint Cerf and Bob Kahn)
 - Backbones: 50Kbps ARPANET - Hosts: >23

PROTOCOLS AND STANDARDS

Protocols and standards. First, we define protocol, which is synonymous with rule. Then we discuss standards, which are agreed-upon rules.

Protocols

- A protocol is a set of rules that governs data communications
- It defines what is communicated, how it is communicated and when it is communicated
- Key elements of a protocol:
 - **Syntax:** Structure or format of data, meaning the order in which they are presented
 - **Semantics:** Refer to the meaning of each section of bits, how a particular pattern is interpreted and what action to be taken
 - **Timing:** Refers to when data should be sent and how fast can they be sent

Standards

- Standards are essential in creating and maintaining an open and competitive market for equipment manufacturers
- Required to guarantee national and international interoperability of data and telecommunications technology and processes
- Categories of data communications standards
 1. De facto: Standards that have not been approved by an organizational body but have been adopted through widespread use, eg. model TCP/IP)
 2. De jure: Those that have been legislated by an official recognized body, eg. OSI model

Standards organizations

- Standards creation committees
 - ISO (International Organization for Standardization)
 - ITU-T (International Telecommunications Union – Telecommunications Standards)
- Initially known as CCITT (Consultative Committee for International Telegraphy and Telephony)
 - ANSI (American National Standards Institute)
 - IEEE (Institute of Electrical and Electronics Engineers)
 - EIA (Electronic Industries Association)
- Forums:
 - Made up of representatives from interested corporations to speed acceptance and use of new technologies in the telecom industry
- Regulatory Agencies
 - Governmental agencies: to protect public interest by regulating radio, TV and wire/cable communications

Internet standards

- An Internet standard is a thoroughly tested specification used by those who work with the Internet
- A specification begins with an Internet draft
 - Working document with no official status and a 6- month lifetime
 - Upon recommendation from the Internet authorities a draft may be published as a Request for Comment(RFC)

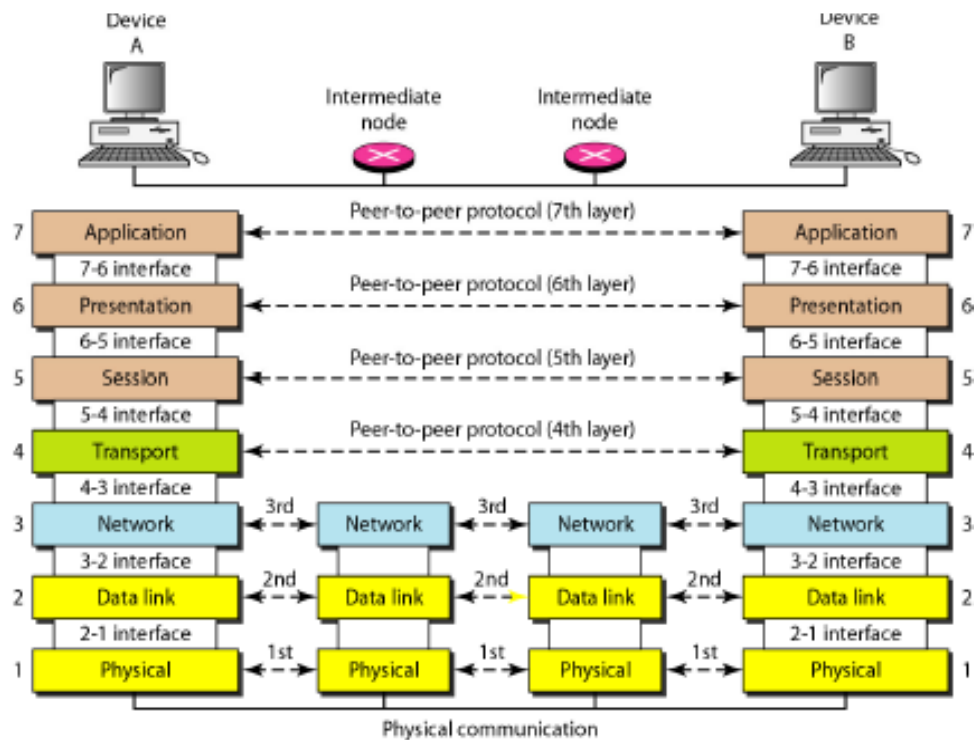
NETWORK MODELS

- A network uses a combination of hardware and software to send data from one location to another
 - Hardware consists of the physical equipment that carries signals from one point of the network to another
 - The task of sending a piece of information from one point in the works to another can be broken into several tasks, each performed by a separate software package
 - Each piece of software uses the services of another software package to do its job
 - At the lowest layer, a signal is sent from the source to the destination computer

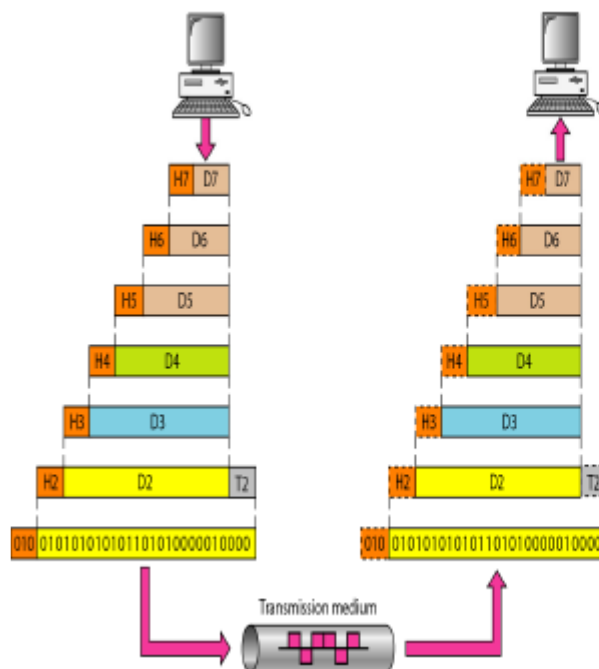
THE OSI MODEL

Established in 1947, the International Standards Organization (ISO) is a multinational body dedicated to world wide agreement on international standards. An ISO standard that covers all aspects of network communications is the Open Systems Interconnection (OSI) model. It was first introduced in the late 1970s.

Peer-to-peer processes



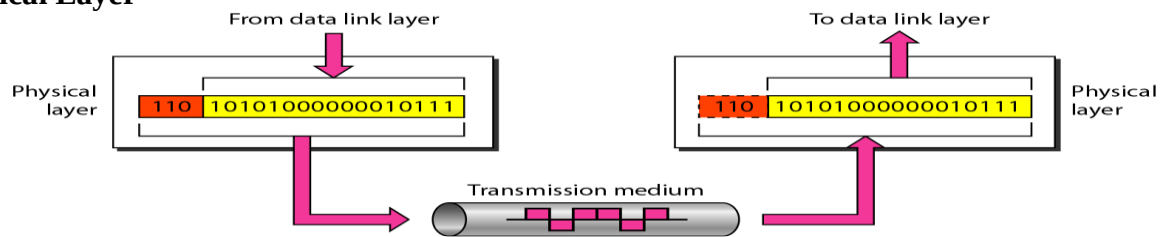
Encapsulation



- The process starts at layer 7 (application layer), then moves from layer to layer in descending, sequential order.
- At each layer, a **header** is added to the data unit.
- At layer 2, a **trailer** is added as well.
- When the formatted data unit passes through the physical layer (layer 1) it is changed into an electromagnetic or optical signal and transported along a physical link
- At the destination the reverse process is performed

THE OSI MODEL AND LAYERS

Physical Layer



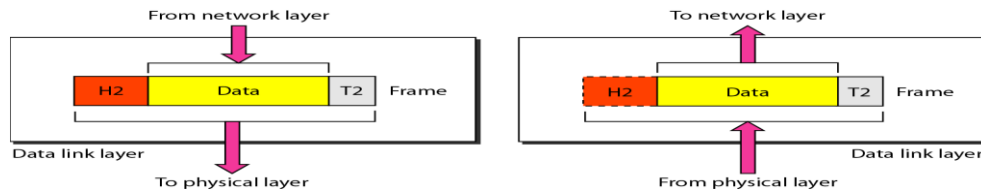
- The physical layer is responsible for movements of individual bits from one hop (node) to the next
- Mechanical and electrical specification, the procedures and functions

Duties:

- Physical characteristics of interfaces and media
- Representation of bits
- Data rate
- Synchronization of bits
- Line configuration
- Physical topology
- Transmission mode

Data link layer

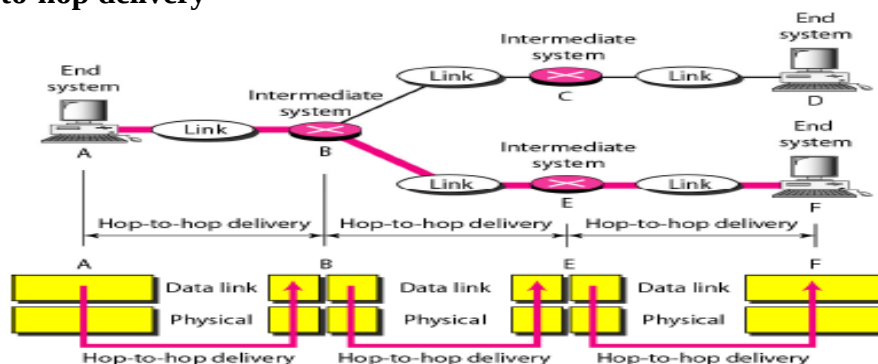
- The data link layer is responsible for moving frames from one hop (node) to the next
- Transform the physical layer to a reliable (error-free) link



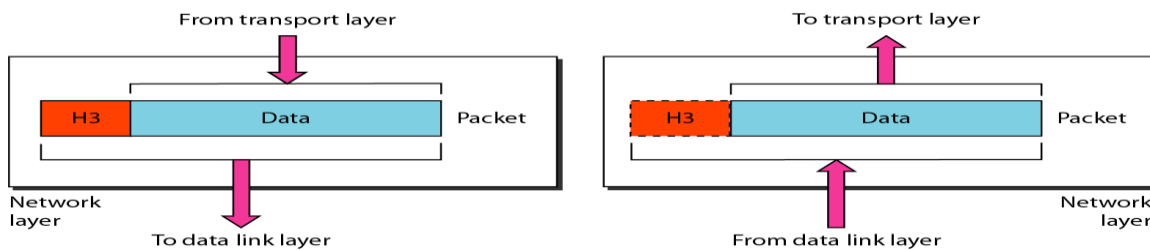
Duties:

- Framing
- Physical addressing
- Flow control
- Error control
- Access control

Hop-to-hop delivery



Network layer

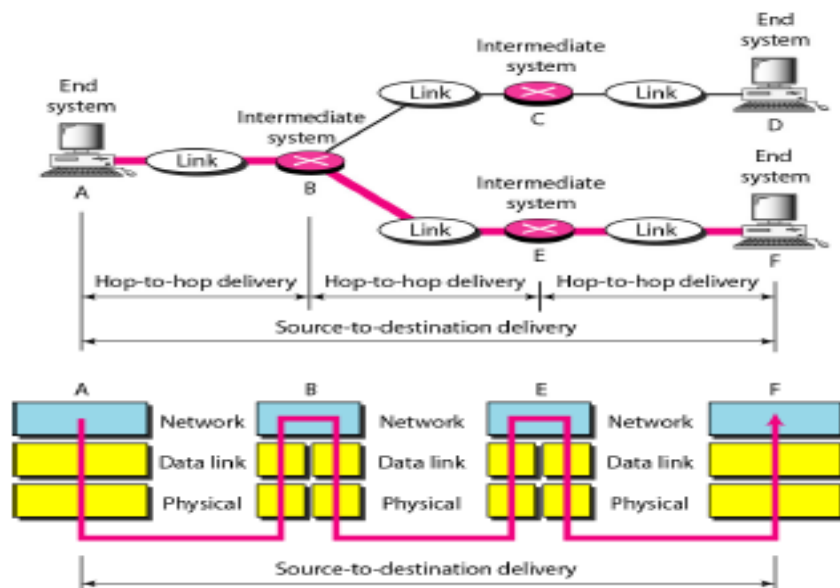


The network layer is responsible for the delivery of individual packets from the source host to the destination host.

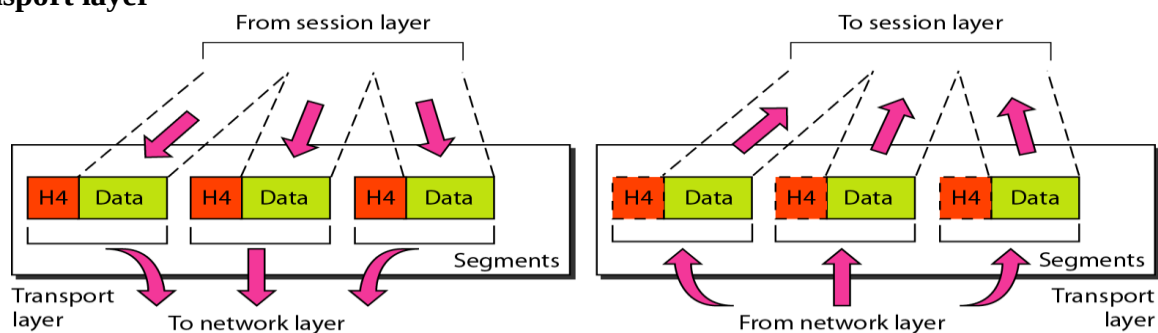
Duties:

- Logical addressing
- Routing

Source-to-destination delivery



Transport layer

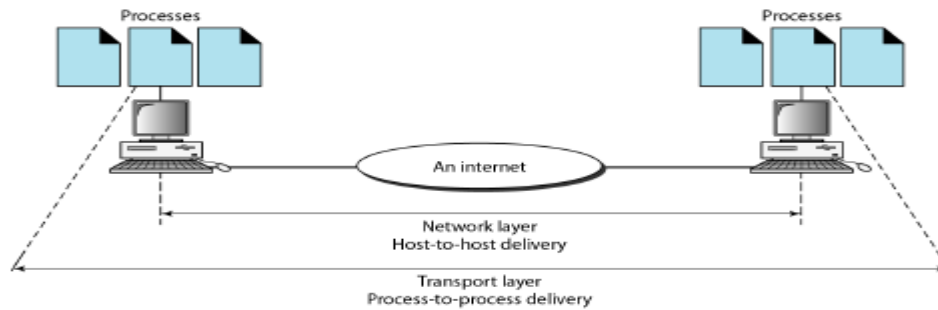


The transport layer is responsible for the delivery of a message from one process to another.

Duties:

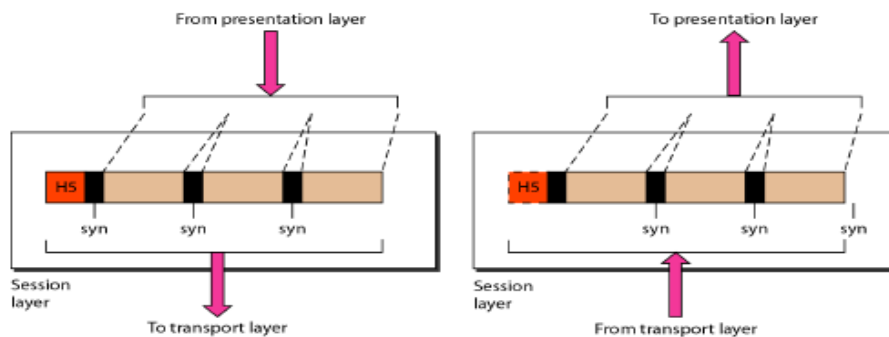
- Service-point (port) addressing
- Segmentation and reassembly
- Connection control
- Flow control
- Error control

Reliable process-to-process delivery



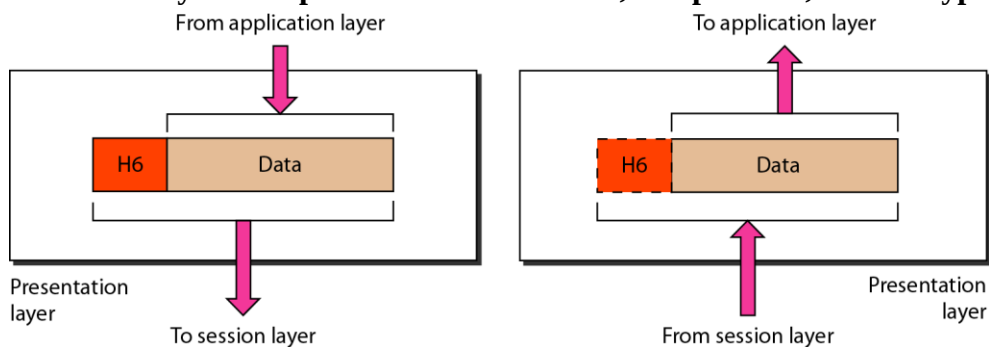
Session layer

The session layer is responsible for dialog control and synchronization.



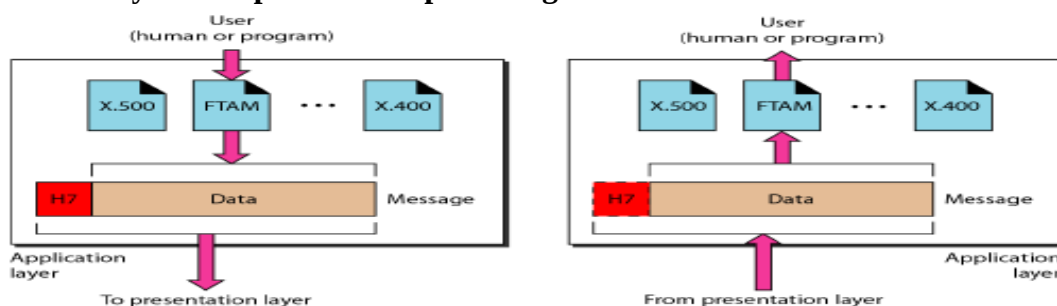
Presentation layer

The presentation layer is responsible for translation, compression, and encryption.



Application layer

The application layer is responsible for providing services to the user.

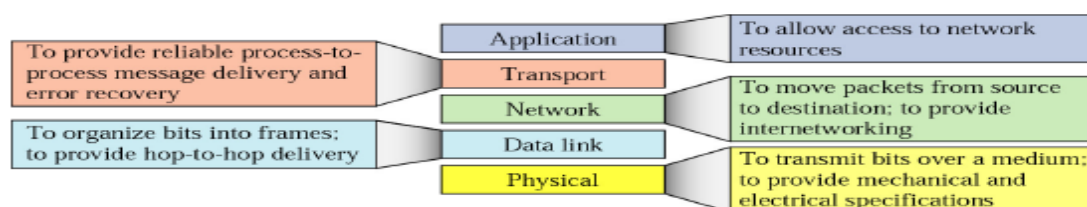
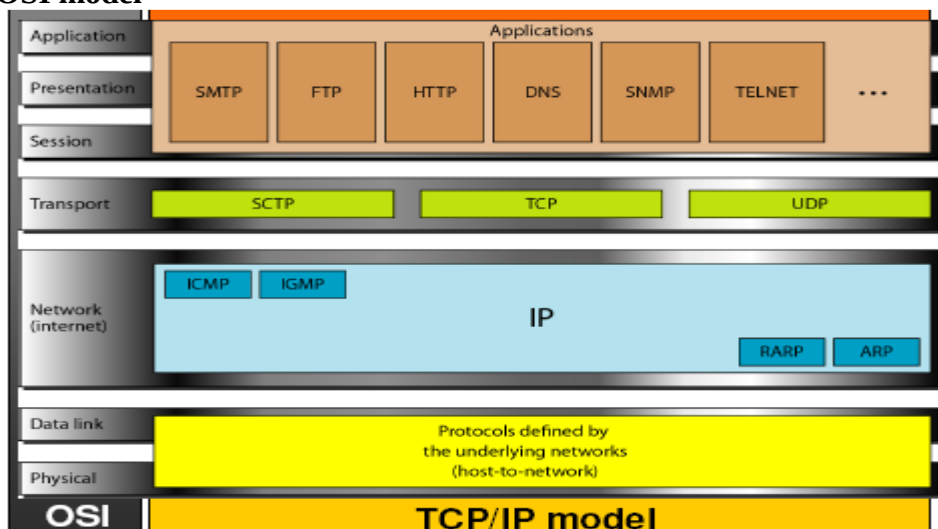


Services:

- Network virtual terminal
- Mail services
- File transfer, access, and management
- Directory services

TCP/IP PROTOCOL SUITE

The layers in the TCP/IP protocol suite do not exactly match those in the OSI model. The original TCP/IP protocol suite was defined as having four layers: host-to-network, internet, transport, and application. However, when TCP/IP is compared to OSI, we can say that the TCP/IP protocol suite is made of five layers: physical, data link, network, transport, and application.

TCP/IP layers**TCP/IP and OSI model****ADDRESSING**

Four levels of addresses are used in an internet employing the TCP/IP protocols: physical, logical, port, and specific.

• Physical address

In computer networks a physical address means a MAC (Medium Access Control) address. Also known as Ethernet Hardware Address (EHA) or hardware address or **adapter address**. It is a number that acts like a name for a particular network adapter, eg. the network cards

• Logical address

—In computer networks, a logical address refers to a network layer address such as an IP address

—An IP address (Internet Protocol address) is a unique address that certain electronic devices use in order to identify and communicate with each other on a computer network utilizing the Internet Protocol standard (IP)

• **Port address**

—TCP and UDP are transport protocols used for communication between computers via ports

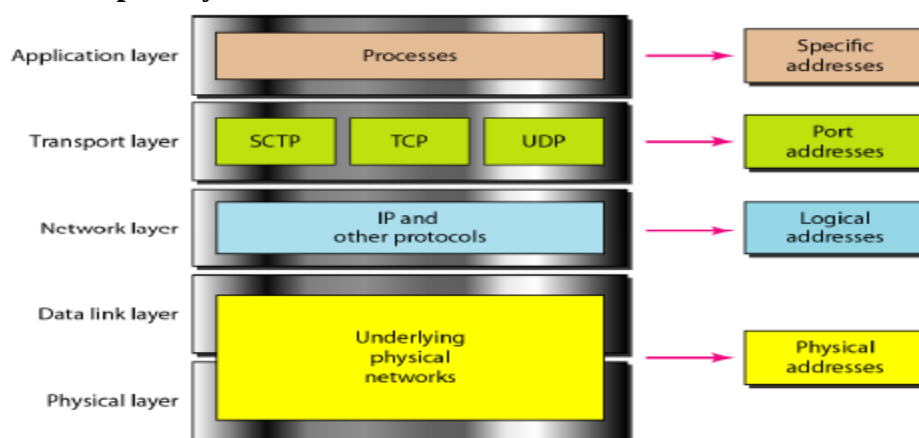
—The port numbers are divided into three ranges.

- The Well Known Ports are those in the range 0–1023.
- The Registered Ports are those in the range 1024–49151.
- The Dynamic and/or Private Ports are those in the range 49152–65535. These ports are not used by any defined application.

• **Specific address**

—This address is used by application processes

Relationship of layers-addresses in TCP/IP



PHYSICAL LAYER- ANALOG AND DIGITAL

Data can be analog or digital. The term analog data refers to information that is continuous; digital data refers to information that has discrete states. Analog data take on continuous values. Digital data take on discrete values.

Note -1

Data can be analog or digital.

Analog data are continuous and take continuous values.

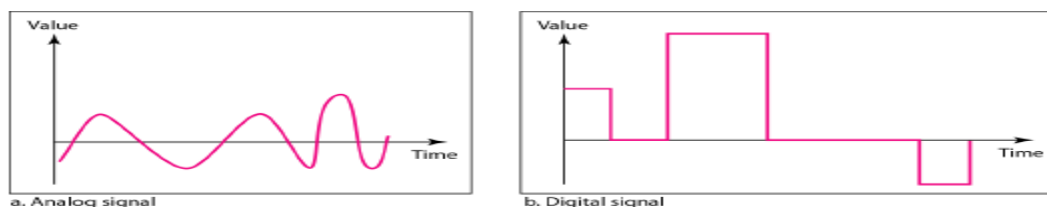
Digital data have discrete states and take discrete values.

Note-2

Signals can be analog or digital.

Analog signals can have an infinite number of values in a range;

Digital signals can have only a limited number of values.



Periodic and Aperiodic signals or NonPeriodic

Both analog and digital signals can take one of two forms

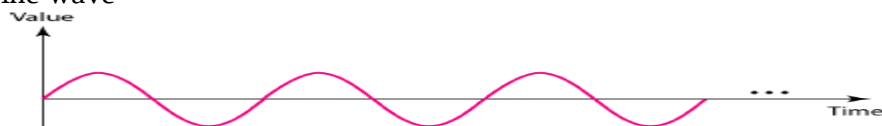
Periodic: completes a pattern within a measurable time frame called a period and repeats that pattern over subsequent identical periods

Nonperiodic: signal changes without exhibiting a pattern or cycle that repeats over time

PERIODIC ANALOG SIGNALS

Periodic analog signals can be classified as simple or composite. A simple periodic analog signal, a sine wave, cannot be decomposed into simpler signals. A composite periodic analog signal is composed of multiple sine waves

A sine wave



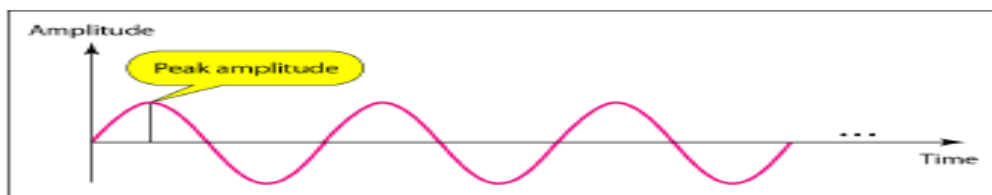
We can mathematically describe the sine wave as

$$s(t) = A \sin(2\pi ft + \phi)$$

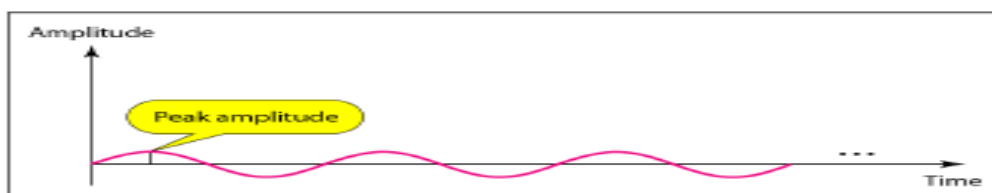
where, s	is the instantaneous amplitude
A	is the peak amplitude
f	is the frequency
ϕ	is the phase
t	is the time
π	is a constant (~ 3.14159)

Two signals

Same phase and frequency, but different amplitudes



a. A signal with high peak amplitude



b. A signal with low peak amplitude

Period and frequency

Period refers to the amount of time, in seconds, a signal needs to complete 1 cycle.

- Denoted by T , measured in seconds.

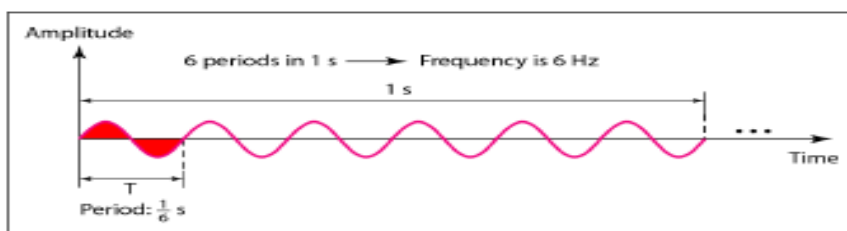
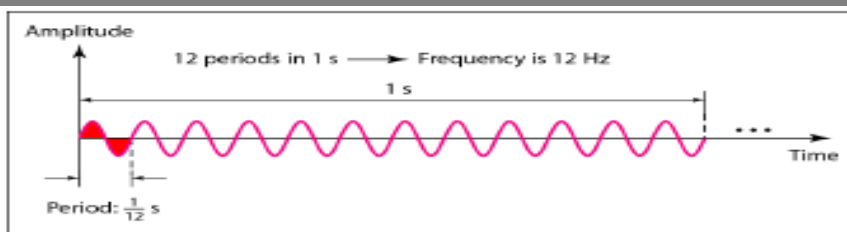
Frequency refers to the number of periods in one second

- Denoted by f , measured in Hertz (Hz)

Frequency and period are the inverse of each other.

$$f = \frac{1}{T} \quad \text{and} \quad T = \frac{1}{f}$$

Two signals-Same amplitude and phase, but different frequencies



Units of period and frequency

Unit	Equivalent	Unit	Equivalent
Seconds (s)	1 s	Hertz (Hz)	1 Hz
Milliseconds (ms)	10^{-3} s	Kilohertz (kHz)	10^3 Hz
Microseconds (μ s)	10^{-6} s	Megahertz (MHz)	10^6 Hz
Nanoseconds (ns)	10^{-9} s	Gigahertz (GHz)	10^9 Hz
Picoseconds (ps)	10^{-12} s	Terahertz (THz)	10^{12} Hz

More about frequency

- **Frequency is the rate of change with respect to time.**
- **Change in a short span of time means high frequency.**
- **Change over a long span of time means low frequency.**

Two extremes

- **If a signal does not change at all, its frequency is zero.**
- **If a signal changes instantaneously, its frequency is infinite.**

Phase

Phase describes the position of the waveform relative to time 0.

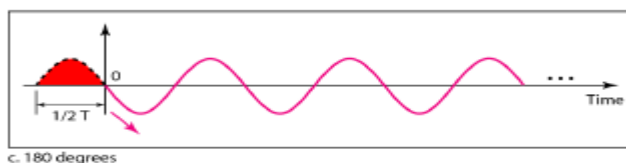
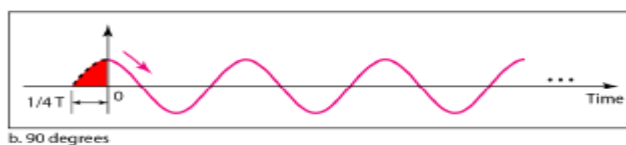
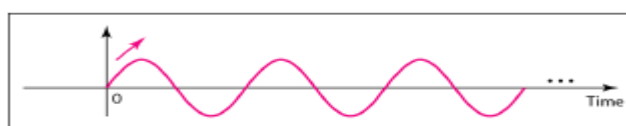
Three sine waves

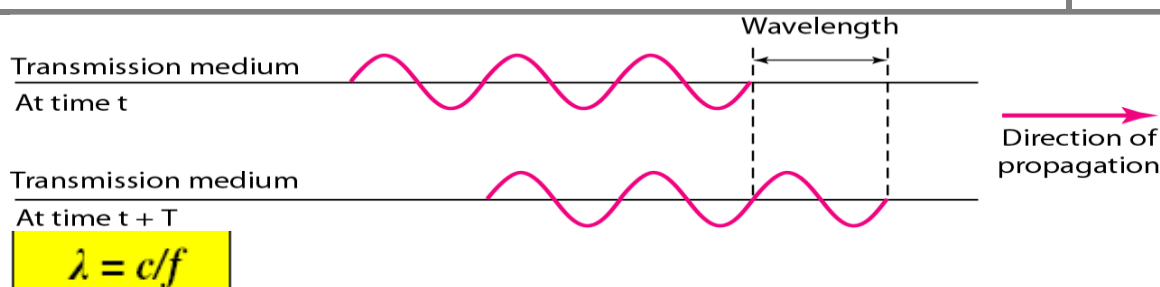
Same amplitude and frequency, but different phases

Wavelength and period

Wavelength is another characteristic of a signal traveling through a transmission medium.

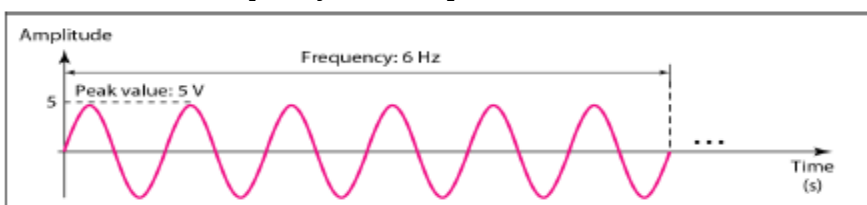
- The wavelength depends on both the frequency and the medium.
- The wavelength is the distance a signal can travel in one period.



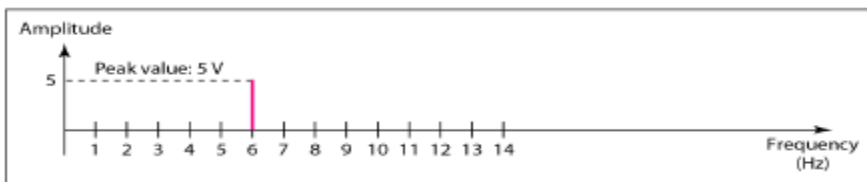


where, λ is the wavelength
 c is the speed of light ($\sim 3 \times 10^8$ m/s)
 f is the frequency

Time-domain and frequency-domain plots of a sine wave



a. A sine wave in the time domain (peak value: 5 V, frequency: 6 Hz)

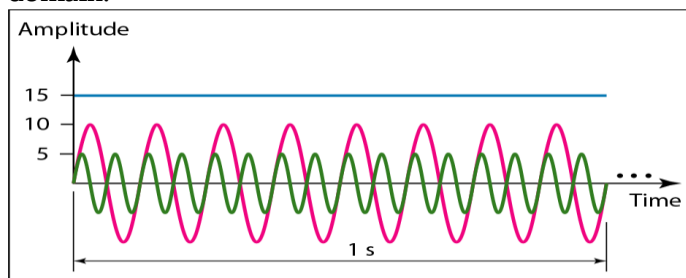


b. The same sine wave in the frequency domain (peak value: 5 V, frequency: 6 Hz)

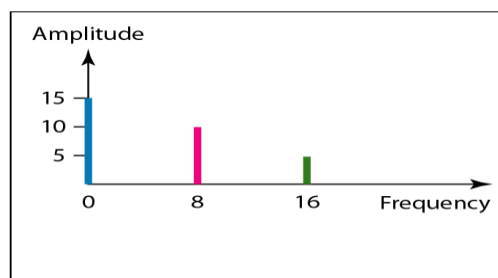
A complete sine wave in the time domain can be represented by one single spike in the frequency domain.

Example

The frequency domain is more compact and useful when we are dealing with more than one sine wave. For example, the following figure shows three sine waves, each with different amplitude and frequency. All can be represented by three spikes in the frequency domain.



a. Time-domain representation of three sine waves with frequencies 0, 8, and 16



b. Frequency-domain representation of the same three signals

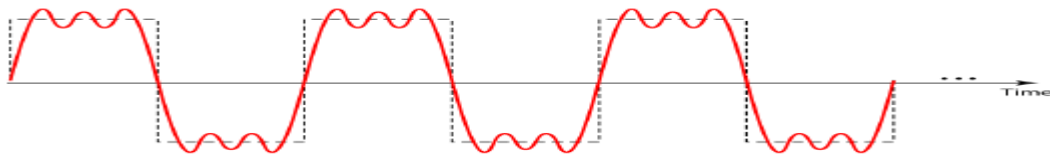
Composite signals

A single-frequency sine wave is not useful in data communications; we need to send a composite signal, a signal made of many simple sine waves.

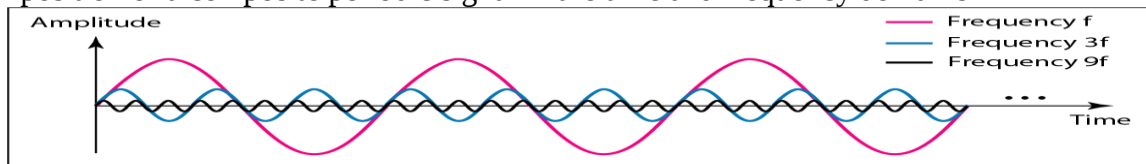
We can use a mathematical technique called **Fourier analysis** to show that any **periodic** signal is made up of an infinite series of sinusoidal frequency components.

If the composite signal is periodic, the decomposition gives a series of signals with discrete frequencies; if the composite signal is nonperiodic, the decomposition gives a combination of sine waves with continuous frequencies.

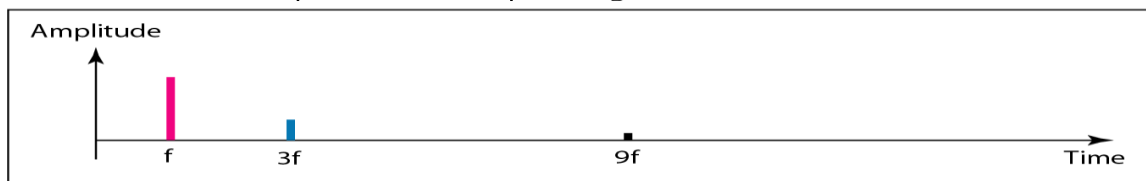
The figure shows a periodic composite signal with frequency f . This type of signal is not typical of those found in data communications. We can consider it to be three alarm systems, each with a different frequency. The analysis of this signal can give us a good understanding of how to decompose signals.



Decomposition of a composite periodic signal in the time and frequency domains

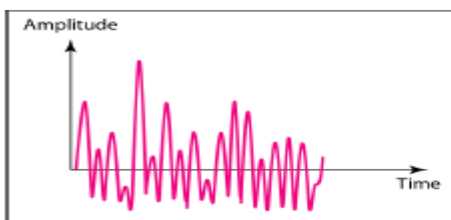


a. Time-domain decomposition of a composite signal

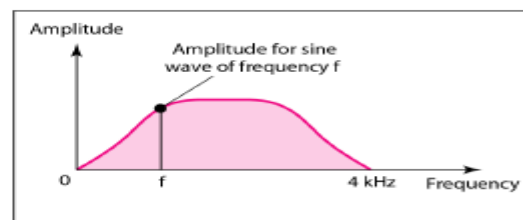


b. Frequency-domain decomposition of the composite signal

Example-The figure shows a nonperiodic composite signal. It can be the signal created by a microphone or a telephone set when a word or two is pronounced. In this case, the composite signal cannot be periodic; because that implies that we are repeating the same word or words with exactly the same tone.



a. Time domain

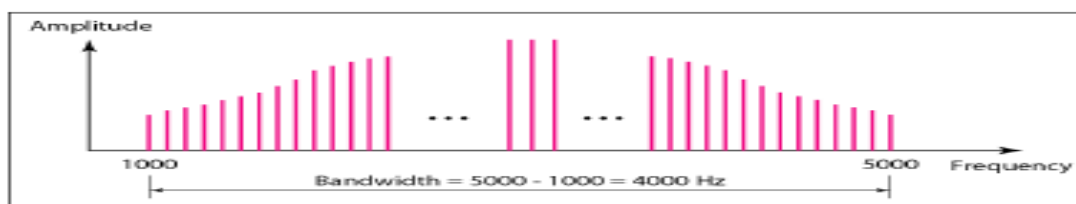


b. Frequency domain

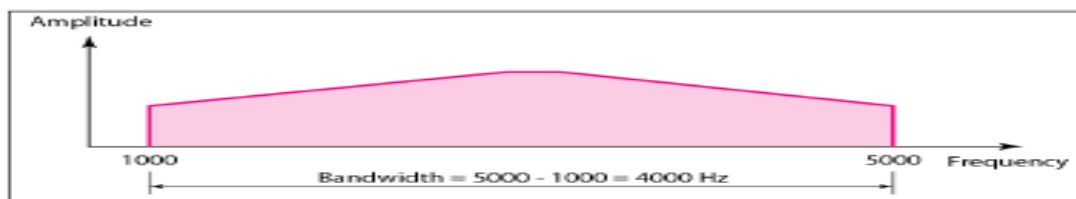
Bandwidth

The bandwidth of a composite signal is the difference between the highest and the lowest frequencies contained in that signal.

The bandwidth of periodic and non periodic composite signals



a. Bandwidth of a periodic signal

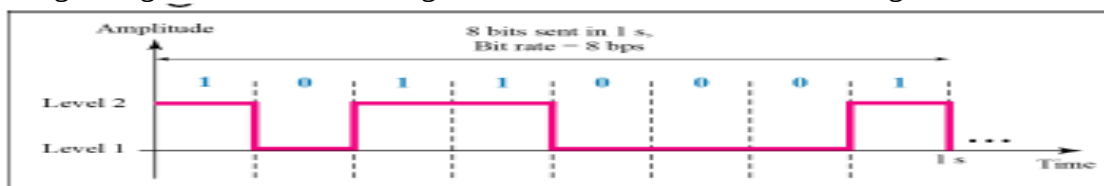


b. Bandwidth of a nonperiodic signal

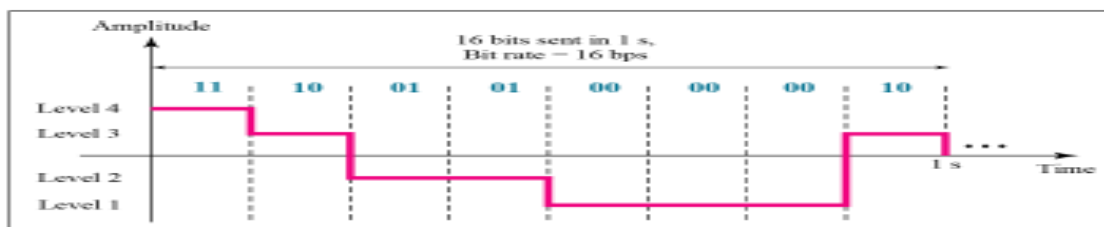
DIGITAL SIGNALS

In addition to being represented by an analog signal information can also be represented by a digital signal. For example, a 1 can be encoded as a positive voltage and a 0 as zero voltage. A digital signal can have more than two levels. In this case,, we can send more than 1 bit for each level.

Two digital signals: one with two signal levels and the other with four signal levels



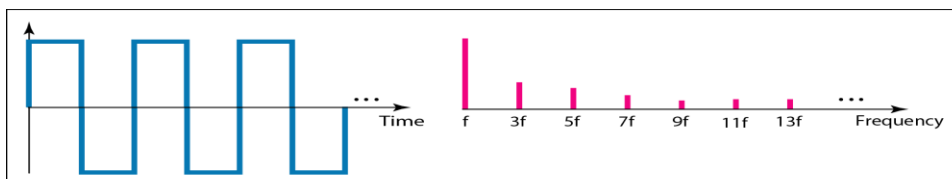
a. A digital signal with two levels



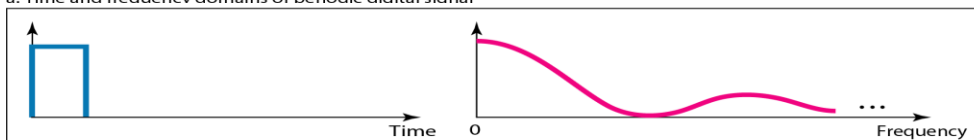
b. A digital signal with four levels

Bit rate and bit interval

Most digital signals are aperiodic, so the period or frequency are not appropriate.



a. Time and frequency domains of periodic digital signal



b. Time and frequency domains of nonperiodic digital signal

- Bit interval (instead of period) and bit rate (instead of frequency) are used to describe digital

signals.

- Bit interval is the time required to send one single bit

- Bit rate is the

number of bit intervals per second—Usually expressed as bits per second (bps)

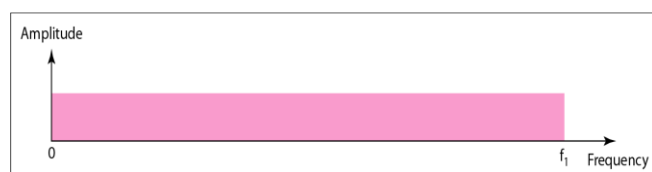
The time and frequency domains of periodic and nonperiodic digital Signals

Baseband transmission

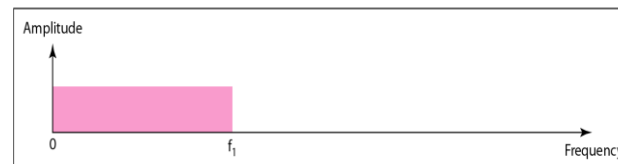
A digital signal is a composite analog signal with an infinite bandwidth.



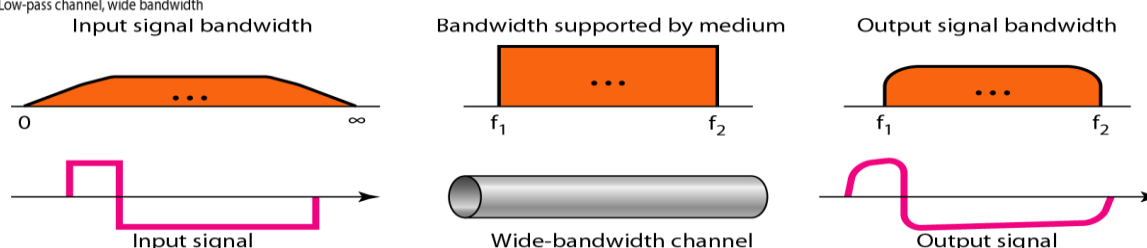
Bandwidths of two low-pass channels



a. Low-pass channel, wide bandwidth

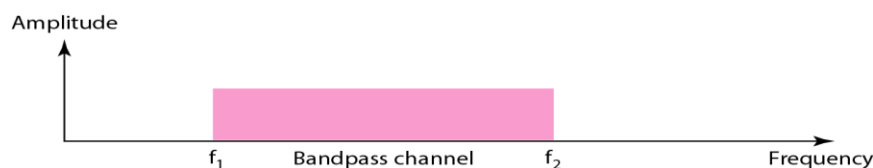


b. Low-pass channel, narrow bandwidth



Baseband transmission using a dedicated medium

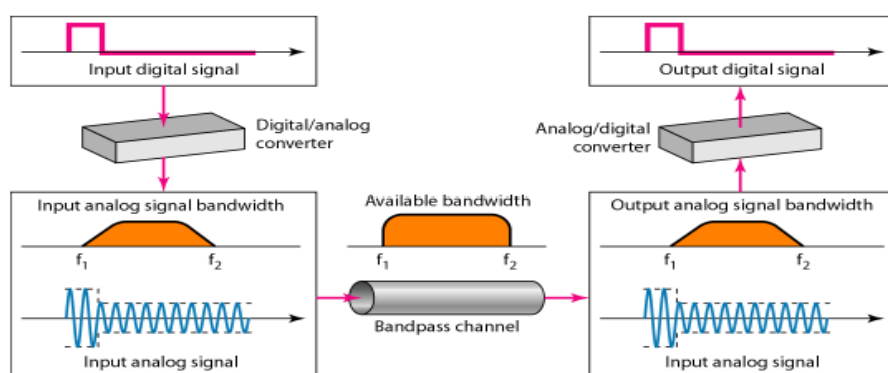
Baseband transmission of a digital signal that preserves the shape of the digital signal is possible only if we have a low-pass channel with an infinite or very wide bandwidth.



we cannot send the digital signal directly to the channel; we need to convert the digital signal to an analog signal before transmission.

Broadband Transmission :

In broadband transmission the signal is converted to analog for transmission. If the available channel is a bandpass channel,



Modulation of a digital signal for transmission on a bandpass channel

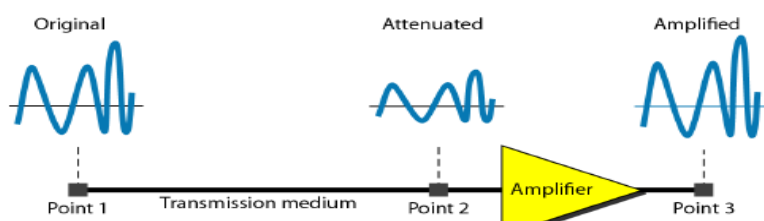
An example of broadband transmission using modulation is the sending of computer data through a telephone subscriber line, the line connecting a resident to the central telephone office.

These lines are designed to carry voice with a limited bandwidth. The channel is considered a bandpass channel. We convert the digital signal from the computer to an analog signal, and send the analog signal. We can install two

converters to change the digital signal to analog and vice versa at the receiving end. The converter, in this case, is called a **modem**.

TRANSMISSION IMPAIRMENT

Signals travel through transmission media, which are not perfect. The imperfection causes signal impairment. This means that the signal at the beginning of the medium is not the same as the signal at the end of the medium. What is sent is not what is received. Three causes of impairment are attenuation, distortion, and noise.



are used to amplify the signal.

Attenuation

Attenuation means loss of energy. When a signal travels through a medium, it loses some of its energy so that it can overcome the resistance of the medium. To compensate for this loss, amplifiers

Decibel

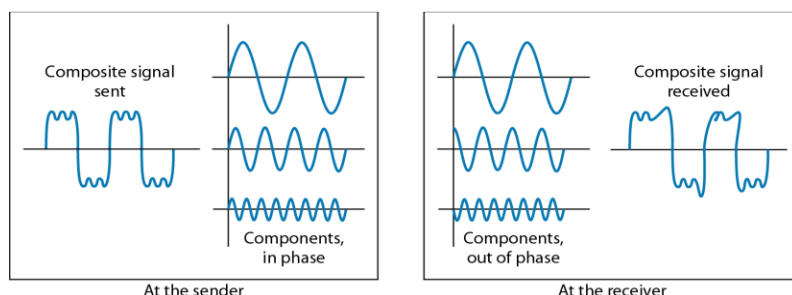
$$\text{dB} = 10 \log_{10} \frac{P_2}{P_1}$$

To show that a signal has lost or gained strength, we use the concept of the decibel (dB).

where P_1 and P_2 are the powers of a signal at points 1 and 2, respectively

- The decibel measures the relative strengths of two signals or a signal at two different points.
- The decibel is negative if a signal is attenuated and positive if a signal is amplified.

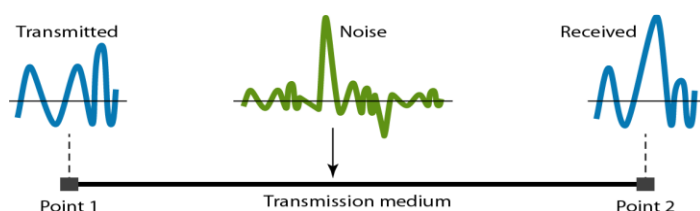
Distortion



Distortion means that the signal changes its form or shape.

- Distortion occurs in a composite signal made of different frequencies.
- Each signal component has its own propagation speed through a medium and therefore its own delay in arriving at the final destination

Noise



Several types of noise such as thermal noise, induced noise, crosstalk and impulse noise may corrupt the signal.

- **Thermal noise** is the random motion of electrons in a wire which creates an extra signal not originally sent by the transmitter.
- **Induced noise** comes from sources such as motors and appliances. These devices act as a sending antenna, and the transmission medium acts as the receiving antenna.
- **Crosstalk** is the effect of one wire on the other. One wire acts as a sending antenna and the other as the receiving antenna.

- **Impulse noise** is a spike (a signal with high energy in a very short time) that comes from power lines, lightning, and so on.

Signal-to-Noise ratio (SNR)

SNR is the statistical ratio of power of the signal to the power of the noise

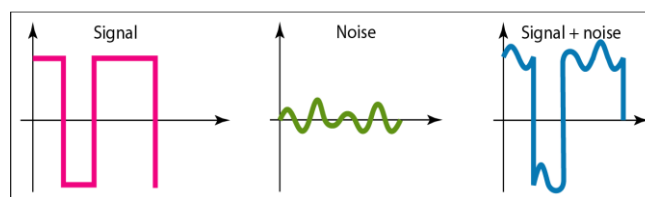
$$\text{SNR} = \frac{\text{average signal power}}{\text{average noise power}}$$

- In decibels it can be expressed as follows:

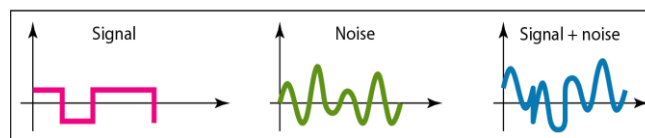
$$\text{SNR}_{\text{dB}} = 10 \log_{10} \text{SNR}$$

Two cases of SNR: a high SNR & a low SNR

A high SNR means the signal is less corrupted by noise; a low SNR means the signal is more corrupted by noise.



a. Large SNR



b. Small SNR

POSSIBLE QUESTIONS**5X8=40 Marks**

- 1) What are the fundamental characteristics and components of data communication
- 2) Describe protocols and standards
- 3) Write short note on analog and digital signals
- 4) Explain types of addressing with neat diagram
- 5) Discuss in briefly about transmission impairments
- 6) Describe categories of networks
- 7) Write short note on network layer in TCP/IP
- 8) Explain OSI reference model with a neat sketch
- 9) what are the four levels of addresses used in internet
- 10) How to measure the network performance explain?



KARPAGAM ACADEMY OF HIGHER EDUCATION
(Deemed to be University Established Under Section 3 of UGC Act 1956)
 Pollachi Main Road, Eachanari Post, Coimbatore – 641 021. INDIA

DEPARTMENT OF COMPUTER SCIENCE
DATA COMMUNICATION NETWORKS (15CSU502)
ONLINE EXAM QUESTION BANK

UNIT-1

Questions	opt1	opt2	opt3	opt4	Answer
Data communication means exchange of data between _____ devices.	one	two	six	four	two
The system must deliver data to the correct destination is called _____	accuracy	jitter	delivery	timeliness	delivery
A _____ is the set of rules.	protocols	transmission medium	networks	ip	protocols
In _____, the communication is unidirection.	duplex mode	full duplex mode	half duplex mode	simplex mode	simplex mode
A _____ is a set of devices connected by communication links.	protocols	networks	computer	printer	networks
A _____ connection provides a dedicated link between two devices.	point-to-point	multi-point	mesh	physical	point-to-point
One long cable acts as a _____ to link all the devices in a network.	bus	mesh	hub	backbone	backbone
MAN stands for _____	metropolitician area network	metropolitan area network	metropolitcal area network	macro area network	metropolitan area network
The term timing refers to _____ characteristics.	two	three	four	six	two
_____ standards are often established originally by manufactures.	de jure	de facto	de fact	semantics	de facto
In physical layer we can transfer data into _____	frame	packet	bit	sp du	bit
Hob to hob delivery is done by the _____	session layer	datalink layer	network layer	transport layer	datalink layer
The _____ layer is responsible for process to process delivery.	physical	presentation	networks	transport	transport

The _____ layer is responsible for dialog control and synchronization.	transport	session	application	presentation	session
Tcp/Ip is a _____ protocol.	hyper text	transfer	internet	hierarchical	hierarchical
Ip is a _____ protocol.	hop to hop	node to node	process to process	host to host	host to host
A set of devices connected by a _____ links	data	networks	communication	application	communication
Bus topology has a long link called _____	backbone	hub	host	hop	backbone
Decibel (dB) = _____	$10 \log_{10} \frac{p_2}{p_1}$	p_1/p_2	$10 \log_{10} \frac{p_1}{p_2}$	$2 \log_{10} \frac{p_1}{p_2}$	$10 \log_{10} \frac{p_2}{p_1}$
Transmission time= _____	message size/birate	distance/band width	message size/distance	message size/bandwidth	message size/bandwidth
_____ and star is a point to point device.	bus	ring	mesh	physical	mesh
Protocols can be classified into _____ key elements	one	three	four	two	three
_____ is a basic key element.	protocols	standards	topology	protocols and standards	protocols and standards
OSI stands for _____	open systems interconnection	open system internetworking	open symantic interconnection	open system internet	open systems interconnection
Net work layer delivers data in the form of _____	frame	bits	data	packet	packet
Session layer provides _____ services.	one	two	three	four	two
UDP _____	user data protocol	user datagram protocol	user defined protocol	user dataframe protocol	user datagram pr
FTP _____	file transmit protocol	file transmission protocol	file transfer protocol	flip transfer protocol	file transfer prote
SMTP _____	single mail transfer protocol	simple mail transfer protocol	simple mail transmissi	single mail transmit protocol	simple mail trans
Complete a cycle is called as _____	period	frequency	non periodic	periodic	period

Jitter is a form of _____	frames	bits	packets	dp tu	packets
Each set is called a _____	node	code	unicode	polar	node
Full duplex also called as _____	simple duplex	single duplex	multiple duplex	duplex	duplex
_____ can be measured in transmit time and response time.	performance	frequency	period	non period	performance
A multipoint is also called as _____	multi line	multi drop	multi level	single level	multi drop
Mesh topology we need _____	$n(n-1)$	$n(n+1)$	$n(n+1)/2$	$n(n-1)/2$	$n(n-1)/2$
A _____ topology on the other hand is multipoint.	star	ring	bus	mesh	bus
A _____ can be hybrid	physical	networks	data	link	networks
A MAN is a network with a size between a _____ and _____.	WAN and LAN	WAN or LAN	LAN	WAN	WAN and LAN
When Two or more networks are connected they become an _____	network	inter network	internet connection	interconnection	inter network
The _____ layer is responsible for providing services to the user.	presentation	datalink	application	network	application
The _____ layer is responsible for translation, compression encryption.	transport	data link	presentation	application	presentation
The _____ layer is responsible for the delivery of a message from one process to another.	data link	transport	presentation	network	transport
A _____ layer is responsible for the delivery of packets from the source to destination.	physical	data link	network	session	network
The _____ layer is responsible for moving frames from one hop to the next.	data link	physical	network	presentation	data link
The _____ layer is responsible for movements of bits from one hop to next.	data link	physical	transport	session	physical

RARP _____	reverse address resolution protocol	reverse address result protocol	reverse address revolutiniz ed protocol	reverse address research protocol	reverse address resolution protocol
_____ does not define any specific protocol.	TCP	HTTP	TCP/IP	SMTP	TCP/IP
The TCP/IP protocol suite was developed prior to the _____ model.	OSI	ISO	TCP	IP	OSI
The _____ layer is responsible for flow control.	session	presentation	application	transport	transport

otocol

ocol

ifer protocol

UNIT-II**UNIT-II-SYLLABUS**

Physical Layer: Multiplexing – Frequency Division Multiplexing-Wavelength Division Multiplexing– Synchronous Time-Division Multiplexing– Statistical Time Division Multiplexing. Transmission Media - Guided Media- Twisted pair and coaxial cable - Fiber optic cable-Unguided Transmission Media Switching – Circuit Switched Networks- Datagram Networks – Virtual Circuit networks.

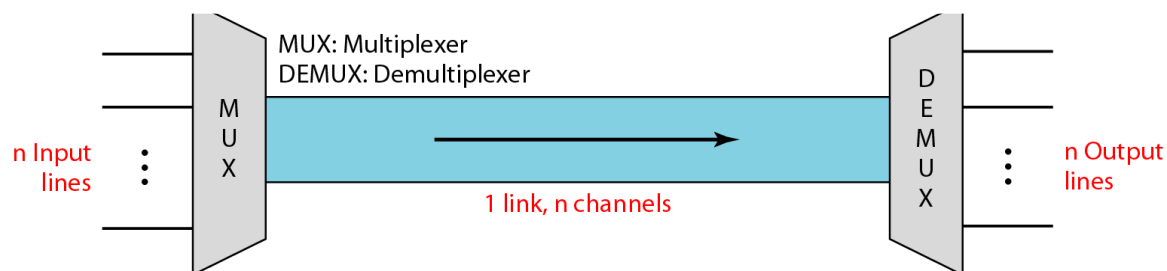
UNIT-II

MULTIPLEXING

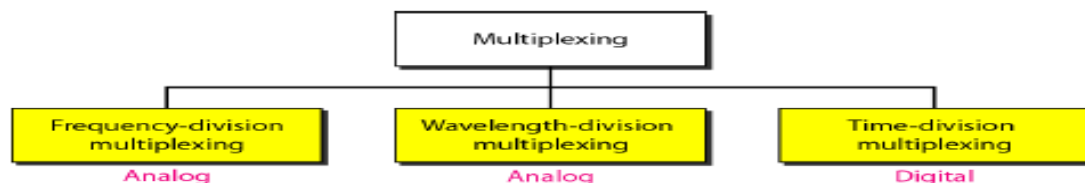
Whenever the bandwidth of a medium linking two devices is greater than the bandwidth needs of the devices, the link can be shared. Multiplexing is the set of techniques that allows the simultaneous transmission of multiple signals across a single data link. As data and telecommunications use increases, so does traffic.

Dividing a link into channels

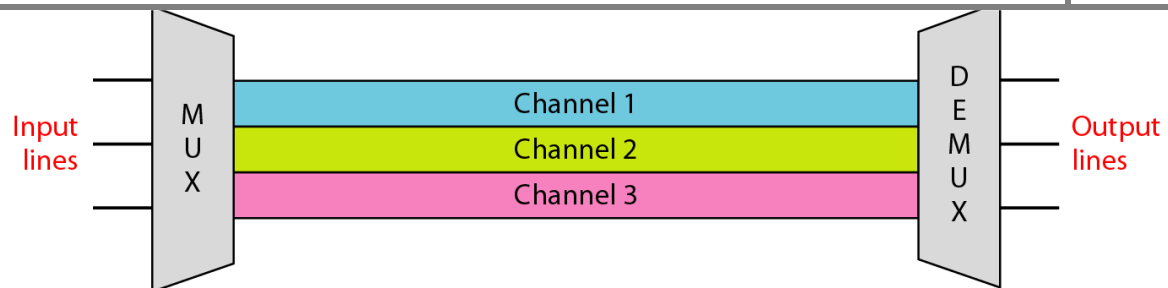
In a multiplexed system, n lines share the bandwidth of one link. Figure shows the basic format of a multiplexed system. The lines on the left direct their transmission streams to a multiplexer (MUX), which combines them into a single stream (many-to-one). At the receiving end, that stream is fed into a demultiplexer (DEMUX), which separates the stream back into its component transmissions (one-to-many) and directs them to their corresponding lines. In the figure, the word link refers to the physical path. The word channel refers to the portion of a link that carries a transmission between a given pair of lines. One link can have many (n) channels.

**Categories of multiplexing**

There are three basic multiplexing techniques: frequency-division multiplexing, wavelength-division multiplexing, and time-division multiplexing. The first two are techniques designed for analog signals, the third, for digital signals.

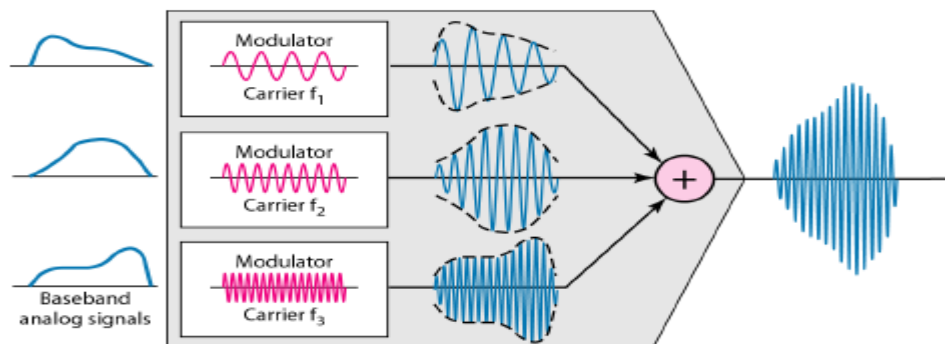
**Frequency Division Multiplexing (FDM)**

Frequency-division multiplexing (FDM) is an analog technique that can be applied when the bandwidth of a link (in hertz) is greater than the combined bandwidths of the signals to be transmitted. In FDM, signals generated by each sending device modulate different carrier frequencies. These modulated signals are then combined into a single composite signal that can be transported by the link. Carrier frequencies are separated by sufficient bandwidth to accommodate the modulated signal. These bandwidth ranges are the channels through which the various signals travel. Channels can be separated by strips of unused bandwidth-guard bands-to prevent signals from overlapping. In addition, carrier frequencies must not interfere with the original data frequencies. Figure gives a conceptual view of FDM. In this illustration, the transmission path is divided into three parts, each representing a channel that carries one transmission.



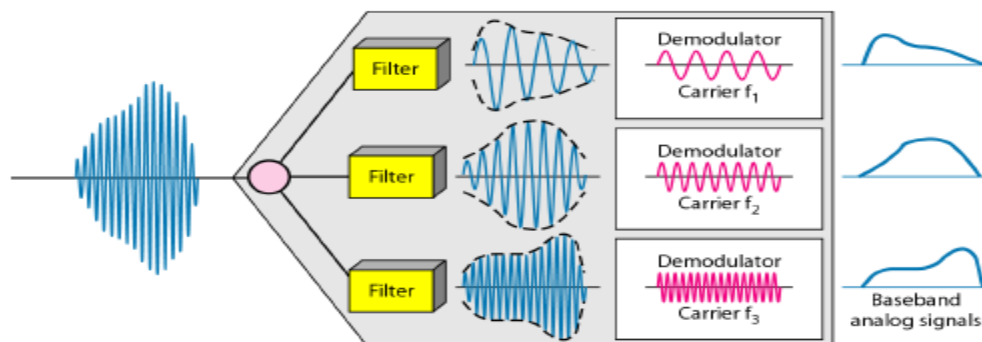
Multiplexing Process

Figure below is a conceptual illustration of the multiplexing process. Each source generates a signal of a similar frequency range. Inside the multiplexer, these similar signals modulates different carrier frequencies. The resulting modulated signals are then combined into a single composite signal that is sent out over a media link that has enough bandwidth to accommodate it.



Demultiplexing Process

The demultiplexer uses a series of filters to decompose the multiplexed signal into its constituent component signals. The individual signals are then passed to a demodulator that separates them from their carriers and passes them to the output lines.

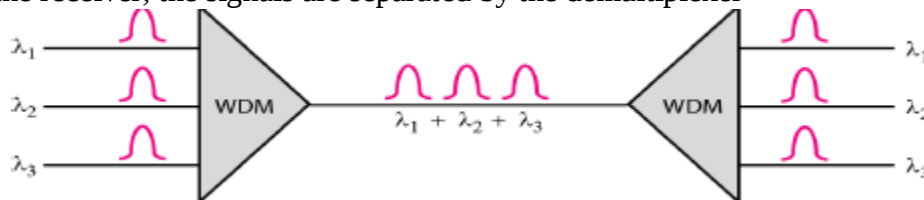


A very common application of FDM is AM and FM radio broadcasting. Radio uses the air as the transmission medium. A special band from 530 to 1700 kHz is assigned to AM radio. All radio stations need to share this band. Each AM station needs 10kHz of bandwidth. Each station uses a different carrier frequency, which means it is shifting its signal and multiplexing. The signal that goes to the air is a combination of signals. A receiver receives all these signals, but filters (by tuning) only the one which is desired. Without multiplexing, only one AM station could broadcast to the common link, the air. The situation is similar in FM broadcasting. However, FM has a wider band of 88 to 108MHz because each station needs a bandwidth of 200 kHz. Another common use of FDM is in television broadcasting. Each TV channel has its own bandwidth of 6 MHz.

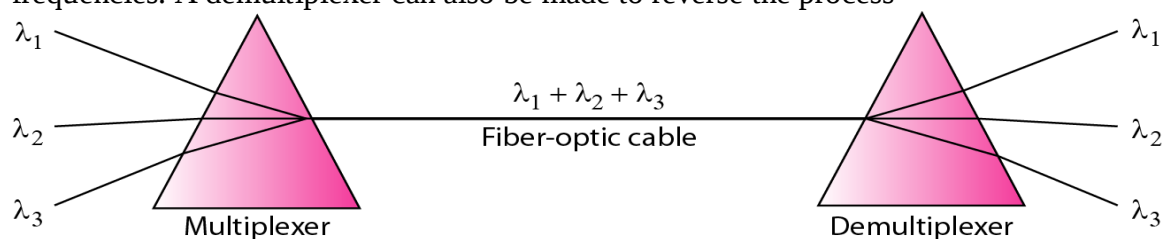
Wavelength Division Multiplexing (WDM)

Wavelength-division multiplexing (WDM) is designed to use the high-data-rate capability of fiber-optic cable. The optical fiber data rate is higher than the data rate of metallic transmission cable. Using a fiber-optic cable for one single line wastes the available bandwidth. Multiplexing allows us to combine several lines into WDM is conceptually the same as FDM, except that the multiplexing and demultiplexing involve optical signals transmitted through fiber-optic channels. The idea is the same: We are combining different signals of different frequencies. The difference is that the frequencies are very high. Figure gives a conceptual view of a WDM multiplexer and demultiplexer.

Very narrow bands of light from different sources are combined to make a wider band of light. At the receiver, the signals are separated by the demultiplexer

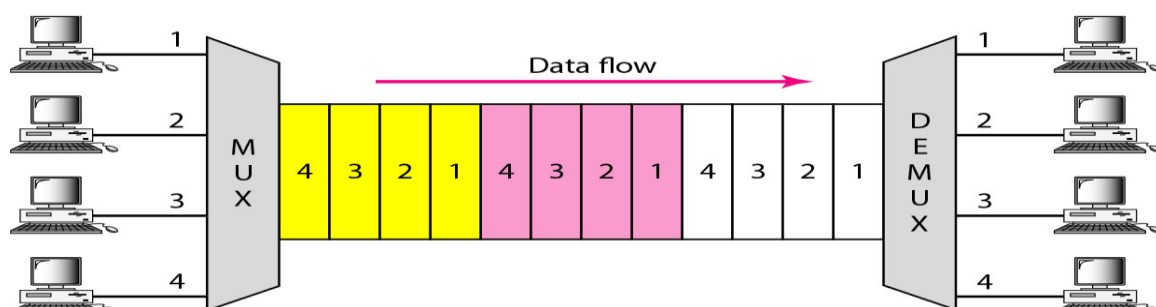


Although WDM technology is very complex, the basic idea is very simple. We want to combine multiple light sources into one single light at the multiplexer and do the reverse at the demultiplexer. The combining and splitting of light sources are easily handled by a prism. Recall from basic physics that a prism bends a beam of light based on the angle of incidence and the frequency. Using this technique, a multiplexer can be made to combine several input beams of light, each containing a narrow band of frequencies, into one output beam of a wider band of frequencies. A demultiplexer can also be made to reverse the process



Time Division Multiplexing (TDM)

Time-division multiplexing (TDM) is a digital process that allows several connections to share the high bandwidth of a line. Instead of sharing a portion of the bandwidth as in FDM, time is shared. Each connection occupies a portion of time in the link. Figure gives a conceptual view of TDM. Note that the same link is used as in FDM; here, however, the link is shown sectioned by time rather than by frequency. In the figure, portions of signals 1, 2, 3, and 4 occupy the link sequentially



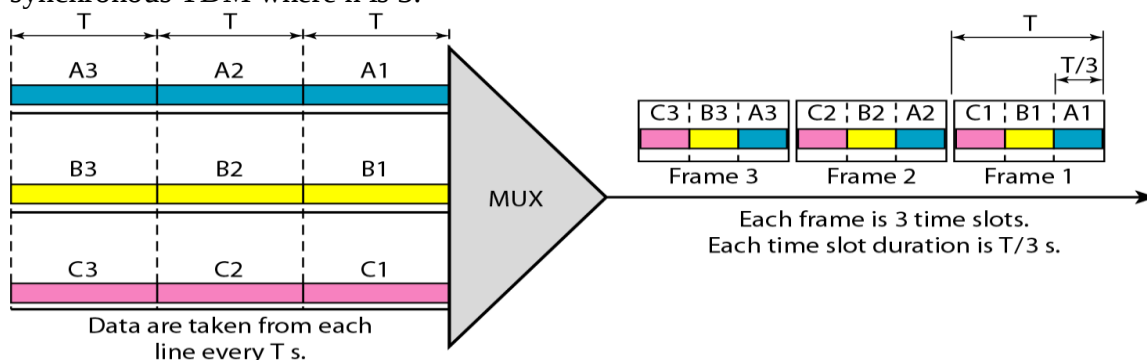
Digital data from different sources are combined into one timeshared link. However, this does not mean that the sources cannot produce analog data; analog data can be sampled, changed to digital data, and then multiplexed by using TDM.

Synchronous Time Division Multiplexing

In synchronous TDM, each input connection has an allotment in the output even if it is not sending data.

Time Slots and Frames

In synchronous TDM, the data flow of each input connection is divided into units, where each input occupies one input time slot. A unit can be 1 bit, one character, or one block of data. Each input unit becomes one output unit and occupies one output time slot. However, the duration of an output time slot is n times shorter than the duration of an input time slot. If an input time slot is T s, the output time slot is T/n s, where n is the number of connections. In other words, a unit in the output connection has a shorter duration; it travels faster. Figure shows an example of synchronous TDM where n is 3.

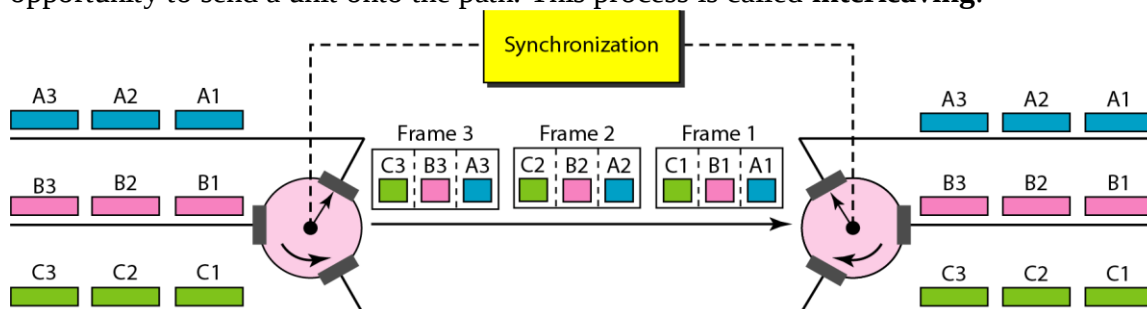


In synchronous TDM, a round of data units from each input connection is collected into a frame. If we have n connections, a frame is divided into n time slots and one slot is allocated for each unit, one for each input line. If the duration of the input unit is T , the duration of each slot is T/n and the duration of each frame is T . The data rate of the output link must be n times the data rate of a connection to guarantee the flow of data. In Figure, the data rate of the link is 3 times the data rate of a connection; likewise, the duration of a unit on a connection is 3 times that of the time slot (duration of a unit on the link). In the figure we represent the data prior to multiplexing as 3 times the size of the data after multiplexing. This is just to convey the idea that each unit is 3 times longer in duration before multiplexing than after.

Time slots are grouped into frames. A frame consists of one complete cycle of time slots, with one slot dedicated to each sending device. In a system with n input lines, each frame has n slots, with each slot allocated to carrying data from a specific input line.

Interleaving

TDM can be visualized as two fast rotating switches, one on the MUX side and the other on the DEMUX side. The switches are synchronized and rotate at the same speed but in opposite directions. On the MUX side, as the switch opens in front of a connection, that connection has the opportunity to send a unit onto the path. This process is called **interleaving**.



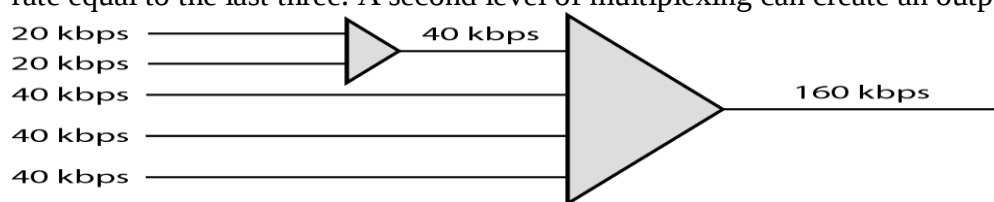
Empty Slots

Synchronous TDM is not as efficient as it could be. If a source does not have data to send, the corresponding slot in the output frame is empty. Statistical TDM can improve the efficiency by removing the empty slots from the frame.

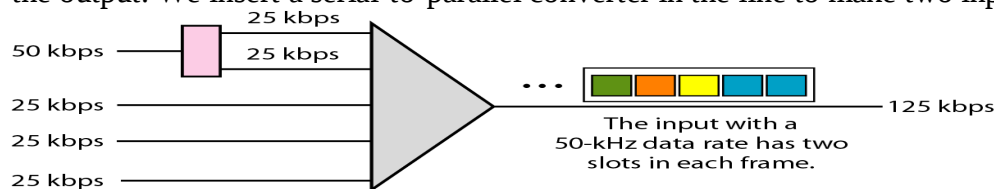
Data Rate Management

One problem with TDM is how to handle a disparity in the input data rates. If data rates are not the same, three strategies, or a combination of them, can be used. The three strategies are **multilevel multiplexing**, **multiple-slot allocation**, and **pulse stuffing**.

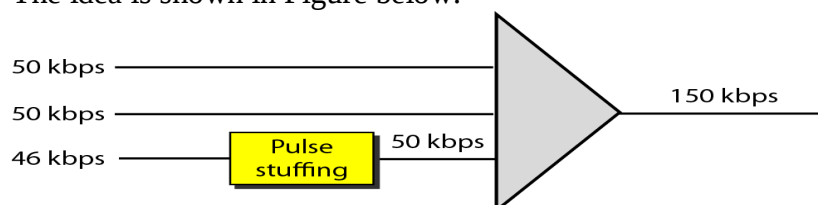
Multilevel Multiplexing Multilevel multiplexing is a technique used when the data rate of an input line is a multiple of others. For example, in Figure below, we have two inputs of 20 kbps and three inputs of 40 kbps. The first two input lines can be multiplexed together to provide a data rate equal to the last three. A second level of multiplexing can create an output of 160 kbps.



Multiple-Slot Allocation Sometimes it is more efficient to allot more than one slot in a frame to a single input line. For example, we might have an input line that has a data rate that is a multiple of another input. In Figure below, the input line with a 50-kbps data rate can be given two slots in the output. We insert a serial-to-parallel converter in the line to make two inputs out of one.

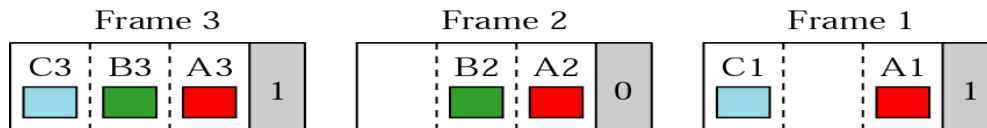


Pulse Stuffing Sometimes the bit rates of sources are not multiple integers of each other. Therefore, neither of the above two techniques can be applied. One solution is to make the highest input data rate the dominant data rate and then add dummy bits to the input lines with lower rates. This will increase their rates. This technique is called pulse stuffing, bit padding, or bit stuffing. The idea is shown in Figure below.

**Frame Synchronizing**

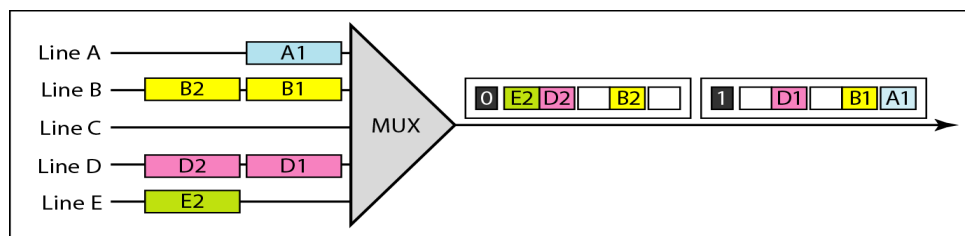
The implementation of TDM is not as simple as that of FDM. Synchronization between the multiplexer and demultiplexer is a major issue. If the multiplexer and the demultiplexer are not synchronized, a bit belonging to one channel may be received by the wrong channel. For this reason, one or more synchronization bits are usually added to the beginning of each frame. These bits, called framing bits, follow a pattern, frame to frame, that allows the demultiplexer to synchronize with the incoming stream so that it can separate the time slots accurately. In most cases, this synchronization information consists of 1 bit per frame, alternating between 0 and 1, as shown in Figure below

Synchronization pattern

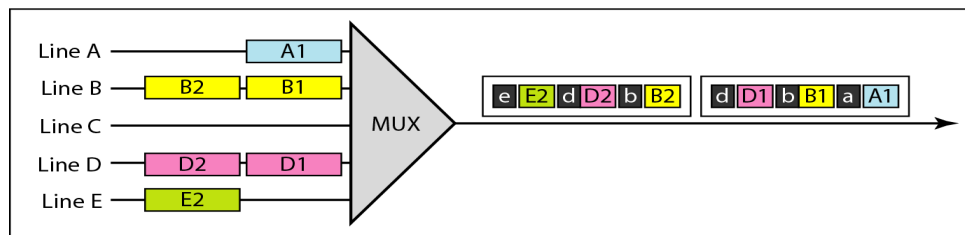
**Statistical Time-Division Multiplexing**

In synchronous TDM, each input has a reserved slot in the output frame. This can be inefficient if some input lines have no data to send. In statistical time-division multiplexing, slots are dynamically allocated to improve bandwidth efficiency. Only when an input line has a slot's worth of data to send is it given a slot in the output frame. In statistical multiplexing, the number of slots in each frame is less than the number of input lines. The multiplexer checks each input line in round-robin fashion; it allocates a slot for an input line if the line has data to send; otherwise, it skips the line and checks the next line.

Figure below shows a synchronous and a statistical TDM example. In the former, some slots are empty because the corresponding line does not have data to send. In the latter, however, no slot is left empty as long as there are data to be sent by any input line.



a. Synchronous TDM



b. Statistical TDM

Addressing

An output slot in synchronous TDM is totally occupied by data; in statistical TDM, a slot needs to carry data as well as the address of the destination. In synchronous TDM, there is no need for addressing; synchronization and preassigned relationships between the inputs and outputs serve as an address. In statistical multiplexing, there is no fixed relationship between the inputs and outputs because there are no preassigned or reserved slots. The addressing in its simplest form can be n bits to define N different output lines with $n = \log_2 N$. For example, for eight different output lines, we need a 3-bit address.

Slot Size

Since a slot carries both data and an address in statistical TDM, the ratio of the data size to address size must be reasonable to make transmission efficient. For example, it would be inefficient to send 1 bit per slot as data when the address is 3 bits. This would mean an overhead

of 300 percent. In statistical TDM, a block of data is usually many bytes while the address is just a few bytes.

No Synchronization Bit

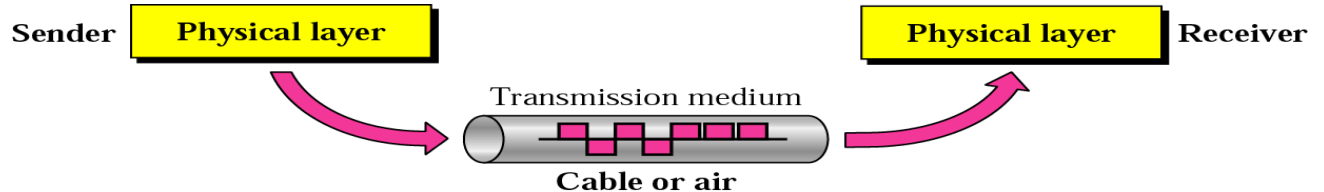
The frames in statistical TDM need not be synchronized, so no need for synchronization bits.

Bandwidth

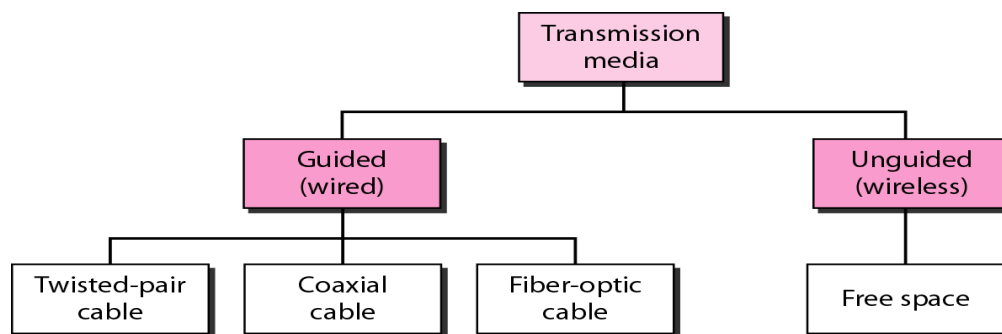
In statistical TDM, the capacity of the link is normally less than the sum of the capacities of each channel.

Transmission Media

A transmission **medium** can be broadly defined as anything that can carry information from a source to a destination.



Transmission media can be divided into two broad categories: guided and unguided. Guided media include twisted-pair cable, coaxial cable, and fiber-optic cable. Unguided medium is free space.

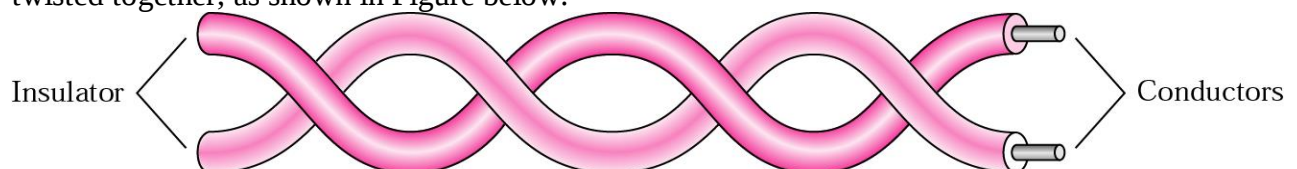


GUIDED MEDIA

Guided media, which are those that provide a conduit from one device to another, include twisted-pair cable, coaxial cable, and fiber-optic cable. A signal traveling along any of these media is directed and contained by the physical limits of the medium. Twisted-pair and coaxial cable use metallic (copper) conductors that accept and transport signals in the form of electric current. Optical fiber is a cable that accepts and transports signals in the form of light.

Twisted-Pair Cable

A twisted pair consists of two conductors (normally copper), each with its own plastic insulation, twisted together, as shown in Figure below.

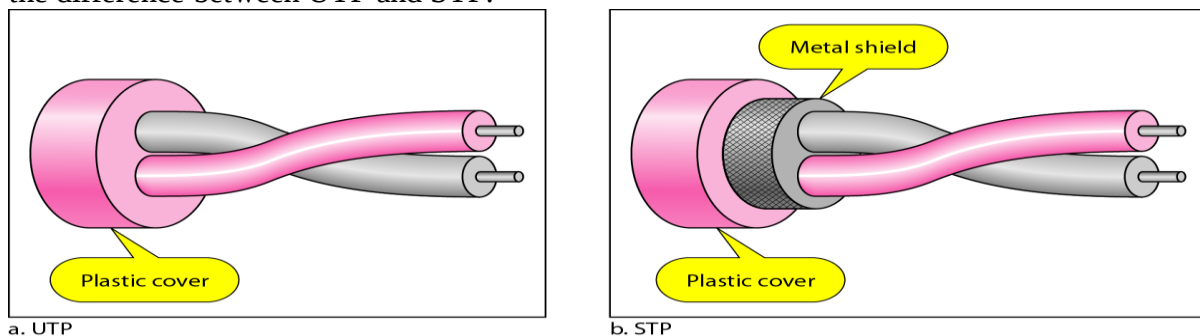


One of the wires is used to carry signals to the receiver, and the other is used only as a ground reference. The receiver uses the difference between the two. In addition to the signal sent by the sender on one of the wires, interference (noise) and crosstalk may affect both wires and create unwanted signals. If the two wires are parallel, the effect of these unwanted signals is not the same in both wires because they are at different locations relative to the noise or crosstalk sources (e.g., one is closer and the other is farther). This results in a difference at the receiver. By twisting the pairs, a balance is maintained.

This means that the receiver, which calculates the difference between the two, receives no unwanted signals. The unwanted signals are mostly canceled out. From the above discussion, it is clear that the number of twists per unit of length (e.g., inch) has some effect on the quality of the cable.

Unshielded Versus Shielded Twisted-Pair Cable

The most common twisted-pair cable used in communications is referred to as unshielded twisted-pair (UTP). IBM has also produced a version of twisted-pair cable for its use called shielded twisted-pair (STP). STP cable has a metal foil or braided mesh covering that encases each pair of insulated conductors. Although metal casing improves the quality of cable by preventing the penetration of noise or crosstalk, it is bulkier and more expensive. Figure shows the difference between UTP and STP.



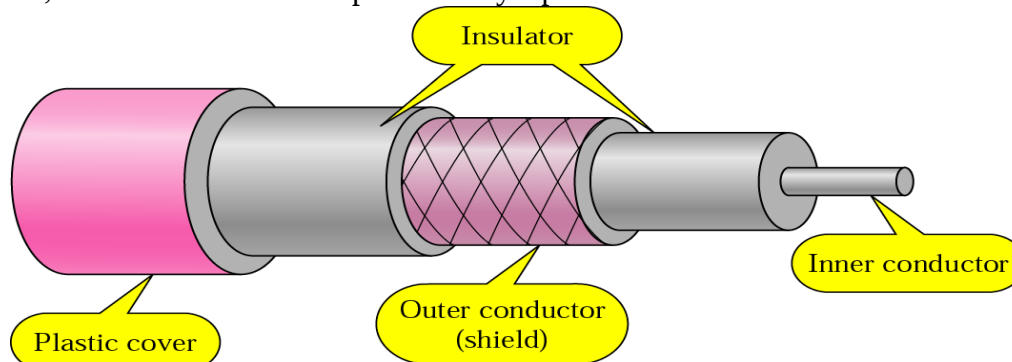
Categories of UTP Cables

Category	Specification	Data Rate (Mbps)	Use
1	Unshielded twisted-pair used in telephone	< 0.1	Telephone
2	Unshielded twisted-pair originally used in T-lines	2	T-1 lines
3	Improved CAT 2 used in LANs	10	LANs
4	Improved CAT 3 used in Token Ring networks	20	LANs
5	Cable wire is normally 24 AWG with a jacket and outside sheath	100	LANs
5E	An extension to category 5 that includes extra features to minimize the crosstalk and electromagnetic interference	125	LANs
6	A new category with matched components coming from the same manufacturer. The cable must be tested at a 200-Mbps data rate.	200	LANs
7	Sometimes called SSTP (shielded screen twisted-pair). Each pair is individually wrapped in a helical metallic foil followed by a metallic foil shield in addition to the outside sheath. The shield decreases the effect of crosstalk and increases the data rate.	600	LANs

Coaxial Cable

Coaxial cable (or *coax*) carries signals of higher frequency ranges than those in twisted pair cable, in part because the two media are constructed quite differently. Instead of having two wires, coax has a central core conductor of solid or stranded wire (usually copper) enclosed in an insulating sheath, which is, in turn, encased in an outer conductor of metal foil, braid, or a combination of the two. The outer metallic wrapping serves both as a shield against noise and as the second

conductor, which completes the circuit. This outer conductor is also enclosed in an insulating sheath, and the whole cable is protected by a plastic cover.



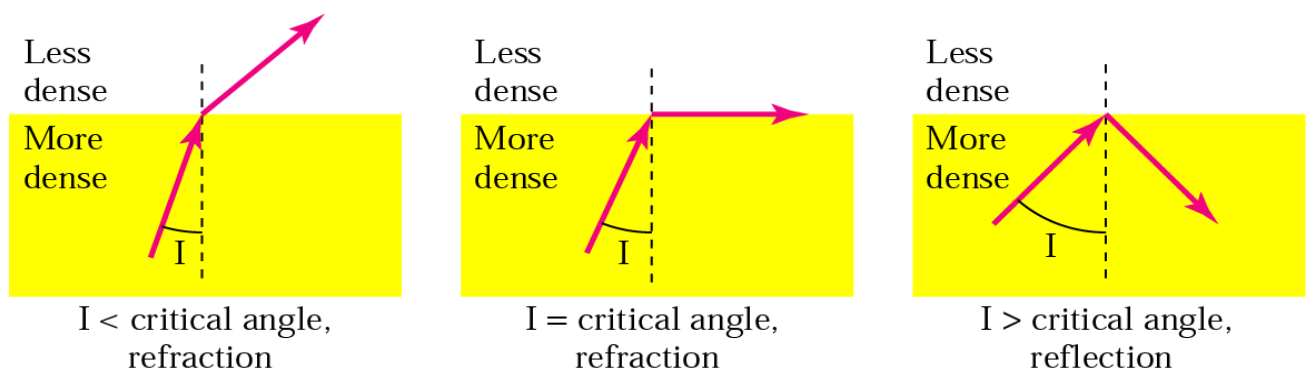
Categories of Coaxial Cables

Category	Impedance	Use
RG-59	75 Ω	Cable TV
RG-58	50 Ω	Thin Ethernet
RG-11	50 Ω	Thick Ethernet

Coaxial cables are categorized by their radio government (RG) ratings. Each RG number denotes a unique set of physical specifications, including the wire gauge of the inner conductor, the thickness and type of the inner insulator, the construction of the shield, and the size and type of the outer casing. Each cable defined by an RG rating is adapted for a specialized function,

Fiber-Optic Cable

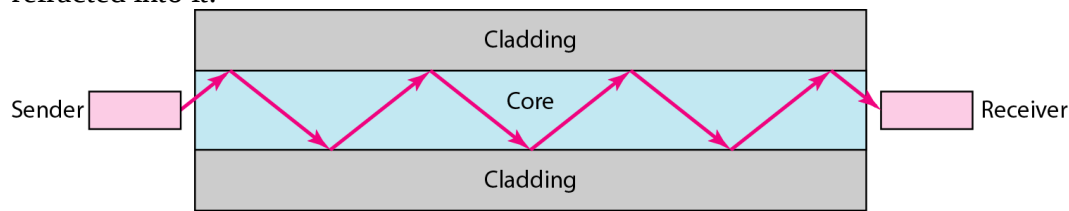
A fiber-optic cable is made of glass or plastic and transmits signals in the form of light. To understand optical fiber, we first need to explore several aspects of the nature of light. Light travels in a straight line as long as it is moving through a single uniform substance. If a ray of light traveling through one substance suddenly enters another substance (of a different density), the ray changes its direction. Figure below shows how a ray of light changes direction when going from a denser to a less dense substance.



As the figure shows, if the angle of incidence I (the angle the ray makes with the line perpendicular to the interface between the two substances) is less than the critical angle, the ray refracts and moves closer to the surface. If the angle of incidence is equal to the critical angle, the light bends along the interface. If the angle is greater than the critical angle, the ray reflects (makes a turn) and travels again in the denser substance. Note that the critical angle is a property of the substance, and its value differs from one substance to another.

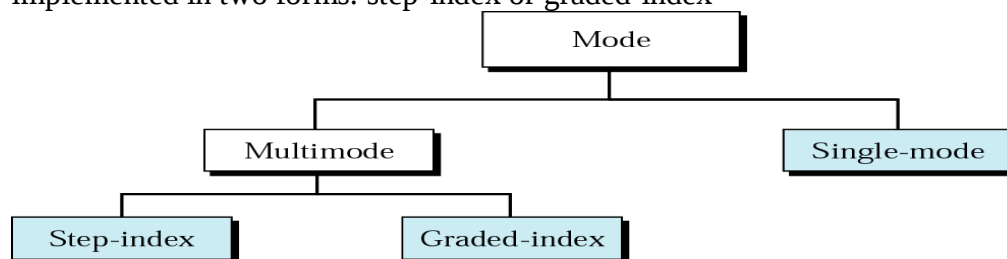
Optical fibers use reflection to guide light through a channel. A glass or plastic core is surrounded by a cladding of less dense glass or plastic. The difference in density of the two materials must be

such that a beam of light moving through the core is reflected off the cladding instead of being refracted into it.



Propagation Modes

Current technology supports two modes (multimode and single mode) for propagating light along optical channels, each requiring fiber with different physical characteristics. Multimode can be implemented in two forms: step-index or graded-index



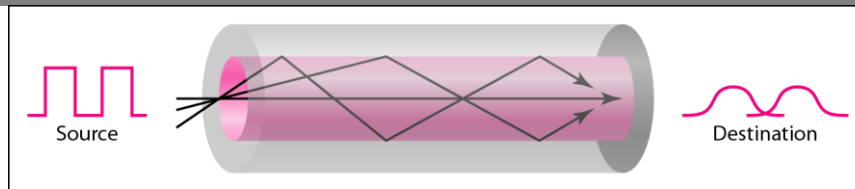
Multimode:

Multimode is so named because multiple beams from a light source move through the core in different paths. How these beams move within the cable depends on the structure of the core. In multimode step-index fiber, the density of the core remains constant from the center to the edges. A beam of light moves through this constant density in a straight line until it reaches the interface of the core and the cladding. At the interface, there is an abrupt change due to a lower density; this alters the angle of the beam's motion. The term *step index* refers to the suddenness of this change, which contributes to the distortion of the signal as it passes through the fiber.

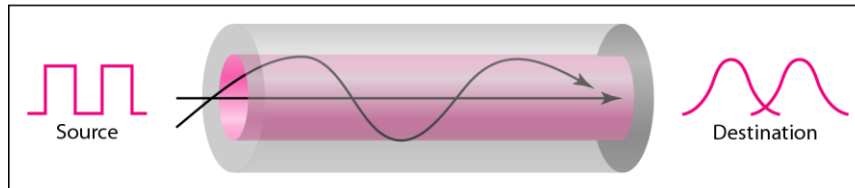
A second type of fiber, called multimode graded-index fiber, decreases this distortion of the signal through the cable. The word *index* here refers to the index of refraction. The index of refraction is related to density. A graded-index fiber, therefore, is one with varying densities. Density is highest at the center of the core and decreases gradually to its lowest at the edge. Figure below shows the impact of this variable density on the propagation of light beams.

Single-Mode:

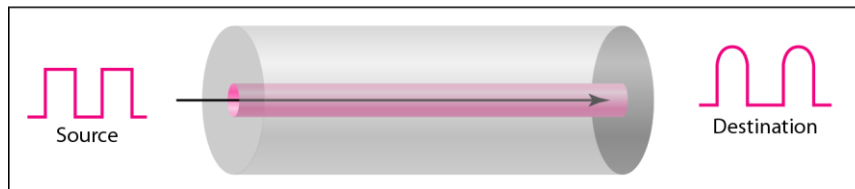
Single-mode uses step-index fiber and a highly focused source of light that limits beams to a small range of angles, all close to the horizontal. The single-mode fiber is manufactured with a much smaller diameter than that of multimode fiber, and with substantially lower density (index of refraction). The decrease in density results in a critical angle that is close enough to 90° to make the propagation of beams almost horizontal. In this case, propagation of different beams is almost identical, and delays are negligible. All the beams arrive at the destination "together" and can be recombined with little distortion to the signal.



a. Multimode, step index



b. Multimode, graded index



c. Single mode

Advantages: Fiber-optic cable has several advantages over metallic cable.

- Higher bandwidth. Fiber-optic cable can support dramatically higher bandwidths (and hence data rates) than either twisted-pair or coaxial cable. Currently, data rates and bandwidth utilization over fiber-optic cable are limited not by the medium but by the signal generation and reception technology available.
- Less signal attenuation. Fiber-optic transmission distance is significantly greater than that of other guided media. A signal can run for 50 km without requiring regeneration. We need repeaters every 5 km for coaxial or twisted-pair cable.
- Immunity to electromagnetic interference. Electromagnetic noise cannot affect fiber-optic cables.
- Resistance to corrosive materials. Glass is more resistant to corrosive materials than copper.
- Light weight. Fiber-optic cables are much lighter than copper cables.
- Greater immunity to tapping. Fiber-optic cables are more immune to tapping than copper cables. Copper cables create antenna effects that can easily be tapped.

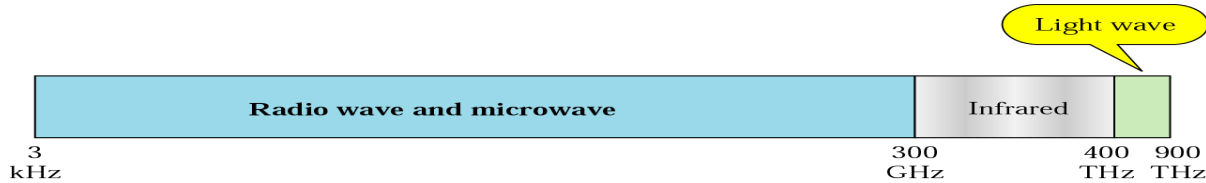
Disadvantages: There are some disadvantages in the use of optical fiber.

- Installation and maintenance. Fiber-optic cable is a relatively new technology. Its installation and maintenance require expertise that is not yet available everywhere.
- Unidirectional light propagation. Propagation of light is unidirectional. If we need bidirectional communication, two fibers are needed.
- Cost. The cable and the interfaces are relatively more expensive than those of other guided media. If the demand for bandwidth is not high, often the use of optical fiber cannot be justified

UNGUIDED MEDIA: WIRELESS

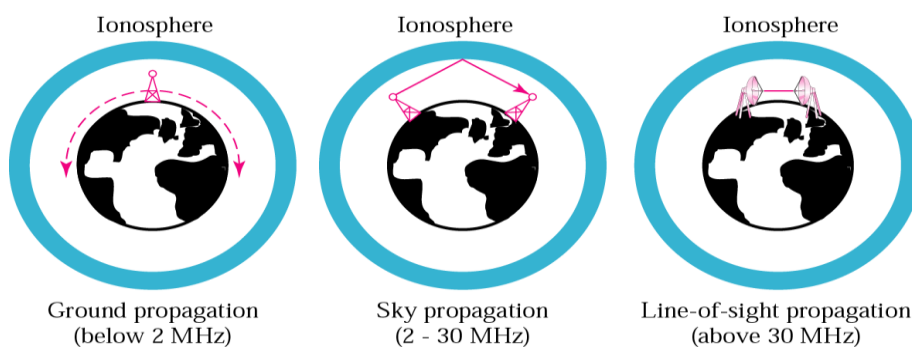
Unguided media transport electromagnetic waves without using a physical conductor. This type of communication is often referred to as wireless communication. Signals are normally broadcast through free space and thus are available to anyone who has a device capable of receiving them.

Figure below shows the part of the electromagnetic spectrum, ranging from 3 kHz to 900 THz, used for wireless communication.



Unguided signals can travel from the source to destination in several ways: ground propagation, sky propagation, and line-of-sight propagation. In ground propagation, radio waves travel through the lowest portion of the atmosphere, hugging the earth. These low-frequency signals emanate in all directions from the transmitting antenna and follow the curvature of the planet. Distance depends

on the amount of power in the signal: The greater the power, the greater the distance. In sky

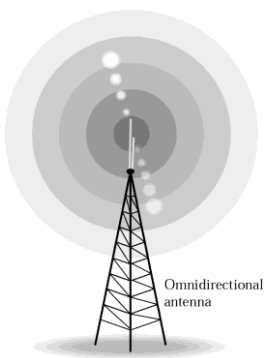


propagation, higher-frequency radio waves radiate upward into the ionosphere (the layer of atmosphere where particles exist as ions) where they are reflected back to earth. This type of transmission allows for greater distances with lower output power. In

line-of-sight propagation, very high-frequency signals are transmitted in straight lines directly from antenna to antenna. Antennas must be directional, facing each other and either tall enough or close enough together not to be affected by the curvature of the earth. Line-of-sight propagation is tricky because radio transmissions cannot be completely focused.

The section of the electromagnetic spectrum defined as radio waves and microwaves is divided into eight ranges, called *bands*, each regulated by government authorities. These bands are rated from *very low frequency* (VLF) to *extremely highfrequency* (EHF). Table lists these bands, their ranges, propagation methods, and some applications

<i>Band</i>	<i>Range</i>	<i>Propagation</i>	<i>Application</i>
VLF (very low frequency)	3–30 kHz	Ground	Long-range radio navigation
LF (low frequency)	30–300 kHz	Ground	Radio beacons and navigational locators
MF (middle frequency)	300 kHz–3 MHz	Sky	AM radio
HF (high frequency)	3–30 MHz	Sky	Citizens band (CB), ship/aircraft communication
VHF (very high frequency)	30–300 MHz	Sky and line-of-sight	VHF TV, FM radio
UHF (ultrahigh frequency)	300 MHz–3 GHz	Line-of-sight	UHF TV, cellular phones, paging, satellite
SHF (superhigh frequency)	3–30 GHz	Line-of-sight	Satellite communication
EHF (extremely high frequency)	30–300 GHz	Line-of-sight	Radar, satellite



Radio Waves

Electromagnetic waves ranging in frequencies between 3 kHz and 1 GHz are normally called radio waves. Radio waves are omnidirectional. When an antenna transmits radio waves, they are propagated in all directions. This means that the sending and receiving antennas do not have to be aligned. A sending antenna sends waves that can be received by any receiving antenna. The omnidirectional property has a disadvantage, too. The radio waves transmitted by one antenna are susceptible to interference by another antenna

that may send signals using the same frequency or band.

Radio waves, particularly those waves that propagate in the sky mode, can travel long distances. This makes radio waves a good candidate for long-distance broadcasting such as AM radio. Radio waves of low and medium frequencies, can penetrate walls. This characteristic can be both an advantage and a disadvantage. It is an advantage because, for example, an AM radio can receive signals inside a building. It is a disadvantage because we cannot isolate a communication to just inside or outside a building.

Omnidirectional Antenna

Radio waves use omnidirectional antennas that send out signals in all directions. Based on the wavelength, strength, and the purpose of transmission, we can have several types of antennas. Figure shows an omnidirectional antenna.

Applications

The omnidirectional characteristics of radio waves make them useful for multicasting, in which there is one sender but many receivers. AM and FM radio, television, maritime radio, cordless phones, and paging are examples of multicasting.

Microwaves

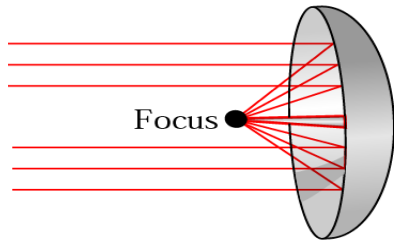
Electromagnetic waves having frequencies between 1 and 300 GHz are called microwaves. Microwaves are unidirectional. When an antenna transmits microwave waves, they can be narrowly focused. This means that the sending and receiving antennas need to be aligned. The unidirectional property has an obvious advantage. A pair of antennas can be aligned without interfering with another pair of aligned antennas. The following describes some characteristics of microwave propagation:

- Microwave propagation is line-of-sight. Since the towers with the mounted antennas need to be in direct sight of each other, towers that are far apart need to be very tall.
- Repeaters are often needed for long-distance communication.

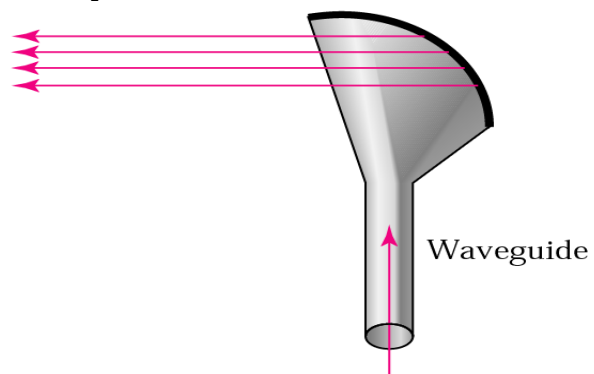
- Very high-frequency microwaves cannot penetrate walls. This characteristic can be a disadvantage if receivers are inside buildings.

Unidirectional Antenna

Microwaves need unidirectional antennas that send out signals in one direction. Two types of antennas are used for microwave communications: the parabolic dish and the horn.



a. Dish antenna



b. Horn antenna

A parabolic dish antenna is based on the geometry of a parabola: Every line parallel to the line of symmetry (line of sight) reflects off the curve at angles such that all the lines intersect in a common point called the focus. The parabolic dish works as a funnel, catching a wide range of waves and directing them to a common point. Outgoing transmissions are broadcast through a horn aimed at the dish. The microwaves hit the dish and are deflected outward in a reversal of the receipt path.

A horn antenna looks like a gigantic scoop. Outgoing transmissions are broadcast up a stem and deflected outward in a series of narrow parallel beams by the curved head. Received transmissions are collected by the scooped shape of the horn, in a manner similar to the parabolic dish, and are deflected down into the stem.

Infrared

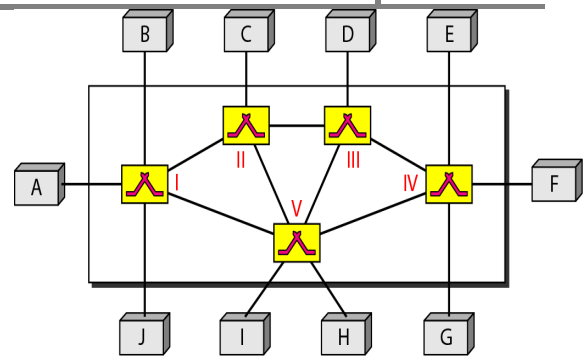
Infrared waves, with frequencies from 300 GHz to 400 THz (wavelengths from 1 mm to 770 nm), can be used for short-range communication. Infrared waves, having high frequencies, cannot penetrate walls. This advantageous characteristic prevents interference between one system and another; a short-range communication system in one room cannot be affected by another system in the next room. When we use our infrared remote control, we do not interfere with the use of the remote by our neighbors. However, this same characteristic makes infrared signals useless for long-range communication. In addition, we cannot use infrared waves outside a building because the sun's rays contain infrared waves that can interfere with the communication.

Applications

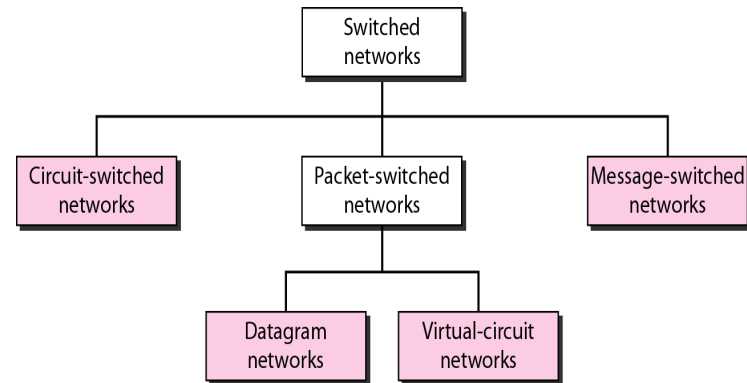
The *Infrared Data Association* (IrDA), an association for sponsoring the use of infrared waves, has established standards for using these signals for communication between devices such as keyboards, mice, PCs, and printers.

Switching

A switched network consists of a series of interlinked nodes, called switches. Switches are devices capable of creating temporary connections between two or more devices linked to the switch. In a switched network, some of these nodes are connected to the end systems (computers or telephones, for example). Others are used only for routing. Figure shows a switched network.



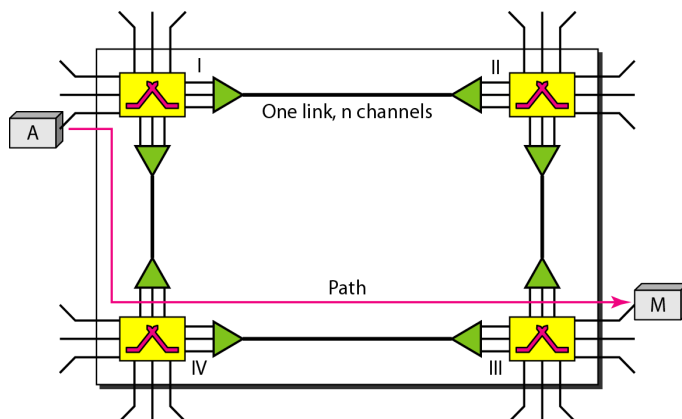
Traditionally, three methods of switching have been important: circuit switching, packet switching, and message switching. Packet-switched networks can further be divided into two subcategories-virtual-circuit networks and datagram networks.



CIRCUIT-SWITCHED NETWORKS

A circuit-switched network consists of a set of switches connected by physical links. A connection between two stations is a dedicated path made of one or more links. However, each connection uses only one dedicated channel on each link. Each link is normally divided into n channels by using FDM or TDM.

Figure below shows a trivial circuit-switched network with four switches and four links. Each link is divided into n (n is 3 in the figure) channels by using FDM or TDM.



When end system A needs to communicate with end system M, system A needs to request a connection to M that must be accepted by all switches as well as by M itself. This is called the setup phase; a circuit (channel) is reserved on each link, and the combination of circuits or channels defines the dedicated path. After the dedicated path made of connected circuits (channels) is established, data transfer can take place. After all data have been transferred, the circuits are torn down.

When end system A needs to communicate with end system M, system A needs to request a connection to M that must be accepted by all switches as well as by M itself. This is called the setup phase; a circuit (channel) is reserved on each link, and the combination of circuits or channels defines the dedicated path. After the dedicated path made of connected circuits (channels) is established, data transfer can take place. After all data have been transferred, the circuits are torn down.

take place. After all data have been transferred, the circuits are torn down.

- Circuit switching takes place at the physical layer.
- Before starting communication, the stations must make a reservation for the resources to be used during the communication. These resources, such as channels, switch buffers, switch processing time, and switch input/output ports, must remain dedicated during the entire duration of data transfer until the teardown phase.
- Data transferred between the two stations are not packetized (physical layer transfer of the signal). The data are a continuous flow sent by the source station and received by the destination station, although there may be periods of silence.
- There is no addressing involved during data transfer. The switches route the data based on their occupied band (FDM) or time slot (TDM)

Three Phases

The actual communication in a circuit-switched network requires three phases: connection setup, data transfer, and connection teardown.

Setup Phase

Before the two parties can communicate, a dedicated circuit needs to be established. The end systems are normally connected through dedicated lines to the switches, so connection setup means creating dedicated channels between the switches. For example, when system A needs to connect to system M, it sends a setup request that includes the address of system M, to switch I. Switch I finds a channel between itself and switch IV that can be dedicated for this purpose. Switch I then sends the request to switch IV, which finds a dedicated channel between itself and switch III. Switch III informs system M of system A's intention at this time. In the next step to making a connection, an acknowledgment from system M needs to be sent in the opposite direction to system A. Only after system A receives this acknowledgment is the connection established.

Data Transfer Phase

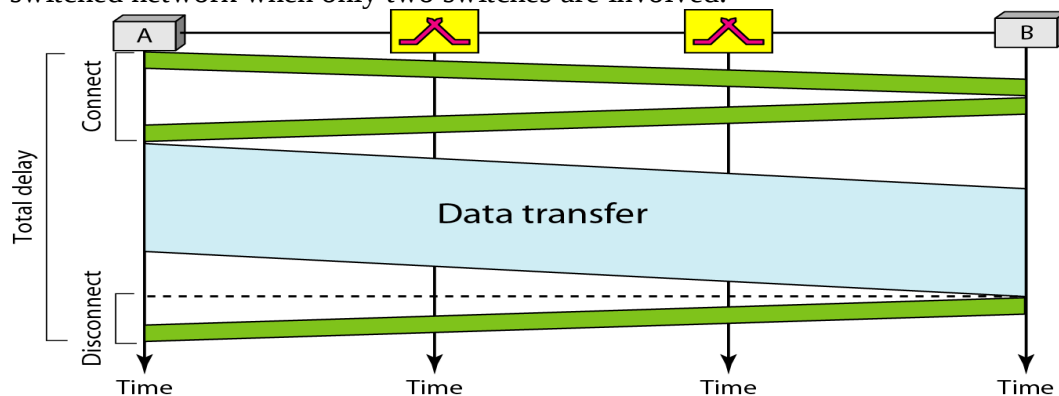
After the establishment of the dedicated circuit (channels), the two parties can transfer data.

Teardown Phase

When one of the parties needs to disconnect, a signal is sent to each switch to release the resources.

Delay

Although a circuit-switched network normally has low efficiency, the delay in this type of network is minimal. During data transfer the data are not delayed at each switch; the resources are allocated for the duration of the connection. Figure shows the idea of delay in a circuit-switched network when only two switches are involved.

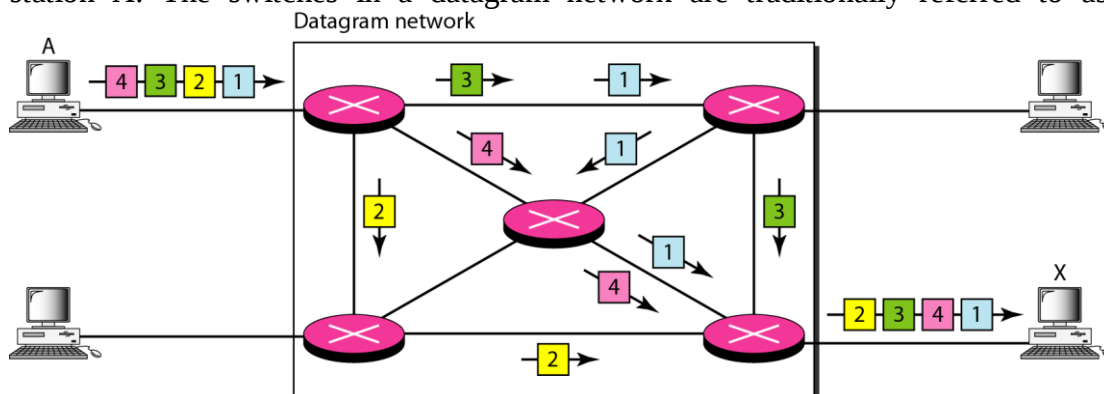


The total delay is due to the time needed to create the connection, transfer data, and disconnect the circuit.

DATAGRAM NETWORKS

If the message is going to pass through a packet-switched network, it needs to be divided into packets of fixed or variable size. The size of the packet is determined by the network and the governing protocol. In packet switching, there is no resource allocation for a packet. This means that there is no reserved bandwidth on the links, and there is no scheduled processing time for each packet. Resources are allocated on demand. The allocation is done on a first-come, first-served basis. In a datagram network, each packet is treated independently of all others. Even if a packet is part of a multipacket transmission, the network treats it as though it existed alone. Packets in this approach are referred to as datagrams. Datagram switching is normally done at the network layer.

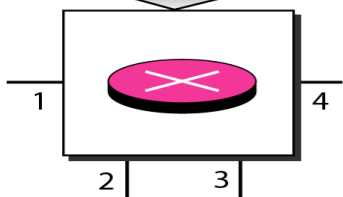
Figure shows how the datagram approach is used to deliver four packets from station A to station X. The switches in a datagram network are traditionally referred to as routers.



In this example, all four packets (or datagrams) belong to the same message, but may travel different paths to reach their destination. This approach can cause the datagrams of a transmission to arrive at their destination out of order with different delays between the packets. Packets may also be lost or dropped because of a lack of resources. In most protocols, it is the responsibility of an upper-layer protocol to reorder the datagrams or ask for lost datagrams before passing them on to the application.

The datagram networks are sometimes referred to as connectionless networks. The term connectionless here means that the switch (packet switch) does not keep information about the connection state. There are no setup or teardown phases. Each packet is treated the same by a switch regardless of its source or destination.

Destination address	Output port
1232	1
4150	2
⋮	⋮
9130	3



Routing Table

In this type of network, each switch (or packet switch) has a routing table which is based on the destination address. The routing tables are dynamic and are updated periodically. The destination addresses and the corresponding forwarding output ports are recorded in the tables. Figure shows the routing table for a switch.

Every packet in a datagram network carries a header that contains, among other information, the destination address of the packet. When the switch receives the packet, this destination address is examined; the routing table is consulted to find the corresponding

port through which the packet should be forwarded. The

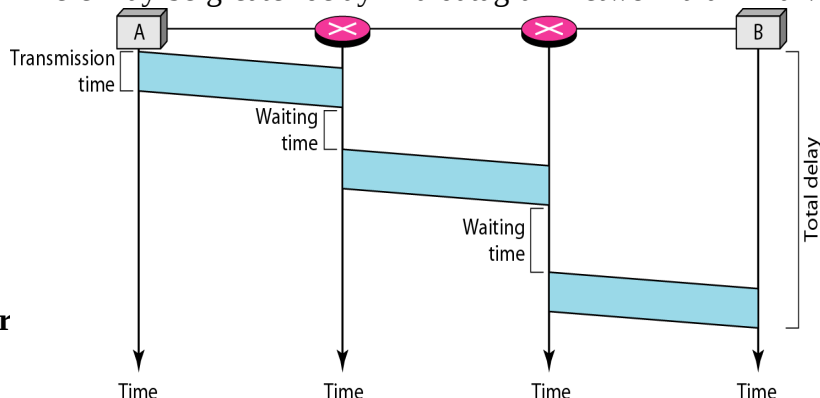
destination address in the header of a packet in a datagram network remains the same during the entire journey of the packet.

Efficiency

The efficiency of a datagram network is better than that of a circuit-switched network; resources are allocated only when there are packets to be transferred

Delay

There may be greater delay in a datagram network than in a virtual-circuit network. Although



there are no setup and teardown phases, each packet may experience a wait at a switch before it is forwarded. In addition, since not all packets in a message

necessarily travel through the same switches, the delay is not uniform for the packets of a message.

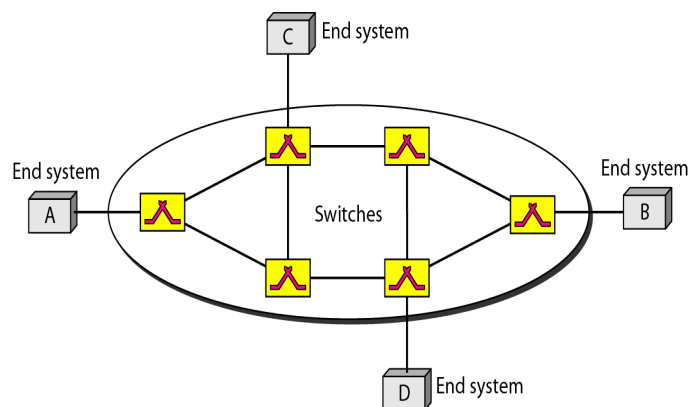
The packet travels through two switches. There are three transmission times ($3T$), three propagation delays (slopes $3t$ of the lines), and two waiting times ($W1 + W2$). The total delay is $\text{Total delay} = 3T + 3t + W1 + W2$

VIRTUAL-CIRCUIT NETWORKS

A virtual-circuit network is a cross between a circuit-switched network and a datagram network. It has some characteristics of both.

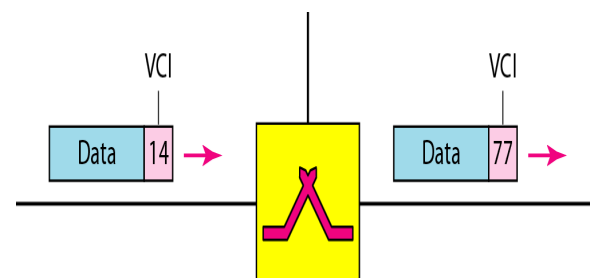
1. As in a circuit-switched network, there are setup and teardown phases in addition to the data transfer phase.
2. Resources can be allocated during the setup phase, as in a circuit-switched network, or on demand, as in a datagram network.
3. As in a datagram network, data are packetized and each packet carries an address in the header. However, the address in the header has local jurisdiction (it defines what should be the next switch and the channel on which the packet is being carried), not end-to-end jurisdiction.
4. As in a circuit-switched network, all packets follow the same path established during the connection.
5. A virtual-circuit network is normally implemented in the data link layer, while a circuit-switched network is implemented in the physical layer and a datagram network in the network layer.

Figure is an example of a virtual-circuit network. The network has switches that allow traffic from sources to destinations. A source or destination can be a computer, packet switch, bridge, or



networks.

any other device that connects other



Addressing

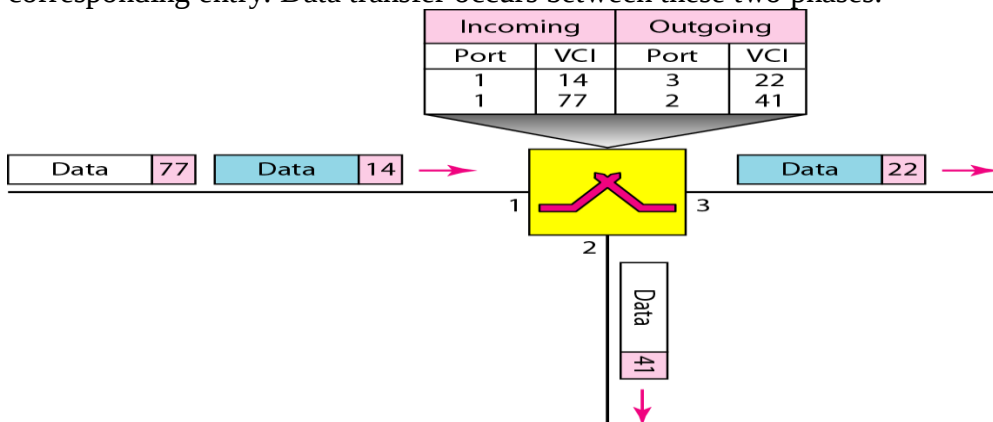
In a virtual-circuit network, two types of addressing are involved: global and local (virtual-circuit identifier-VCI).

- Global address is used only to create a VCI
- Virtual Circuit Identifier is a small number that has only switch scope

- It is used by a frame between two switches.
- When a frame arrives at a switch, it has a VCI; when it leaves, it has a different VCI.

Three Phases

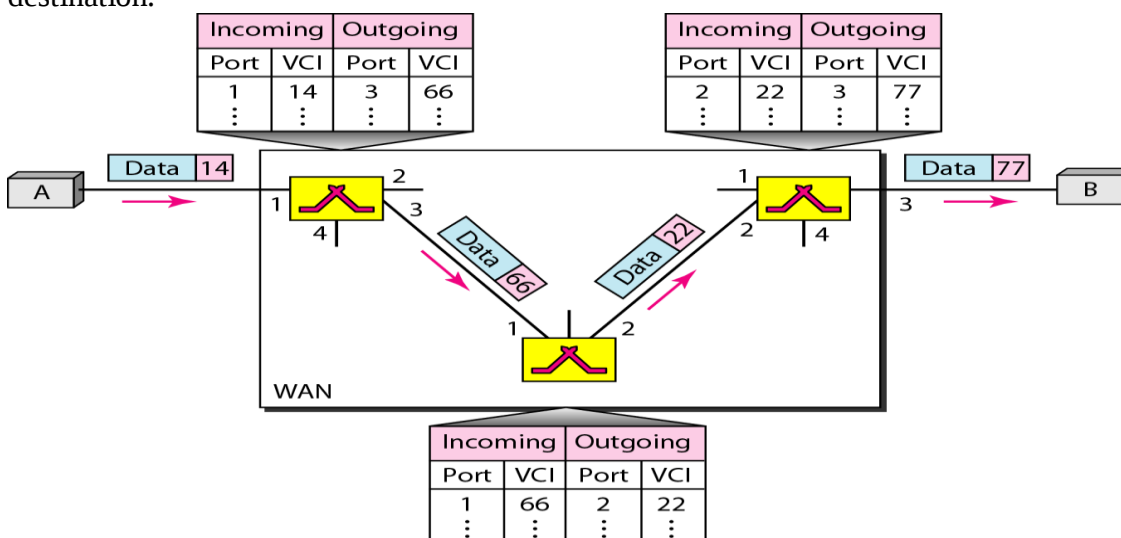
As in a circuit-switched network, a source and destination need to go through three phases in a virtual-circuit network: setup, data transfer, and teardown. In the setup phase, the source and destination use their global addresses to help switches make table entries for the connection. In the teardown phase, the source and destination inform the switches to delete the corresponding entry. Data transfer occurs between these two phases.



Data Transfer Phase

To transfer a frame from a source to its destination, all switches need to have a table entry for this virtual circuit. The table, in its simplest form, has four columns. This means that the switch holds four pieces of information for each virtual circuit that is already set up. Figure 8.12 shows such a switch and its corresponding table.

Figure above shows a frame arriving at port 1 with a VCI of 14. When the frame arrives, the switch looks in its table to find port 1 and a VCI of 14. When it is found, the switch knows to change the VCI to 22 and send out the frame from port 3. Figure below shows how a frame from source A reaches destination B and how its VCI changes during the trip. Each switch changes the VCI and routes the frame. The data transfer phase is active until the source sends all its frames to the destination. The procedure at the switch is the same for each frame of a message. The process creates a virtual circuit, not a real circuit, between the source and destination.

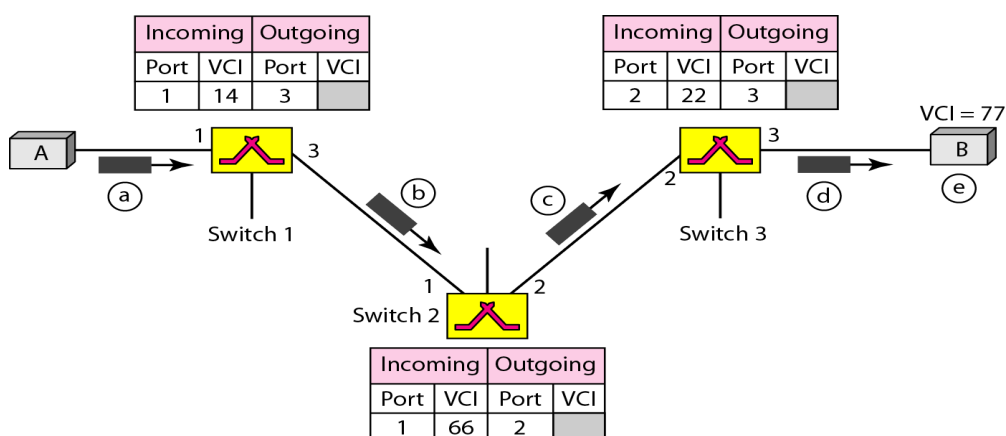


Setup Phase

In the setup phase, a switch creates an entry for a virtual circuit. Two steps are required: the setup request and the acknowledgment.

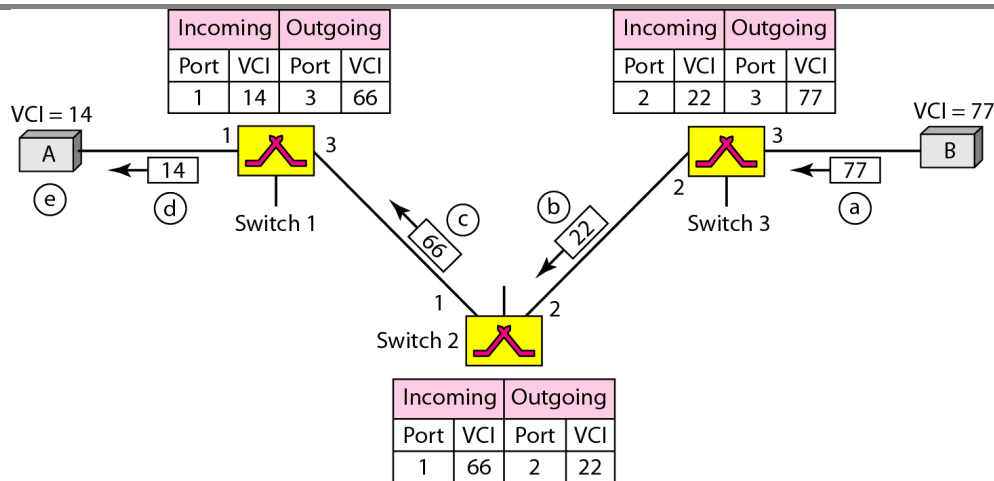
Setup Request: A setup request frame is sent from the source to the destination. Figure shows the process.

- Source A sends a setup frame to switch 1.
- Switch 1 receives the setup request frame. The switch assigns the incoming port (1) and chooses an available incoming VCI (14) and the outgoing port (3). It does not yet know the outgoing VCI, which will be found during the acknowledgment step. The switch then forwards the frame through port 3 to switch 2.
- Switch 2 receives the setup request frame. The same events happen here as at switch 1; three columns of the table are completed: in this case, incoming port (1), incoming VCI (66), and outgoing port (2).
- Switch 3 receives the setup request frame. Again, three columns are completed: incoming port (2), incoming VCI (22), and outgoing port (3).
- Destination B receives the setup frame, and if it is ready to receive frames from A, it assigns a VCI to the incoming frames that come from A, in this case 77. This VCI lets the destination know that the frames come from A, and not other sources.



Acknowledgment: A special frame, called the acknowledgment frame, completes the entries in the switching tables. Figure 8.15 shows the process.

- The destination sends an acknowledgment to switch 3. The acknowledgment carries the global source and destination addresses so the switch knows which entry in the table is to be completed. The frame also carries VCI 77, chosen by the destination as the incoming VCI for frames from A. Switch 3 uses this VCI to complete the outgoing VCI column for this entry. Note that 77 is the incoming VCI for destination B, but the outgoing VCI for switch 3.
- Switch 3 sends an acknowledgment to switch 2 that contains its incoming VCI in the table, chosen in the previous step. Switch 2 uses this as the outgoing VCI in the table.
- Switch 2 sends an acknowledgment to switch 1 that contains its incoming VCI in the table, chosen in the previous step. Switch 1 uses this as the outgoing VCI in the table.
- Finally switch 1 sends an acknowledgment to source A that contains its incoming VCI in the table, chosen in the previous step.
- The source uses this as the outgoing VCI for the data frames to be sent to destination B.



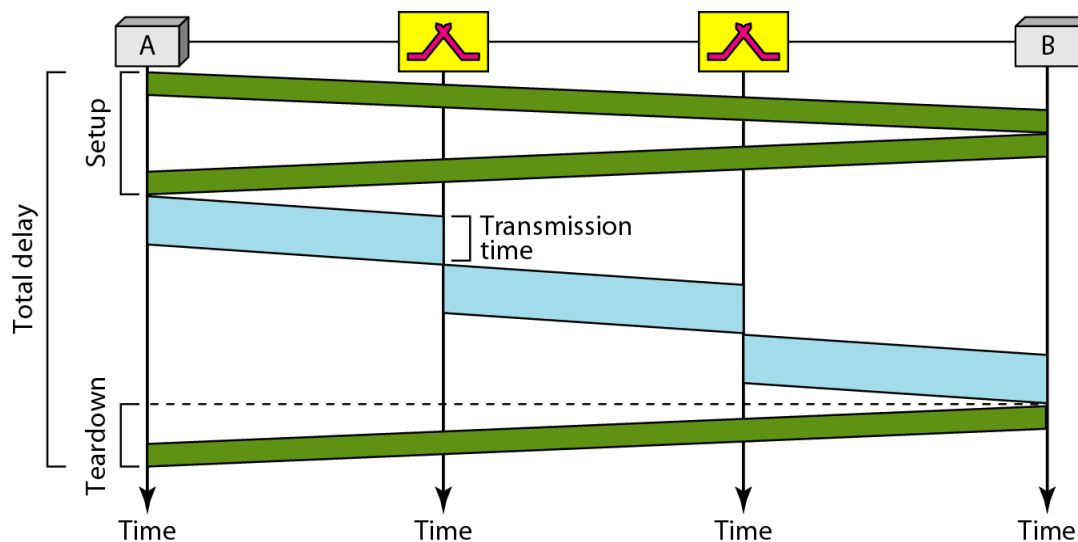
Teardown Phase

In this phase, source A, after sending all frames to B, sends a special frame called a *teardown request*. Destination B responds with a *teardown confirmation* frame. All switches delete the corresponding entry from their tables.

Delay in Virtual-Circuit Networks

In a virtual-circuit network, there is a one-time delay for setup and a one-time delay for teardown. If resources are allocated during the setup phase, there is no wait time for individual packets. In virtual-circuit switching, all packets belonging to the same source and destination travel the same path; but the packets may arrive at the destination with different delays if resource allocation is on demand.

The total delay time is $= 3T + 3t + \text{setup delay} + \text{teardown delay}$



POSSIBLE QUESTIONS**5X8=40 Marks**

1. Elucidate the principle of frequency division multiplexing with a neat diagram.
2. Explain the principle of multiplexing for light signals with a neat diagram.
3. Illustrate the multiplexing of a digital signal with a neat diagram.
4. Discuss in detail about virtual- circuit networks with neat sketches.
5. Elucidate the types of unguided media with a neat sketch.
6. Draw the construction of Twisted-pair cables and Co-Axial Cables and explain their working.
7. Explain the principle behind the fiber-optic cable.
8. What is a datagram network? Explain its working principle.
9. Compare and contrast Circuit switched networks and virtual- circuit networks
10. Discuss the advantages of statistical time division multiplexing over synchronous time division multiplexing.



KARPAGAM ACADEMY OF HIGHER EDUCATION
(Deemed to be University Established Under Section 3 of UGC Act 195
Pollachi Main Road, Eachanari Post, Coimbatore – 641 021. INDIA
DEPARTMENT OF COMPUTER SCIENCE
DATA COMMUNICATION NETWORKS (15CSU502)
ONLINE EXAM QUESTION BANK

UNIT-2

Questions	opt1	opt2	opt3	opt4
Before data can be transmitted, they must be transformed to _____	periodic signals	electromagnetic signals	Aperiodic signals	low frequency sine waves
Which of the following can be determined from a frequency_domain graph of a signal?	frequency	phase	power	amplitude
Which of the following can be determined from a frequency_domain graph of a signal?	bandwidth	phase	power	frequency
In a frequency_domain plot, the vertical axis measures the _____	peak amplitude	frequency	phase	slope
In a frequency_domain plot, the horizontal axis measures the _____	peak amplitude	frequency	phase	slope
As frequency increases, the period _____	decreases	increases	remains the same	doubles
A sine wave is _____	periodic and continuous	aperiodic and continuous	periodic and discrete	aperiodic and discrete
_____ is a type of transmission impairment in which the signal loses strength due to the resistance of the transmission medium	attenuation	distortion	noise	decibel
_____ is a type of transmission impairment in which the signal loses strength due to different propagation speeds of each frequency that makes up the signal	attenuation	distortion	noise	decibel

_____ is a type of transmission impairment in which an outside source such as crosstalk corrupts a signal	attenuation	distortion	noise	decibel
Propagation time is _____ proportional to distance and _____ proportional to propagation speed	inversely; directly	directly; inversely	inversely; inversely	directly; directly
The wavelength of a signal depend on the _____	frequency of the signal	medium	phase of signal	Frequency and Phase of signal
Unipolar, bipolar and polar encoding are types of _____ encoding	line	block	NRZ	manchester
If a symbol is composed of 3bits ther are _____ data levels	2	4	8	16
_____ encoding has a transition at the middle of each bit	RZ	manchester	differential manchester	all the above
_____ encoding has a transition at the begining of each 0 bit	RZ	manchester	differential manchester	all the above
PCM is an example of _____ conversion	digital-to-digital	digital-to-analog	analog-to-analog	analog-to-digital
_____ transmission, bits are transmitted simultaneously, each across the own wire	asynchronous serial	synchronous serial	parallel	perpendicular
In _____ transmission, a start bit and a stop bit frame a character byte	asynchronous serial	synchronous serial	parallel	fixed
In asynchronous transmission, the gap tim between bytes is _____	fixed	variable	a function of the data rate	zero
ASK, PSK, FSK and QAM are examples of _____ modulation	digital-to-digital	digital-to-analog	analog-to-analog	analog-to-digital
AM and FM are examples of modulation	digital-to-digital	digital-to-analog	analog-to-analog	analog-to-digital
In QAM, both phase and _____ of a carrier frequency are varied	amplitude	frequency	bit rate	baud rate

Which multiplexing technique transmits analog signals	FDM	TDM	WDM	FDM and WDM
Which multiplexing technique transmits digital signals	FDM	TDM	WDM	FDM and WDM
Which multiplexing technique shifts each signal to a different carrier frequency?	FDM	TDM	WDM	FDM and WDM
In TDM, for n signal sources of the same data rate, each frame contains _____ slots	n	n+1	n-1	0 to n
Guard bands increases the bandwidth for _____	FDM	TDM	WDM	TDM and WDM
Which multiplexing technique involves signals composed of light beams?	FDM	TDM	WDM	FDM and TDM
Transmission media are usually categorized as _____	fixed or unfixed	guided or unguided	determinate or indeterminate	metallic or non-metallic
Transmission media are usually categorized as _____	physical	network	transport	application
Category 1 UTP cable is most often used in _____ networks	fast ethernet	traditional ethernet	infrared	telephone
BNC connectors are used by _____ cables	UTP	STP	coaxial	fiber-optic
In fiber optics, the signal source is _____ waves	light	radio	infrared	very low frequency
A parabolic dish antenna is a(n) _____ antenna	omni directional	bi directional	uni directional	horn
A telephone network is an example of a _____ network	packet switching	circuit switched	message switched	packet and message switched
Periodic analog signals can be classified into _____	simple	composite	simple or composite	simple and composite
Period and frequency has the following formula.	$f=1/t$ and $t=1/f$	$t=1/f$ or $f=1/t$	$c=t/f$	$t=c/f$

Wavelength is _____	propagation speed	propagation speed * frequency	propagation speed/period	propagation speed/frequency
Composite signal can be classified into _____ types	five	three	four	two
The range of frequency contained in a _____ signal is its bandwidth.	simple	composite	periodic	non periodic
The bandwidth of the composite signal is the difference between the _____	highest	highest or lowest	highest and lowest	lowest
The _____ is the number of bits sent in a second.	bit length	bandpass	bandwidth	bit rate
Bit length is _____	propagation speed/period	propagation speed * frequency	bit	propagation speed*bit duration
A _____ signal is a composite analog signal with an infinite bandwidth	simple	composite	digital	analog
Bit rate=_____	$4 \cdot BW \cdot \log_2 L$	$2 \cdot BW \cdot \log_2 L$	$4 \cdot BW / L$	$2 \cdot BW \cdot \log_2 4L$
The term _____ data refers to information continuous	analog	digital	physical	analog and digital
The sine wave is the most fundamental form of a _____ analog signal.	composite	single	periodic	non periodic

6)

answer
electromagnetic signals
frequency
bandwidth
peak amplitude
frequency
increases
periodic and continuous
attenuation
distortion

noise
directly; inversely
Frequency and Phase of signal
line
8
manchester
RZ
analog-to-digital
parallel
synchronous serial
fixed
digital-to-analog
analog-to-analog
amplitude

FDM and WDM
TDM
FDM
n
FDM
WDM
guided or unguided
physical
telephone
coaxial
light
omni directional
circuit switched
simple or composite
$f=1/t$ and $t=1/f$

propagation speed/frequency

two

composite

highest and lowest

bit rate

propagation speed*bit duration

digital

$2 \cdot BW \cdot \log_2 L$

analog

periodic

UNIT-III-SYLLABUS

Data Link Layer: Error Correction and Detection – Framing – Flow and Error Control – Protocols – Noisy and Noiseless Channel – Multiple Access.

UNIT-III DATA LINK LAYER

ERROR DETECTION & CORRECTION

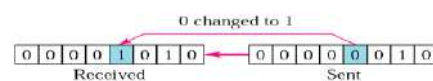
Networks must be able to transfer data from one device to another with acceptable accuracy. For most applications, a system must guarantee that the data received are identical to the data transmitted. Data can be corrupted during transmission. Some applications require that errors be detected and corrected.

Types of Errors

Whenever bits flow from one point to another, they are subject to unpredictable changes because of interference. This interference can change the shape of the signal. In a single-bit error, a 0 is changed to a 1 or a 1 to 0. In a burst error, multiple bits are changed.

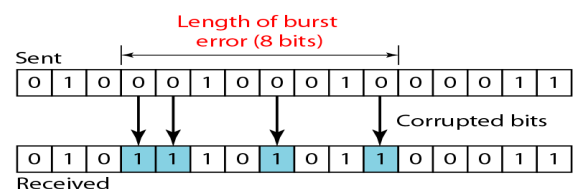
Single-bit error

The term single-bit error means that only 1 bit of a given data unit (such as a byte, character, or packet) is changed from 1 to 0 or from 0 to 1. Single-bit errors are the least likely type of error in serial data transmission.



Burst Error

The term burst error means that 2 or more bits in the data unit have changed from 1 to 0 or from 0 to 1. A burst error is more likely to occur than a single-bit error. The duration of noise is normally longer than the duration of 1 bit, which means that when noise affects data, it affects a set of bits. The number of bits affected depends on the data rate and duration of noise.



Redundancy

The central concept in detecting or correcting errors is redundancy. To be able to detect or correct errors, we need to send some extra bits with our data. These redundant bits are added by the sender and removed by the receiver. Their presence allows the receiver to detect or correct corrupted bits.

Detection Versus Correction

The correction of errors is more difficult than the detection. In error detection, we are looking only to see if any error has occurred. The answer is a simple yes or no. We are not even interested in the number of errors. A single-bit error is the same for us as a burst error.

In error correction, we need to know the exact number of bits that are corrupted and more importantly, their location in the message. The number of the errors and the size of the message are important factors.

Forward Error Correction Versus Retransmission

There are two main methods of error correction. Forward error correction is the process in which the receiver tries to guess the message by using redundant bits. This is possible, as we see later, if the number of errors is small. Correction by retransmission is a technique in which the receiver detects the occurrence of an error and asks the sender to resend the message. Resending is repeated until a message arrives that the receiver believes is error-free (usually, not all errors can be detected).

BLOCK CODING

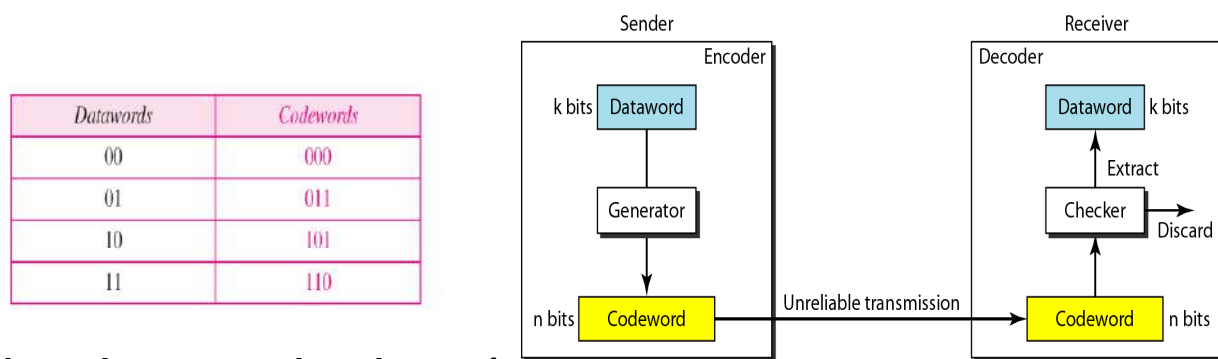
In block coding, we divide our message into blocks, each of k bits, called datawords. Add r redundant bits to each block to make the length $n = k + r$. The resulting n -bit blocks are called codewords.

Error Detection in block coding

If the following two conditions are met, the receiver can detect a change in the original codeword.

1. The receiver has (or can find) a list of valid codewords.

2. The original codeword has changed to an invalid one.



The sender creates codewords out of datawords by using a generator that applies the rules and procedures of encoding. Each codeword sent to the receiver may change during transmission. If the received codeword is the same as one of the valid code-words, the word is accepted; the corresponding dataword is extracted for use. If the received codeword is not valid, it is discarded. However, if the codeword is corrupted during transmission but the received word still matches a valid codeword, the error remains undetected. This type of coding can detect only single errors. Two or more errors may remain undetected.

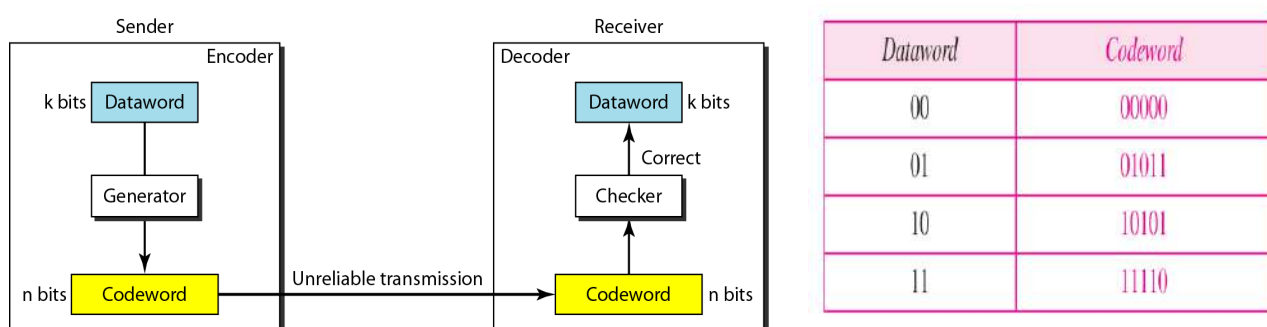
Example: Assume the sender encodes the dataword 01 as 011 and sends it to the receiver. Consider the following cases:

1. The receiver receives 011 which is a valid codeword. The receiver extracts the dataword 01 from it.
2. The codeword is corrupted during transmission, and 111 is received. This is not a valid codeword and is discarded.
3. The codeword is corrupted during transmission, and 000 is received. This is a valid codeword. The receiver incorrectly extracts the dataword 00. Two corrupted bits have made the error undetectable.

→ An error-detecting code can detect only the types of errors for which it is designed; other types of errors may remain undetected

Error Correction Using Block Coding

In error detection, the receiver needs to know only that the received codeword is invalid; in error correction the receiver needs to find (or guess) the original codeword sent. More redundant bits are



needed for error correction than for error detection.

Assume the dataword is 01. The sender creates the codeword 01011. The codeword is corrupted during transmission, and 01001 is received. First, the receiver finds that the received codeword is not in the table. This means an error has occurred. The receiver, assuming that there is only 1 bit corrupted, uses the following strategy to guess the correct dataword

1. Comparing the received codeword with the first codeword in the table (01001 versus 00000), the receiver decides that the first codeword is not the one that was sent because there are two different bits. (the same for third or fourth one in the table)
- 2.. The original codeword must be the second one in the table because this is the only one that differs from the received codeword by 1 bit.

Hamming Distance

The Hamming distance between two words is the number of differences between corresponding bits. Example: Hamming distance $d(10101, 11110)$ is 3

The minimum Hamming distance is the smallest Hamming distance between all possible pairs in a set of words

$$d_{\min} = 2$$

Hamming distance and error

When a codeword is corrupted during transmission, the Hamming distance between the sent and received code-words is the number of bits affected by the error. In other words, the Hamming distance between the received codeword and the sent codeword is the number of bits that are corrupted during transmission. For example, if the codeword 00000 is sent and 01101 is received, 3 bits are in error and the Hamming distance between the two is $d(00000, 01101) = 3$.

Minimum Distance for Error Detection

If s errors occur during transmission, the Hamming distance between the sent codeword and received codeword is s . If our code is to detect up to s errors, the minimum distance between the valid codes must be $s + 1$, so that the received codeword does not match a valid codeword. To guarantee the detection of up to s errors in all cases, the minimum Hamming distance in a block code must be $d_{\min} = s + 1$.

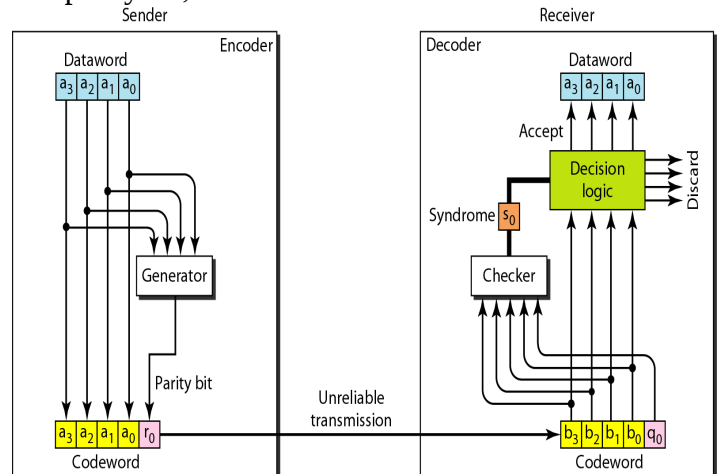
LINEAR BLOCK CODES

- Almost all block codes used today belong to a subset called linear block codes.
- A linear block code is a code in which the exclusive OR (addition modulo-2) of two valid codewords creates another valid codeword
- The minimum Hamming distance is the number of 1s in the nonzero valid codeword with the smallest number of 1s

Simple Parity-Check Code

A simple parity-check code is a single-bit error-detecting code in which $n = k + 1$ with $d_{\min} = 2$. In this code, a k -bit dataword is changed to an n -bit codeword where $n = k + 1$. The extra bit, called the parity bit, is selected to make the total number

Datawords	Codewords	Datawords	Codewords
0000	00000	1000	10001
0001	00011	1001	10010
0010	00101	1010	10100
0011	00110	1011	10111
0100	01001	1100	11000
0101	01010	1101	11011
0110	01100	1110	11101
0111	01111	1111	11110



of 1s in the codeword even.

The sender sends the codeword which may be corrupted during transmission. The receiver receives a 5-bit word. The checker at the receiver does the same thing as the generator in the sender with one

exception: The addition is done over all 5 bits. The result, which is called the syndrome, is just 1 bit. The syndrome is 0 when the number of 1s in the received codeword is even; otherwise, it is 1. The syndrome is passed to the decision logic analyzer. If the syndrome is 0, there is no error in the received codeword; the data portion of the received codeword is accepted as the dataword; if the syndrome is 1, the data portion of the received codeword is discarded. The dataword is not created.

Example:

Let us look at some transmission scenarios. Assume the sender sends the dataword 1011. The codeword created from this dataword is 10111, which is sent to the receiver. We examine five cases:

1. No error occurs; the received codeword is 10111. The syndrome is 0. The dataword 1011 is created.
2. One single-bit error changes a_1 . The received codeword is 10011. The syndrome is 1. No dataword is created.
3. One single-bit error changes r_0 . The received codeword is 10110. The syndrome is 1. No dataword is created. Note that although none of the dataword bits are corrupted, no dataword is created because the code is not sophisticated enough to show the position of the corrupted bit.
4. An error changes r_0 and a second error changes a_3 . The received codeword is 00110. The syndrome is 0. The dataword 0011 is created at the receiver. Note that here the dataword is wrongly created due to the syndrome value. The simple parity-check decoder cannot detect an even number of errors. The errors cancel each other out and give the syndrome a value of 0.
5. Three bits- a_3 , a_2 , and a_1 -are changed by errors. The received codeword is 01011. The syndrome is 1. The dataword is not created.

This shows that the simple parity check, guaranteed to detect one single error, can also find any odd number of errors.

A better approach is the two-dimensional parity check. In this method, the dataword is organized in a table (rows and columns). The data to be sent, five 7-bit bytes, are put in separate rows. For each row and each column, 1 parity-check bit is calculated. The whole table is then sent to the receiver, which finds the syndrome for each row and each column. the two-dimensional parity check can detect up to three errors that occur anywhere in the table (arrows point to the locations of the created nonzero syndromes). However, errors affecting 4 bits may not be detected.

1	1	0	0	1	1	1	1	Row parities
1	0	1	1	1	0	1	1	
0	1	1	1	0	0	1	0	
0	1	0	1	0	0	1	1	
0	1	0	1	0	1	0	1	
Column parities								

a. Design of row and column parities

1	1	0	0	1	1	1	1	Row parities
1	0	1	1	1	0	1	1	
0	1	1	1	0	0	1	0	
0	1	0	1	0	0	1	1	
0	1	0	1	0	1	0	1	
Column parities								

b. One error affects two parities

1	1	0	0	1	1	1	1	Row parities
1	0	1	1	1	0	1	1	
0	1	1	1	0	0	1	0	
0	1	0	1	0	0	1	1	
0	1	0	1	0	1	0	1	
Column parities								

c. Two errors affect two parities

1	1	0	0	1	1	1	1	Row parities
1	0	1	1	1	0	1	1	
0	1	1	1	0	0	1	0	
0	1	0	1	0	0	1	1	
0	1	0	1	0	1	0	1	
Column parities								

d. Three errors affect four parities

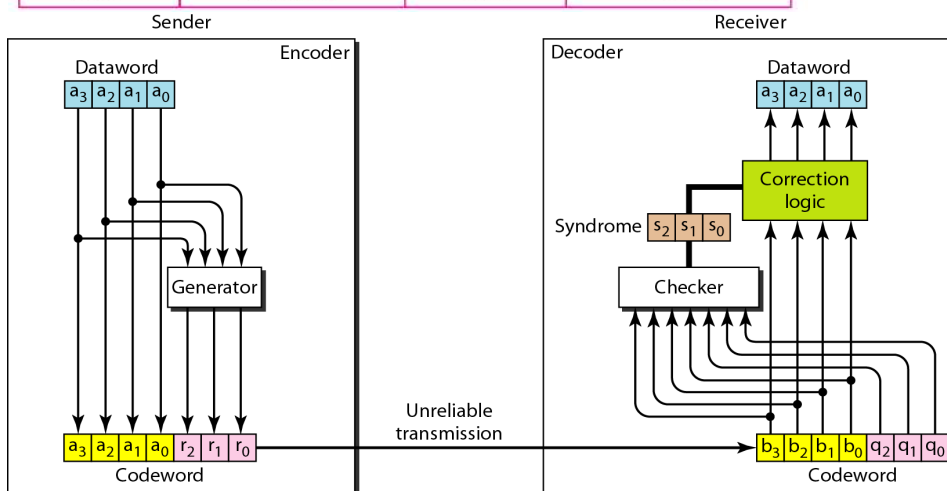
1	1	0	0	1	1	1	1	Row parities
1	0	1	1	1	0	1	1	
0	1	1	1	0	0	1	0	
0	1	0	1	0	0	1	1	
0	1	0	1	0	1	0	1	
Column parities								

e. Four errors cannot be detected

Hamming Codes

Hamming codes were originally designed with $d_{\min} = 3$, which means that they can detect up to two errors or correct one single error. The relationship between m and n in these codes is $n = 2^m - 1$

Datawords	Codewords	Datawords	Codewords
0000	0000000	1000	1000110
0001	0001101	1001	1001011
0010	0010111	1010	1010001
0011	0011010	1011	1011100
0100	0100011	1100	1100101
0101	0101110	1101	1101000
0110	0110100	1110	1110010
0111	0111001	1111	1111111



- $r_0 = a_2 + a_1 + a_0$ modulo-2
- $r_1 = a_3 + a_2 + a_1$ modulo-2
- $r_2 = a_1 + a_0 + a_3$ modulo-2
- $s_0 = b_2 + b_1 + b_0 + q_0$ modulo-2
- $s_1 = b_3 + b_2 + b_1 + q_1$ modulo-2
- $s_2 = b_1 + b_0 + b_3 + q_2$ modulo-2

A Hamming code can only correct a single error or detect a double error.

Logical Decision by Decoder is

Syndrome	000	001	010	011	100	101	110	111
Error	None	q_0	q_1	b_2	q_2	b_0	b_3	b_1

Let us trace the path of three datawords from the sender to the destination:

- 1 The dataword 0100 becomes the codeword 0100011. The codeword 0100011 is received. The syndrome is 000, the final dataword is 0100.
- 2 The dataword 0111 becomes the codeword 0111001. The codeword 0011001 received. The syndrome is 011. After flipping b_2 (changing the 1 to 0), the final dataword is 0111.
- 3 The dataword 1101 becomes the codeword 1101000. The codeword 0001000 received (two errors). The syndrome is 101. After flipping b_0 , we get 0000, the wrong dataword. This shows that our code cannot correct two errors.

CYCLIC CODES

Cyclic codes are special linear block codes with one extra property. In a cyclic code, if a codeword is cyclically shifted (rotated), the result is another codeword.

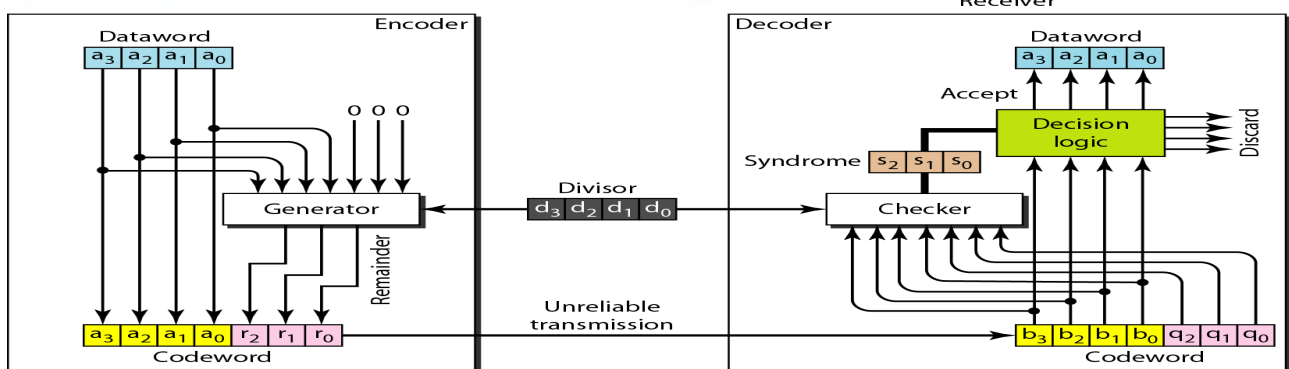
Cyclic Redundancy Check

Cyclic codes called the cyclic redundancy check (CRC) is used in networks such as LANs and

WANs. Table shows an example of a CRC code. We can see both the linear and cyclic properties of this code.

In the encoder, the dataword has k bits (4 here); the codeword has n bits (7 here). The size of the dataword is augmented by adding $n - k$ (3 here) 0s to the right-hand side of the word. The n -bit result is fed into the generator. The generator uses a divisor of size $n - k + 1$ (4 here), predefined and agreed upon. The generator divides the augmented dataword by the divisor (modulo-2 division). The quotient of the division is discarded; the remainder (r₂r₁r₀) is appended to the dataword to create the codeword.

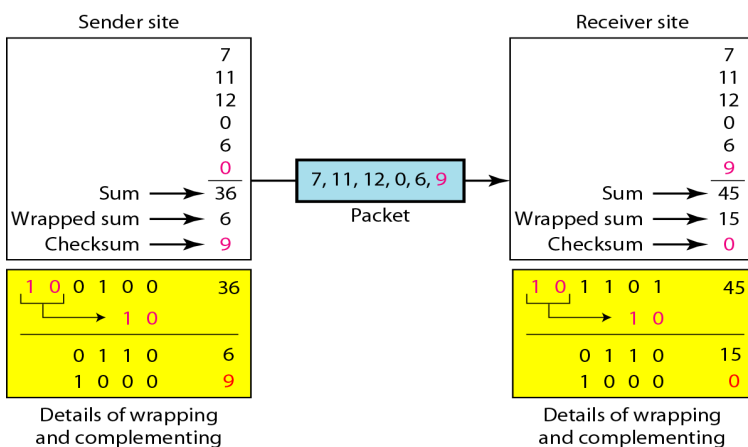
Dataword	Codeword	Dataword	Codeword
0000	0000000	1000	1000101
0001	0001011	1001	1001110
0010	0010110	1010	1010011
0011	0011101	1011	1011000
0100	0100111	1100	1100010
0101	0101100	1101	1101001
0110	0110001	1110	1110100
0111	0111010	1111	1111111



The decoder receives the possibly corrupted codeword. A copy of all n bits is fed to the checker which is a replica of the generator. The remainder produced by the checker is a syndrome of $n - k$ (3 here) bits, which is fed to the decision logic analyzer. The analyzer has a simple function. If the syndrome bits are all as, the 4 leftmost bits of the codeword are accepted as the dataword (interpreted as no error); otherwise, the 4 bits are discarded (error).

CHECKSUM

The last error detection method we discuss here is called the checksum. The checksum is used in the Internet by several protocols although not at the data link layer.



- Tendency is to replace the checksum with a CRC
- Not as strong as CRC in error-checking capability
- One's complement arithmetic
 - We can represent unsigned numbers between 0 and $2^n - 1$ using only n bits
 - If the number has more than n bits, the extra leftmost bits need to be added to the n rightmost bits (wrapping)

- A negative number can be represented by inverting all bits. It is the same as subtracting the number from $2^n - 1$

Example: The sender initializes the checksum to 0 and adds all data items and the checksum. However, 36 cannot be expressed in 4 bits. The extra two bits are wrapped and added with the sum to create the wrapped sum value 6. The sum is then complemented, resulting in the checksum value

9 (15 - 6 = 9).

FRAMING

The data link layer needs to pack bits into frames, so that each frame is distinguishable from another. Framing in the data link layer separates a message from one source to a destination, or from other messages to other destinations, by adding a sender address and a destination address. Although the whole message could be packed in one frame, that is not normally done. One reason is that a frame can be very large, making flow and error control very inefficient. When a message is carried in one very large frame, even a single-bit error would require the retransmission of the whole message. When a message is divided into smaller frames, a single-bit error affects only that small frame.

Fixed-Size Framing

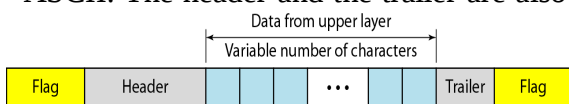
Frames can be of fixed or variable size. In fixed-size framing, there is no need for defining the boundaries of the frames; the size itself can be used as a delimiter.

Variable-Size Framing

In variable-size framing, we need a way to define the end of the frame and the beginning of the next. Two approaches were used for this purpose: a character-oriented approach and a bit-oriented approach.

Character-Oriented Protocols

In a character-oriented protocol, data to be carried are 8-bit characters from a coding system such as ASCII. The header and the trailer are also multiples of 8 bits. To separate one frame from the next, an 8-bit (1-byte) flag is added at the beginning and the end of a frame. The flag, composed of protocol-dependent special characters, signals the start or end of

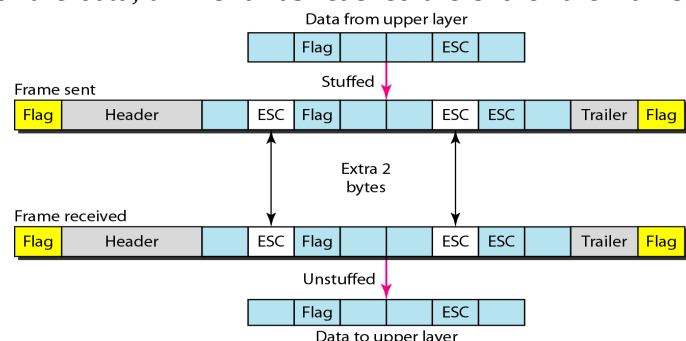


a frame.

Any pattern used for the flag could also be part of the information. If this happens, the receiver, when it encounters this pattern in the middle of the data, thinks it has reached the end of the frame.

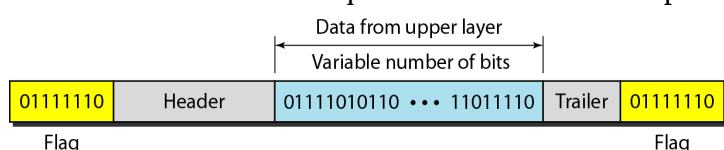
To fix this problem, a byte-stuffing strategy was added to character-oriented framing. In byte stuffing (or character stuffing), a special byte is added to the data section of the frame when there is a character with the same pattern as the flag. The data section is stuffed with an extra byte. This byte is usually called the escape character (ESC), which has a predefined bit pattern. Whenever the receiver encounters the ESC character, it removes it from the data section and treats the next character as data, not a delimiting flag.

If the text contains one or more escape characters followed by a flag, the receiver removes the escape character, but keeps the flag, which is incorrectly interpreted as the end of the frame. To solve this problem, the escape characters that are part of the text must also be marked by another escape character. In other words, if the escape character is part of the text, an extra one is added to show that the second one is part of the text.

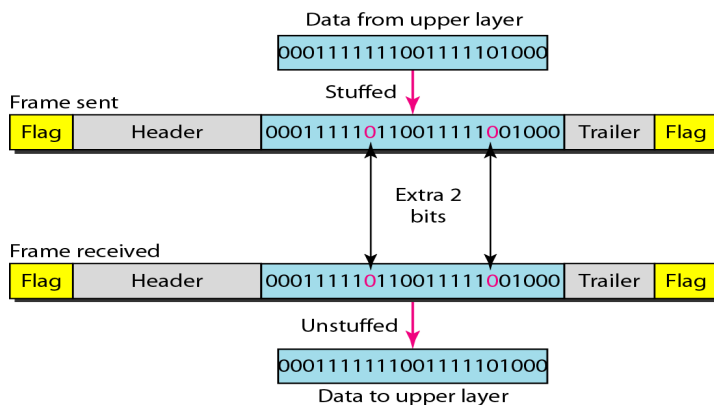


Bit-Oriented Protocols

In a bit-oriented protocol, the data section of a frame is a sequence of bits to be interpreted by the upper layer as text, graphic, audio, video, and so on. Most protocols use a special 8-bit pattern flag 01111110 as the delimiter to define the beginning and the end of the frame.



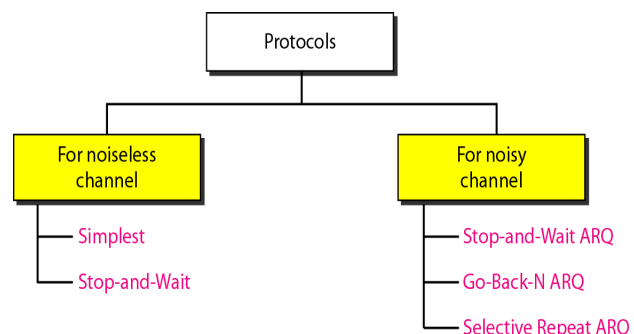
This flag can create the same type of problem the byte-oriented protocols. That is, if the flag pattern appears in the data, we need to somehow inform the receiver that this is not the end of the frame. This is done by stuffing 1 single bit (instead of 1 byte) to prevent the pattern from looking like a flag. The strategy is called bit stuffing. In bit stuffing, if a 0 and five consecutive 1 bits are encountered, an extra 0 is added. This extra stuffed bit is eventually removed from the data by the receiver. Note that the extra bit is added after one 0 followed by five 1s regardless of the value of the next bit. This



guarantees that the flag field sequence does not inadvertently appear in the frame.

Flow and Error Control

- *Data link control* = *flow control* + *error control*
- *Flow control* refers to a set of procedures used to restrict the amount of data that the sender can send before waiting for acknowledgement
- *Error control* in the data link layer is based on *automatic repeat request* (ARQ), which is the retransmission of data



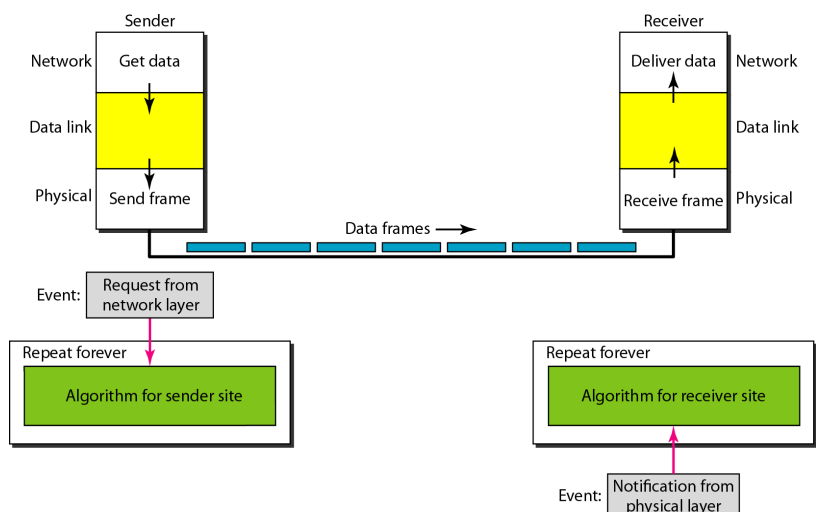
All the protocols discussed here are unidirectional in the sense that the data frames travel from one node, called the sender, to another node, called the receiver. Although special frames, called acknowledgment (ACK) and negative acknowledgment (NAK) can flow in the opposite direction for flow and error control purposes, data flow in only one direction. In a real-life network, the data link protocols are implemented as bidirectional; data flow in both directions. In these protocols the flow and error control information such as ACKs and NAKs is included in the data frames in a technique called piggybacking. Bidirectional protocols are more complex than unidirectional ones.

NOISELESS CHANNELS

A Noiseless channel is an ideal channel in which no frames are lost, duplicated, or corrupted. There are two protocols for this type of channel. The first is a protocol that does not use flow control; the second is the one that does.

Simplest Protocol

The Simplest Protocol has no flow or error control. It is a unidirectional protocol in which data frames are traveling in only one direction—from the sender to receiver. We assume that the receiver can immediately handle any frame it receives with a processing time that is small enough to be negligible. The data link layer of the receiver immediately removes the header



from the frame and hands the data packet to its network layer, which can also accept the packet immediately. In other words, the receiver can never be overwhelmed with incoming frames.

Design

There is no need for flow control in this scheme. The data link layer at the sender site gets data from its network layer, makes a frame out of the data, and sends it. The data link layer at the receiver site receives a frame from its physical layer, extracts data from the frame, and delivers the data to its network layer. The data link layers of the sender and receiver provide transmission services for their network layers. The data link layers use the services provided by their physical layers (such as signaling, multiplexing, and so on) for the physical transmission of bits.

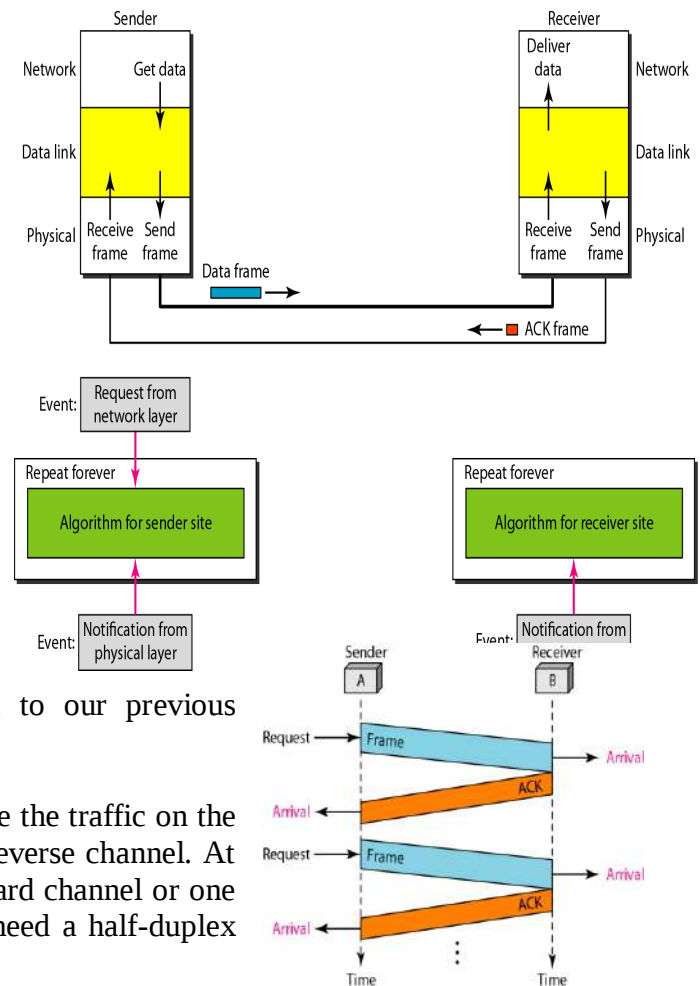
Stop-and-Wait Protocol

If data frames arrive at the receiver site faster than they can be processed, the frames must be stored until their use. Normally, the receiver does not have enough storage space, especially if it is receiving data from many sources. This may result in either the discarding of frames or denial of service. To prevent the receiver from becoming overwhelmed with frames, we somehow need to tell the sender to slow down. There must be feedback from the receiver to the sender.

The protocol we discuss now is called the Stop-and-Wait Protocol because the sender sends one frame, stops until it receives confirmation from the receiver (okay to go ahead), and then sends the next frame. We still have unidirectional communication for data frames, but auxiliary ACK frames (simple tokens of acknowledgment) travel from the other direction. Flow control is added to our previous protocol.

Design

Comparing this figure with Figure above, we can see the traffic on the forward channel (from sender to receiver) and the reverse channel. At any time, there is either one data frame on the forward channel or one ACK frame on the reverse channel. We therefore need a half-duplex link.



NOISY CHANNELS

Although the Stop-and-Wait Protocol gives us an idea of how to add flow control to its predecessor, noiseless channels are nonexistent. There are three protocols in this section that use error control

Stop-and-Wait Automatic Repeat Request

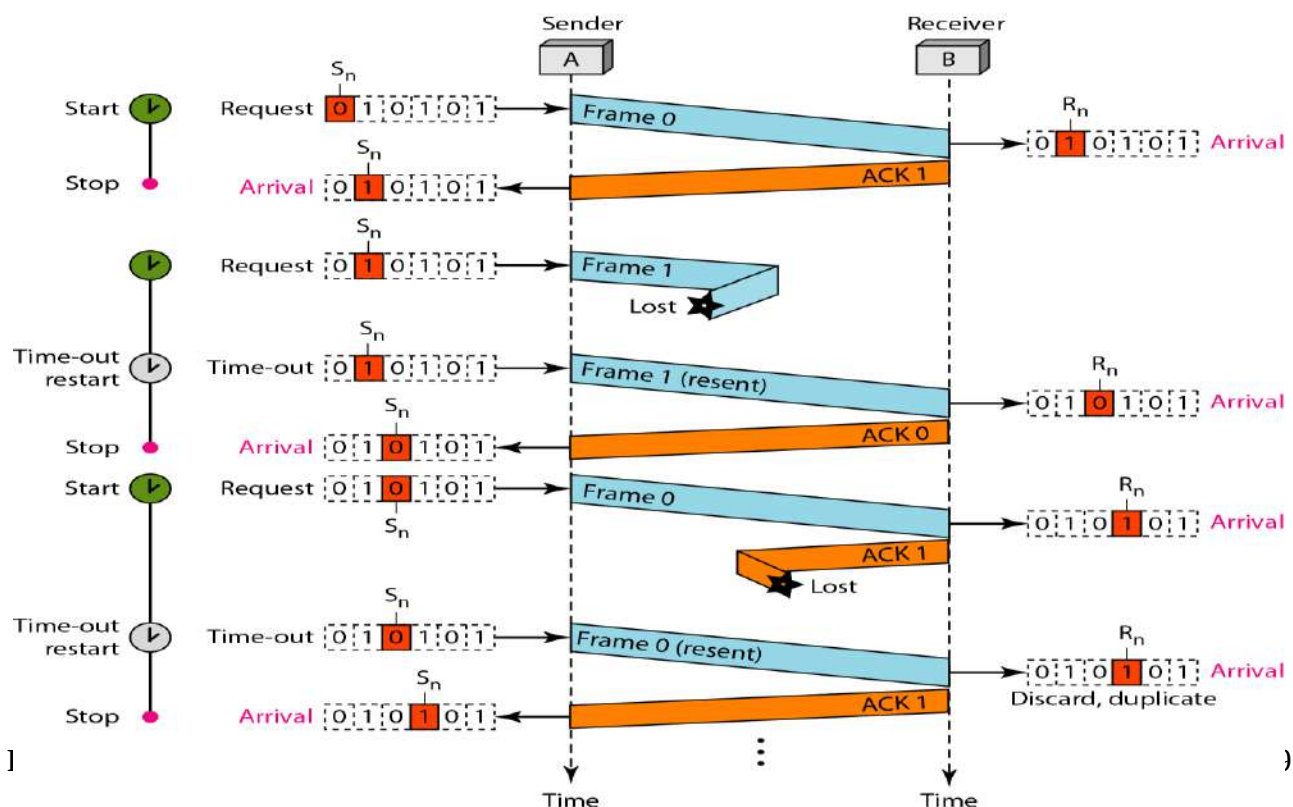
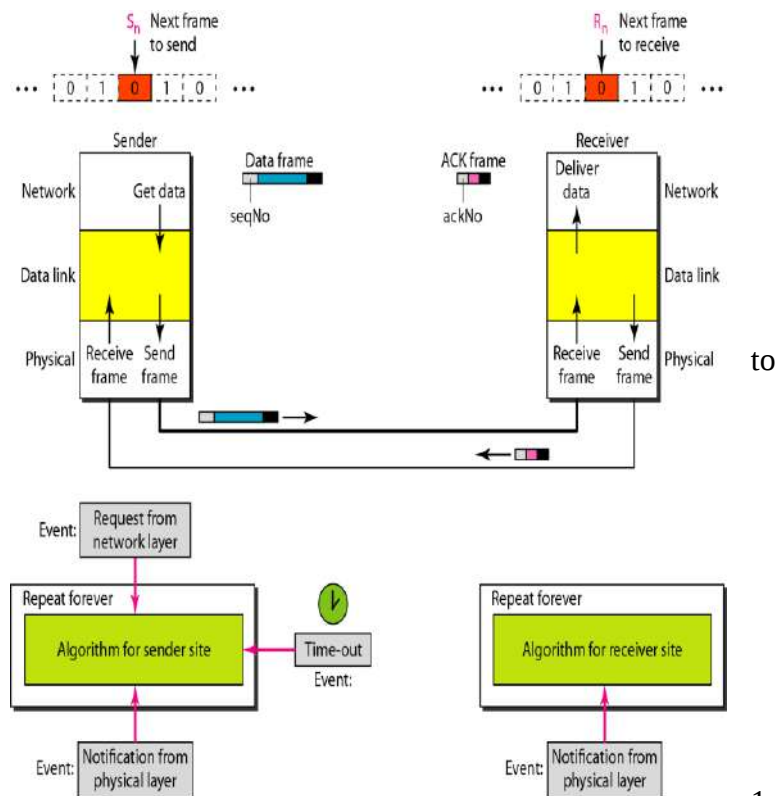
The Stop-and-Wait Automatic Repeat Request (Stop-and-Wait ARQ), adds a simple error control mechanism to the Stop-and-Wait Protocol. When the frame arrives at the receiver site, it is checked and if it is corrupted, it is silently discarded. The detection of errors in this protocol is manifested by the silence of the receiver.

Lost frames are more difficult to handle than corrupted ones. The received frame could be the correct one, or a duplicate, or a frame out of order. The solution is to number the frames. When the receiver receives a data frame that is out of order, this means that frames were either lost or duplicated. The corrupted and lost frames need to be resent in this protocol. If the receiver does not respond when there is an error, how can the sender know which frame to resend? To remedy this

problem, the sender keeps a copy of the sent frame. At the same time, it starts a timer. If the timer expires and there is no ACK for the sent frame, the frame is resent, the copy is held, and the timer is restarted. Since the protocol uses the stop-and-wait mechanism, there is only one specific frame that needs an ACK even though several copies of the same frame can be in the network. Since an ACK frame can also be corrupted and lost, it too needs redundancy bits and a sequence number. The ACK frame for this protocol has a sequence number field. In this protocol, the sender simply discards a corrupted ACK frame or ignores an out-of-order one. The sequence numbers of frames are based on modulo-2 arithmetic. The acknowledgment number always announces in modulo-2 arithmetic the sequence number of the next frame expected.

Design

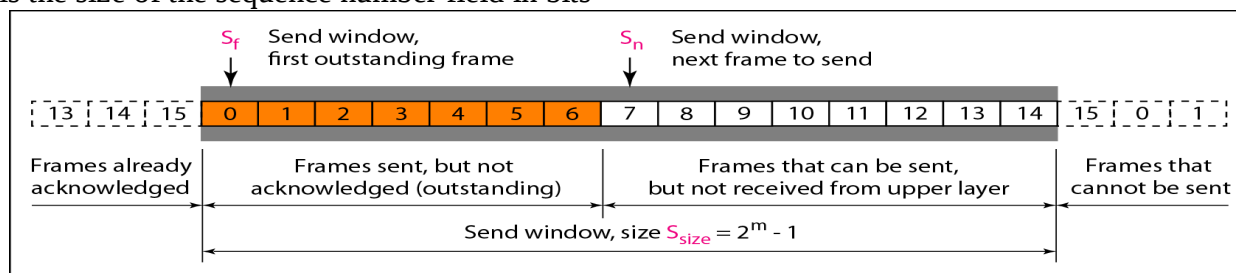
Figure below shows the design of the Stop-and-WaitARQ Protocol. The sending device keeps a copy of the last frame transmitted until it receives an acknowledgment for that frame. A data frame uses a seqNo (sequence number); an ACK frame uses an ackNo (acknowledgment number). The sender has a control variable, which we call S_n (sender, next frame send), that holds the sequence number for the next frame to be sent (0 or 1). The receiver has a control variable, which we call R_n (receiver, next frame expected), that holds the number of the next frame expected. When a frame is sent, the value of S_n is incremented (modulo-2), which means if it is 0, it becomes 1 and vice versa. When a frame is received, the value of R_n is incremented (modulo-2), which means if it is 0, it becomes



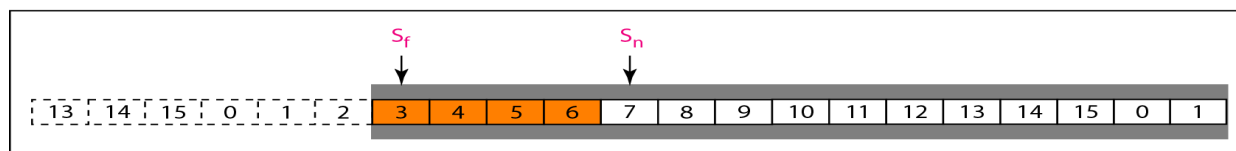
and vice versa. Variable S_n points to the slot that matches the sequence number of the frame that has been sent, but not acknowledged; R_n points to the slot that matches the sequence number of the expected frame.

Go-Back-N Automatic Repeat Request

To improve the efficiency of transmission (filling the pipe), multiple frames must be in transition while waiting for acknowledgment. Go-Back-N Automatic Repeat Request protocol can send several frames before receiving acknowledgments; copies of these frames are kept until the acknowledgments arrive. In the Go-Back-N Protocol, the sequence numbers are modulo 2^m , where m is the size of the sequence number field in bits



a. Send window before sliding



b. Send window after sliding

In this protocol (and the next), the sliding window is an abstract concept that defines the range of sequence numbers that is the concern of the sender and receiver. In other words, the sender and receiver need to deal with only part of the possible sequence numbers. The range which is the concern of the sender is called the send sliding window; the range that is the concern of the receiver is called the receive sliding window.

The send window is an imaginary box covering the sequence numbers of the data frames which can be in transit. In each window position, some of these sequence numbers define the frames that have been sent; others define those that can be sent. The maximum size of the window is $2^m - 1$

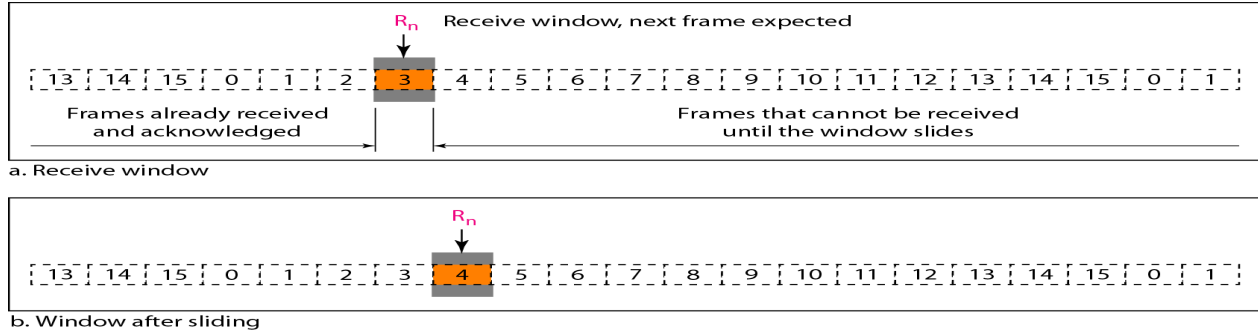
The window at any time divides the possible sequence numbers into four regions.

- The first region, from the far left to the left wall of the window, defines the sequence numbers belonging to frames that are already acknowledged. The sender does not worry about these frames and keeps no copies of them.
- The second region, colored defines the range of sequence numbers belonging to the frames that are sent and have an unknown status. The sender needs to wait to find out if these frames have been received or were lost. We call these outstanding frames.
- The third range, white in the figure, defines the range of sequence numbers for frames that can be sent; however, the corresponding data packets have not yet been received from the network layer.
- Finally, the fourth region defines sequence numbers that cannot be used until the window slides.

The window itself is an abstraction; three variables define its size and location at any time.

- S_f (send window, the first outstanding frame),
- S_n (send window, the next frame to be sent), and
- S_{size} (send window, size).

The receive window makes sure that the correct data frames are received and that the correct acknowledgments are sent. The size of the receive window is always 1. The receiver is always looking for the arrival of a specific frame. Any frame arriving out of order is discarded and needs to be resent. The sequence numbers to the left of the window belong to the frames already received and acknowledged; the sequence numbers to the right of this window define the frames that cannot be received. Any received frame with a sequence number in these two regions is discarded. Only a



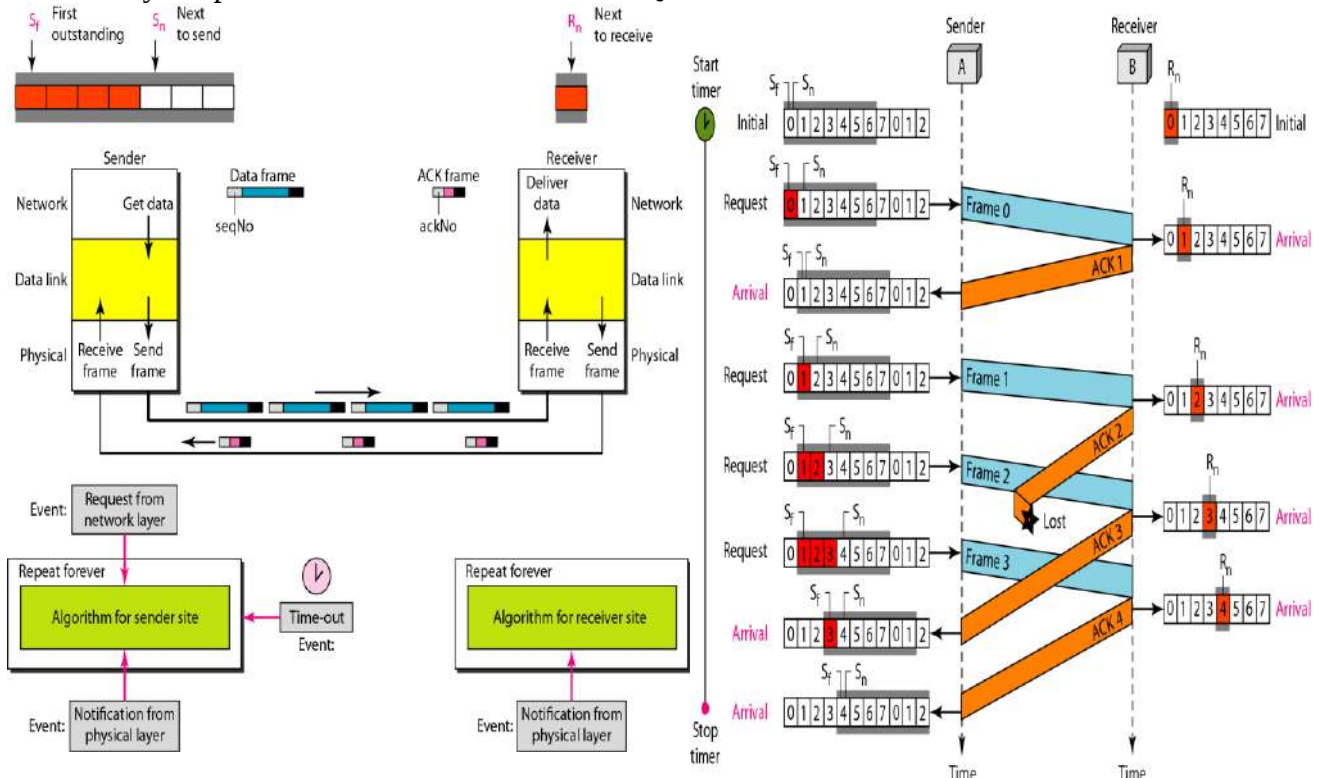
frame with a sequence number matching the value of R_n is accepted and acknowledged.

Timers: Although there can be a timer for each frame that is sent, in our protocol we use only one. The reason is that the timer for the first outstanding frame always expires first; we send all outstanding frames when this timer expires.

Acknowledgment

The receiver sends a positive acknowledgment if a frame has arrived safe and sound and in order. If a frame is damaged or is received out of order, the receiver is silent and will discard all subsequent frames until it receives the one it is expecting. The silence of the receiver causes the timer of the unacknowledged frame at the sender site to expire. This, in turn, causes the sender to go back and resend all frames, beginning with the one with the expired timer. The receiver does not have to acknowledge each frame received. It can send one cumulative acknowledgment for several frames.

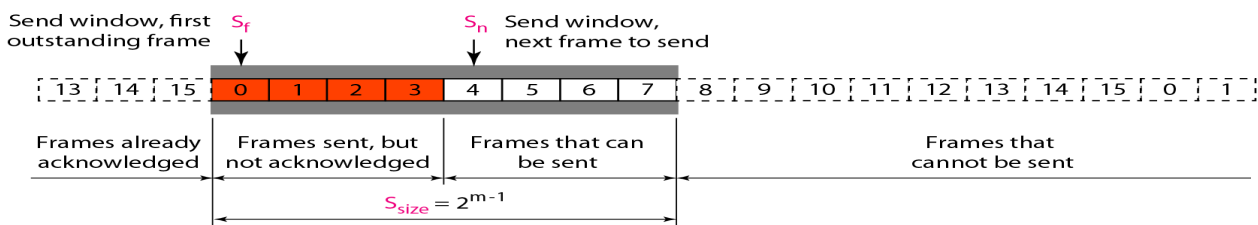
Resending a Frame When the timer expires, the sender resends all outstanding frames. For example, suppose the sender has already sent frame 6, but the timer for frame 3 expires. This means that frame 3 has not been acknowledged; the sender goes back and sends frames 3, 4, 5, and 6 again. That is why the protocol is called *Go-Back-N ARQ*.



Selective Repeat Automatic Repeat Request

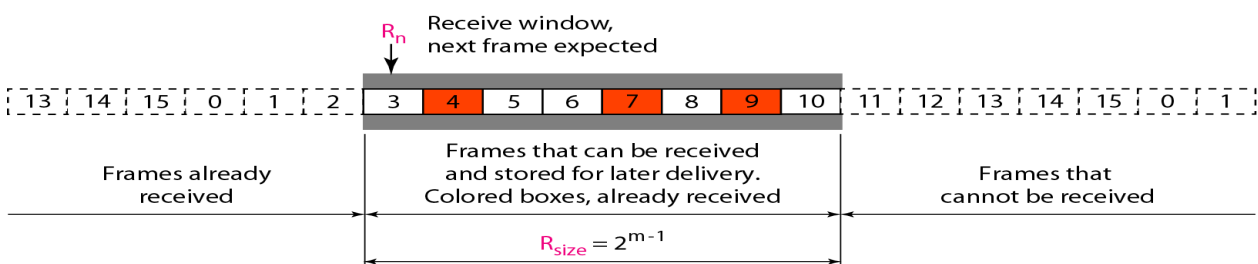
Go-Back-N ARQ simplifies the process at the receiver site. The receiver keeps track of only one variable, and there is no need to buffer out-of-order frames; they are simply discarded. However, this protocol is very inefficient for a noisy link. In a noisy link a frame has a higher probability of damage, which means the resending of multiple frames. This resending uses up the bandwidth and slows down the transmission. For noisy links, there is another mechanism that does not resend N frames when just one frame is damaged; only the damaged frame is resent. This mechanism is called Selective Repeat ARQ. It is more efficient for noisy links, but the processing at the receiver is more complex.

Windows:



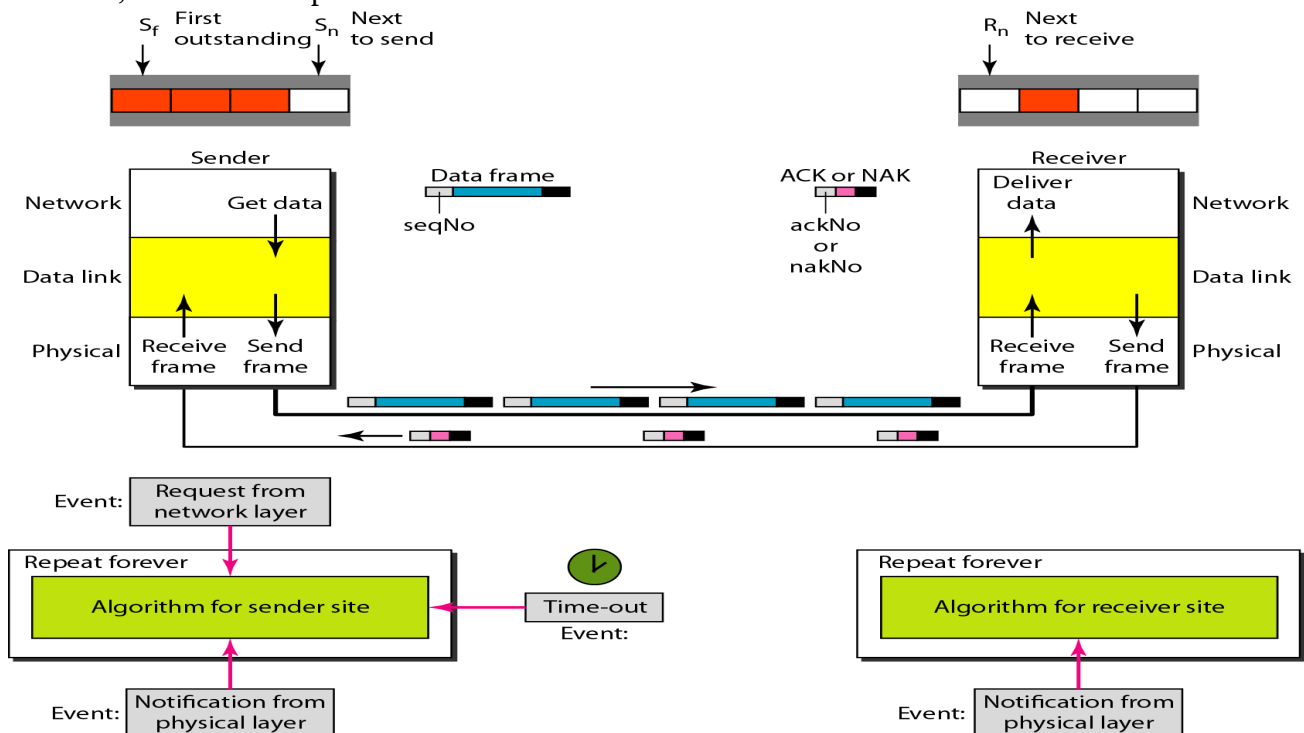
The Selective Repeat Protocol also uses two windows: a send window and a receive window. However, there are differences between the windows in this protocol and the ones in *Go-Back-N*. First, the size of the send window is much smaller; it is 2^{m-1} . The reason for this will be discussed later. Second, the receive window is the same size as the send window. The send window maximum size can be 2^{m-1} . For example, if $m = 4$, the sequence numbers go from 0 to 15, but the size of the window is just 8 (it is 15 in the *Go-Back-N* Protocol). The smaller window size means less efficiency in filling the pipe, but the fact that there are fewer duplicate frames can compensate for this. The protocol uses the same variables as we discussed for *Go-Back-N*.

The receive window in Selective Repeat is totally different from the one in *Go-Back-N*. First, the size of the receive window is the same as the size of the send window (2^{m-1}). The Selective Repeat Protocol allows as many frames as the size of the receive window to arrive out of order and be kept until there is a set of in-order frames to be delivered to the network



layer. Because the sizes of the send window and receive window are the same, all the frames in the send frame can arrive out of order and be stored until they can be delivered.

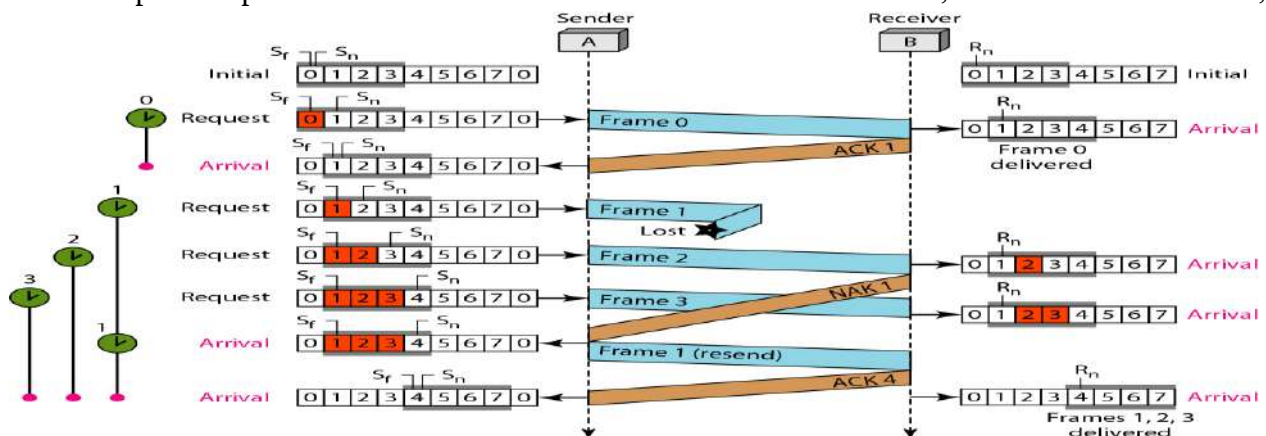
The design is shown below. Its similar some extent similar to the one we described for the GO-Back-N, but more complicated.



One main difference is the number of timers. Here, each frame sent or resent needs a timer, which means that the timers need to be numbered (0, 1, 2, and 3). The timer for frame 0 starts at the first request, but stops when the ACK for this frame arrives. The timer for frame 1 starts at the second request, restarts when a NAK arrives, and finally stops when the last ACK arrives. The other two timers start when the corresponding frames are sent and stop at the last arrival event.

At the receiver site we need to distinguish between the acceptance of a frame and its delivery to the network layer. At the second arrival, frame 2 arrives and is stored and marked (colored slot), but it cannot be delivered because frame 1 is missing. At the next arrival, frame 3 arrives and is marked and stored, but still none of the frames can be delivered. Only at the last arrival, when finally a copy of frame 1 arrives, can frames 1, 2, and 3 be delivered to the network layer. There are two conditions for the delivery of frames to the network layer: First, a set of consecutive frames must have arrived. Second, the set starts from the beginning of the window. After the first arrival, there was only one frame and it started from the beginning of the window. After the last arrival, there are three frames and the first one starts from the beginning of the window.

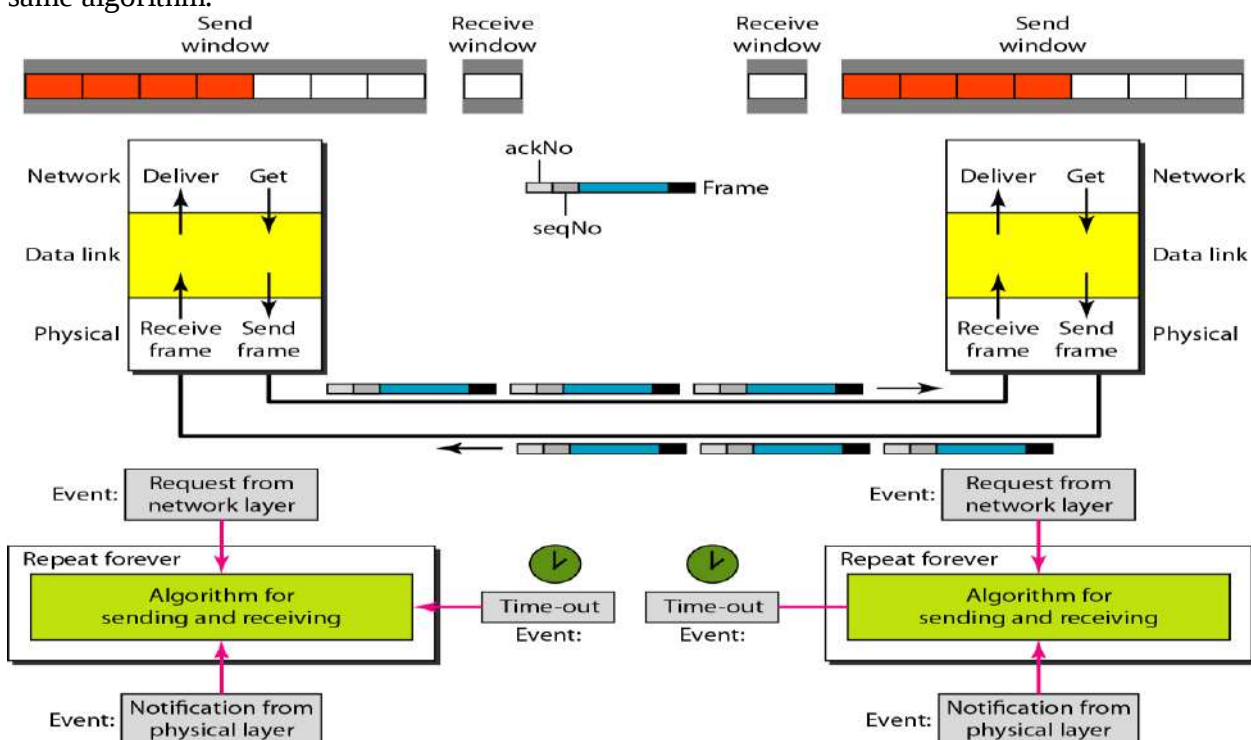
Another important point is that a NAK is sent after the second arrival, but not after the third,



although both situations look the same. The reason is that the protocol does not want to crowd the network with unnecessary NAKs and unnecessary resent frames. The second NAK would still be NAKI to inform the sender to resend frame 1 again; this has already been done. The first NAK sent is remembered (using the nakSent variable) and is not sent again until the frame slides. A NAK is sent once for each window position and defines the first slot in the window. The next point is about the ACKs. Notice that only two ACKs are sent here. The first one acknowledges only the first frame; the second one acknowledges three frames. In Selective Repeat, ACKs are sent when data are delivered to the network layer. If the data belonging to n frames are delivered in one shot, only one ACK is sent for all of them.

Piggybacking

The three protocols in this section are all unidirectional: data frames flow in only one direction although control information such as ACK and NAK frames can travel in the other direction. In real life, data frames are normally flowing in both directions: from node A to node B and from node B to node A. This means that the control information also needs to flow in both directions. A technique called **piggybacking** is used to improve the efficiency of the bidirectional protocols. When a frame is carrying data from A to B, it can also carry control information about arrived (or lost) frames from B; when a frame is carrying data from B to A, it can also carry control information about the arrived (or lost) frames from A. the design below is for a Go-Back-N ARQ using piggybacking in. Note that each node now has two windows: one send window and one receive window. Both also need to use a timer. Both are involved in three types of events: request, arrival, and time-out. However, the arrival event here is complicated; when a frame arrives, the site needs to handle control information as well as the frame itself. Both of these concerns must be taken care of in one event, the arrival event. The request event uses only the send window at each site; the arrival event needs to use both windows. An important point about piggybacking is that both sites must use the same algorithm.

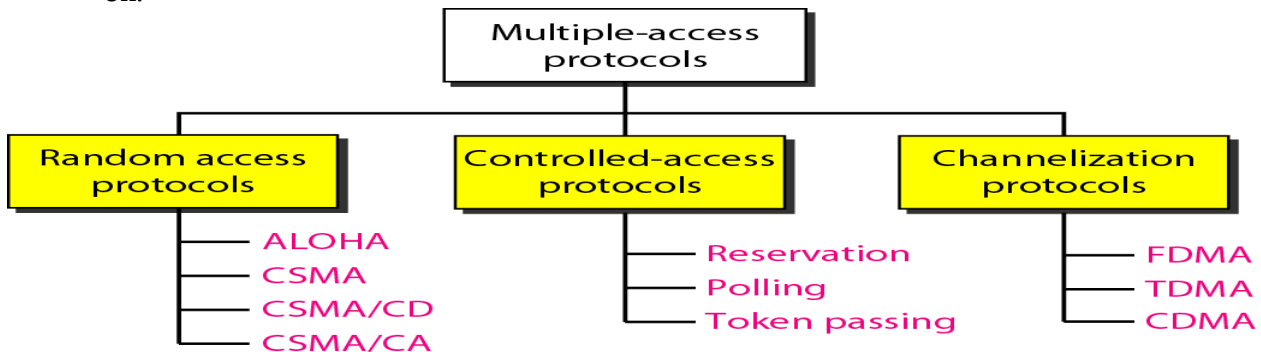


MULTIPLE ACCESS

When nodes are connected to a common link, called multipoint or broadcast link, we need a multiple-access protocol to coordinate access to the link.

Multiple-Access Protocols

- Nodes or stations are connected to or use a common link, called a multipoint or broadcast link.
- Problem of controlling the access to the medium is similar to the rules of speaking in an assembly.
- The procedures guarantee that the right to speak is upheld and ensure that two people do not speak at the same time, do not interrupt each other, do not monopolize the discussion, and so on.



RANDOM ACCESS

In random access or contention methods, no station is superior to another station and none is assigned the control over another. No station permits, or does not permit, another station to send. At each instance, a station that has data to send uses a procedure defined by the protocol to make a decision on whether or not to send. This decision depends on the state of the medium (idle or busy). In other words, each station can transmit when it desires on the condition that it follows the predefined procedure, including the testing of the state of the medium.

Two features give this method its name.

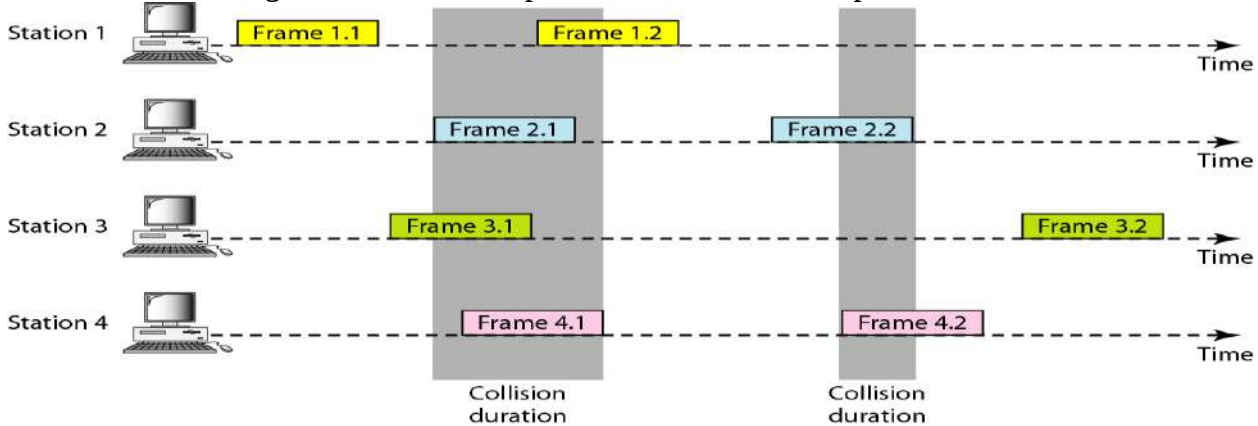
- First, there is no scheduled time for a station to transmit. Transmission is random among the stations. That is why these methods are called *random access*.
- Second, no rules specify which station should send next. Stations compete with one another to access the medium. That is why these methods are also called *contention* methods.
- In a random access method, each station has the right to the medium without being controlled by any other station. However, if more than one station tries to send, there is an access conflict-collision-and the frames will be either destroyed or modified.
- To avoid access conflict or to resolve it when it happens, each station follows a procedure that answers the following questions:
 - When can the station access the medium?
 - What can the station do if the medium is busy?
 - How can the station determine the success or failure of the transmission?
 - What can the station do if there is an access conflict?

ALOHA

ALOHA, the earliest random access method, was developed at the University of Hawaii in early 1970. It was designed for a radio (wireless) LAN, but it can be used on any shared medium. It is obvious that there are potential collisions in this arrangement. The medium is shared between the stations. When a station sends data, another station may attempt to do so at the same time. The data from the two stations collide and become garbled.

Pure ALOHA The original ALOHA protocol is called pure ALOHA. This is a simple, but elegant protocol. The idea is that each station sends a frame whenever it has a frame to send. However, since there is only one channel to share, there is the possibility of collision between frames from

different stations. Figure shows an example of frame collisions in pure ALOHA.



There are four stations (unrealistic assumption) that contend with one another for access to the shared channel. The figure shows that each station sends two frames; there are a total of eight frames on the shared medium. Some of these frames collide because multiple frames are in contention for the shared channel.

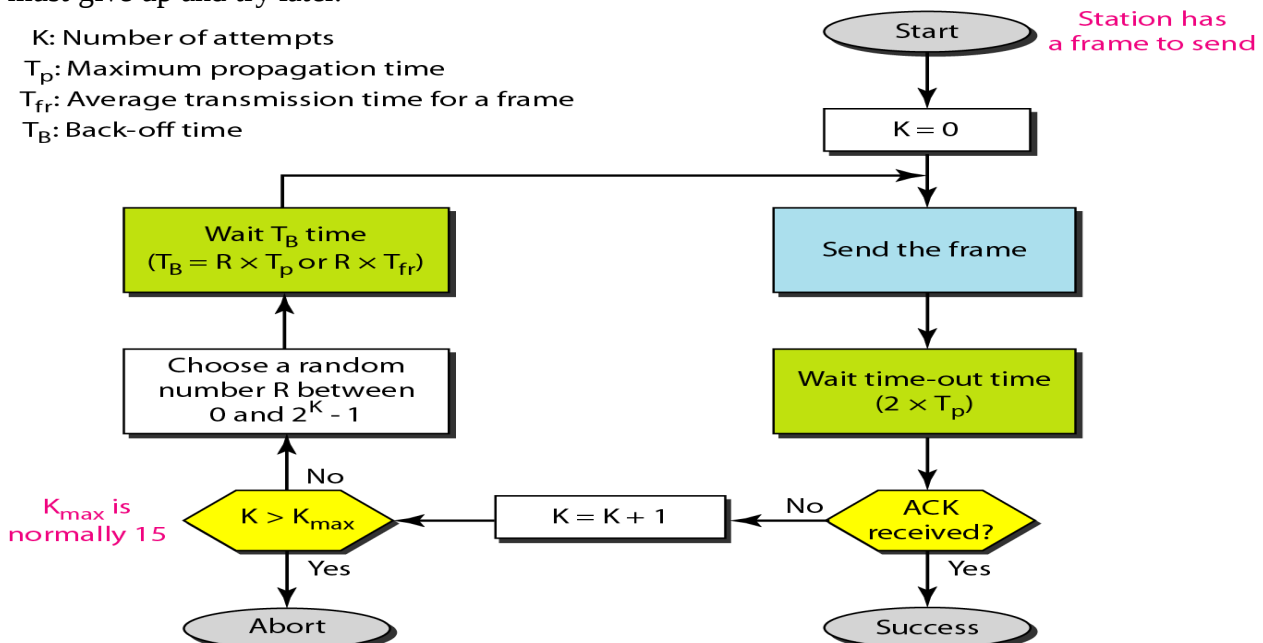
When a station sends a frame, it expects the receiver to send an acknowledgment. If the acknowledgment does not arrive after a time-out period, the station assumes that the frame (or the acknowledgment) has been destroyed and resends the frame. A collision involves two or more stations. If all these stations try to resend their frames after the time-out, the frames will collide again. Pure ALOHA dictates that when the time-out period passes, each station waits a random amount of time before resending its frame. The randomness will help avoid more collisions. We call this time the back-off time T_B . Pure ALOHA has a second method to prevent congesting the channel with retransmitted frames. After a maximum number of retransmission attempts K_{max} a station must give up and try later.

K : Number of attempts

T_p : Maximum propagation time

T_{fr} : Average transmission time for a frame

T_B : Back-off time

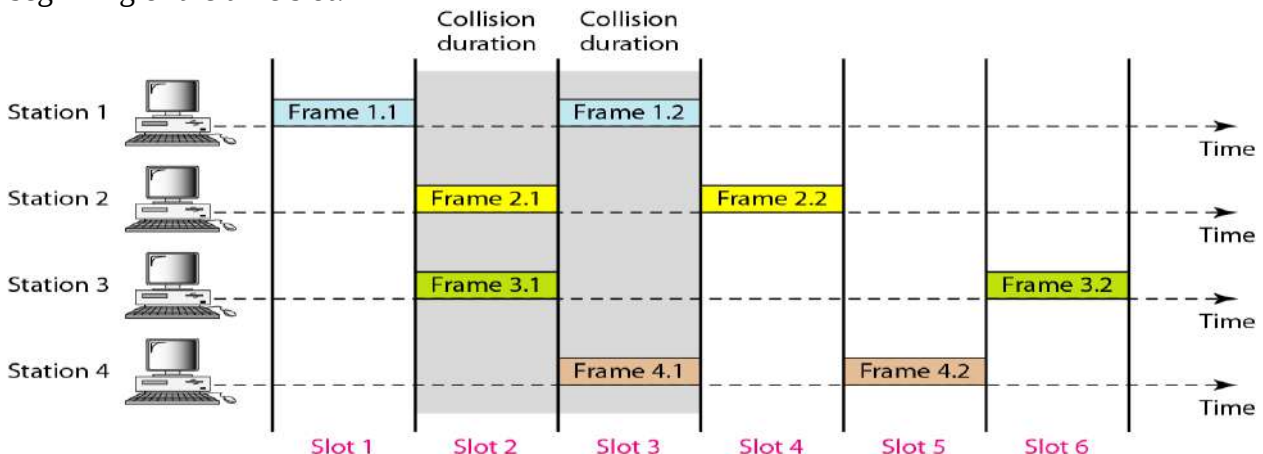


The time-out period is equal to the maximum possible round-trip propagation delay, which is twice the amount of time required to send a frame between the two most widely separated stations ($2 \times T_p$). The back-off time T_B is a random value that normally depends on K (the number of attempted unsuccessful transmissions).

Slotted ALOHA

Pure ALOHA has a vulnerable time of $2 \times T_{fr}$. This is so because there is no rule that defines when the station can send. A station may send soon after another station has started or soon before another

station has finished. Slotted ALOHA was invented to improve the efficiency of pure ALOHA. In slotted ALOHA we divide the time into slots of T_{fr} seconds and force the station to send only at the beginning of the time slot.

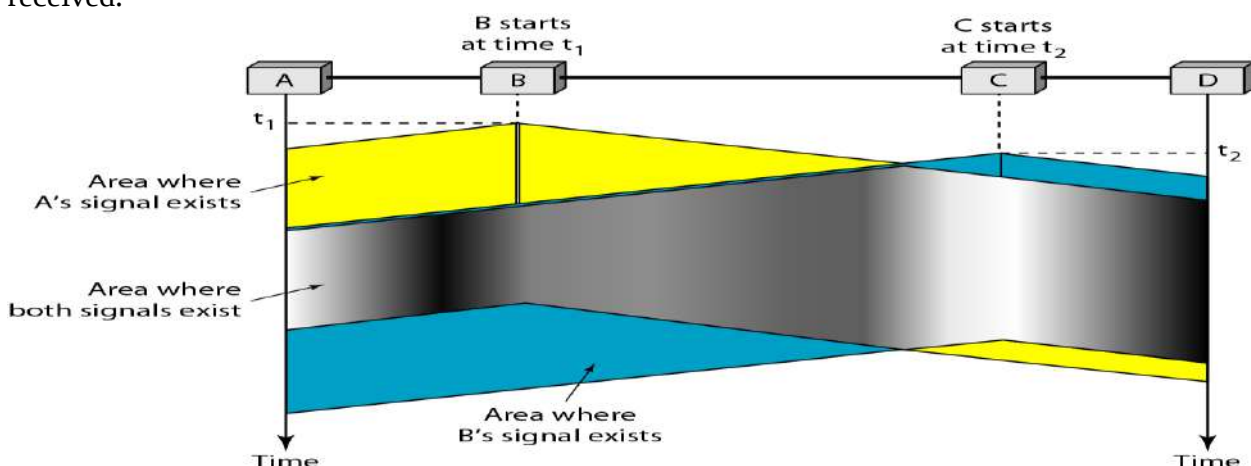


Because a station is allowed to send only at the beginning of the synchronized time slot, if a station misses this moment, it must wait until the beginning of the next time slot. This means that the station which started at the beginning of this slot has already finished sending its frame. Of course, there is still the possibility of collision if two stations try to send at the beginning of the same time slot. However, the vulnerable time is now reduced to one-half, equal to T_{fr} .

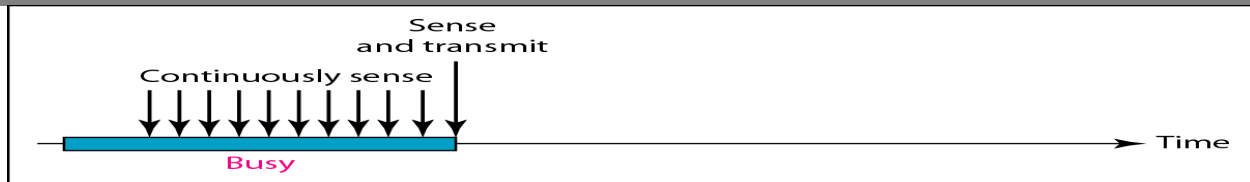
Carrier Sense Multiple Access (CSMA)

The chance of collision can be reduced if a station senses the medium before trying to use it. Carrier sense multiple access (CSMA) requires that each station first listen to the medium (or check the state of the medium) before sending. In other words, CSMA is based on the principle "sense before transmit" or "listen before talk."

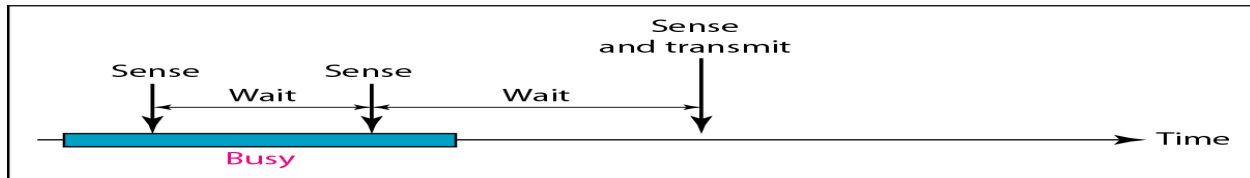
CSMA can reduce the possibility of collision, but it cannot eliminate it. The possibility of collision still exists because of propagation delay; when a station sends a frame, it still takes time (although very short) for the first bit to reach every station and for every station to sense it. In other words, a station may sense the medium and find it idle, only because the first bit sent by another station has not yet been received.



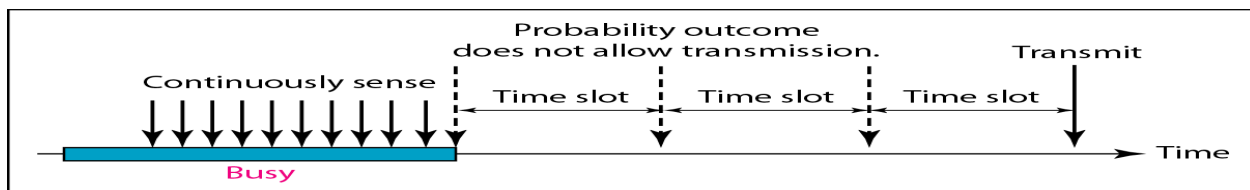
Persistence Methods Three methods have been devised to answer these questions: the I-persistent method, the nonpersistent method, and the p-persistent method. Figure below shows the behavior of three persistence methods when a station finds a channel busy.



a. 1-persistent

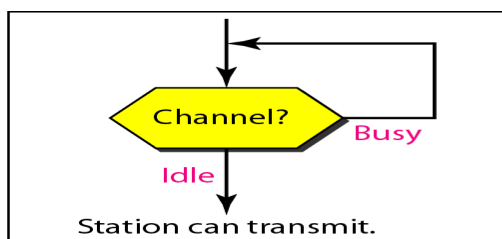


b. Nonpersistent

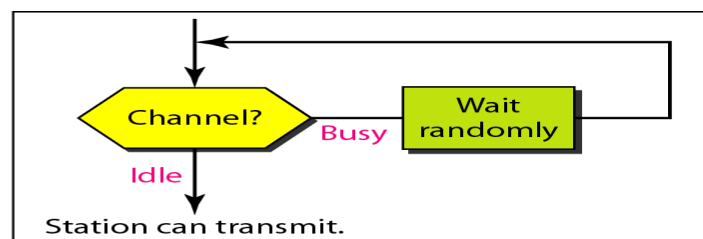


c. p-persistent

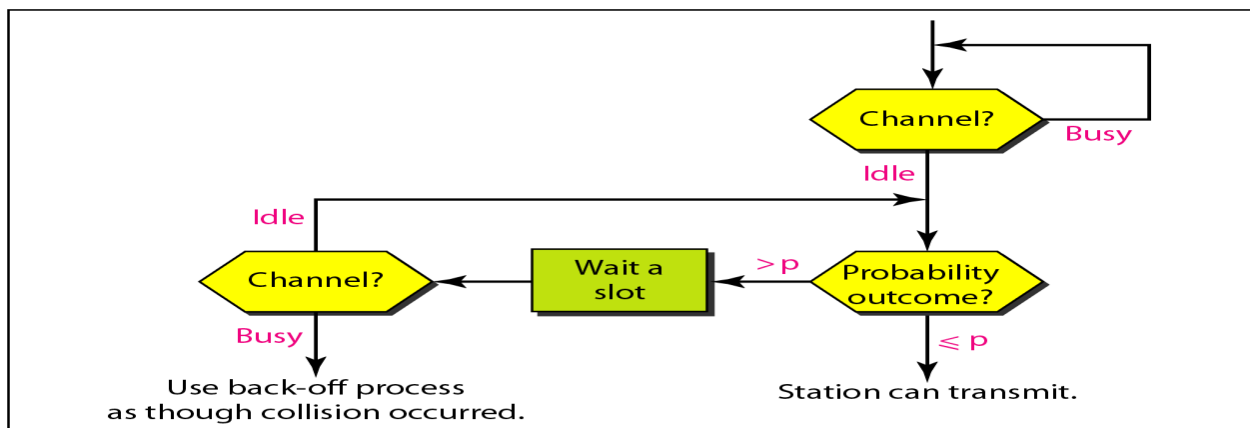
Flow diagram for 1-persistent, Nonpersistent, p-persistent method



a. 1-persistent



b. Nonpersistent



c. p-persistent

I-Persistent The **I-persistent method** is simple and straightforward. In this method, after the station finds the line idle, it sends its frame immediately (with probability 1). This method has the highest chance of collision because two or more stations may find the line idle and send their frames immediately.

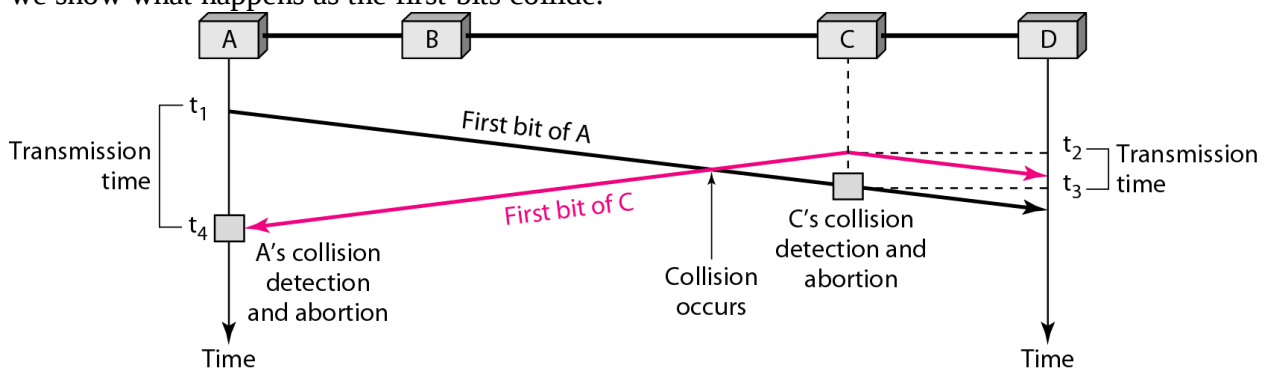
Nonpersistent In the **nonpersistent method**, a station that has a frame to send senses the line. If the line is idle, it sends immediately. If the line is not idle, it waits random amount of time and then senses the line again. The nonpersistent approach reduces the chance of collision because it is unlikely that two or more stations will wait the same amount of time and retry to send simultaneously. However, this method reduces the efficiency of the network because the medium

remains idle when there may be stations with frames to send. **p-Persistent** The **p-persistent method** is used if the channel has time slots with a slot duration equal to or greater than the maximum propagation time. The p-persistent approach combines the advantages of the other two strategies. It reduces the chance of collision and improves efficiency. In this method, after the station finds the line idle it follows these steps:

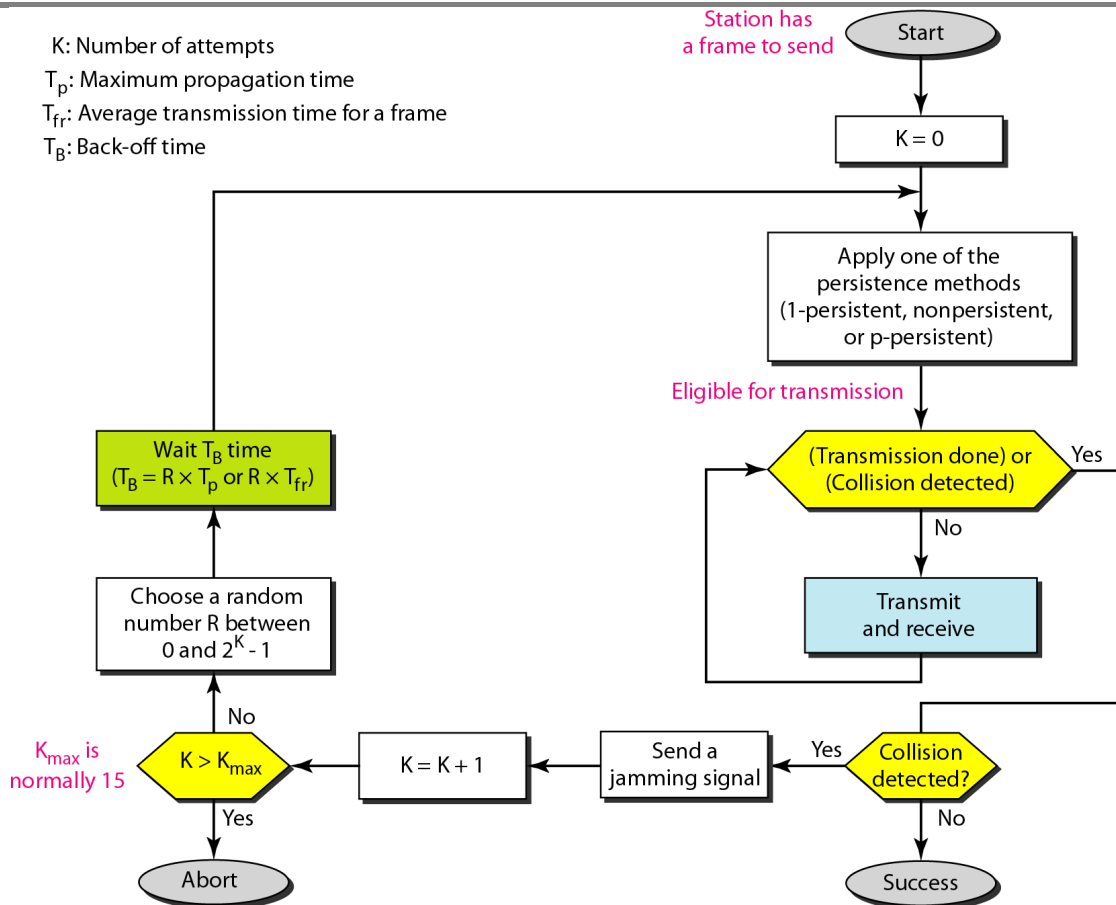
1. With probability p , the station sends its frame.
2. With probability $q = 1 - p$, the station waits for the beginning of the next time slot and checks the line again.
 - a. If the line is idle, it goes to step 1.
 - b. If the line is busy, it acts as though a collision has occurred and uses the backoff procedure.

Carrier Sense Multiple Access with Collision Detection (CSMA/CD)

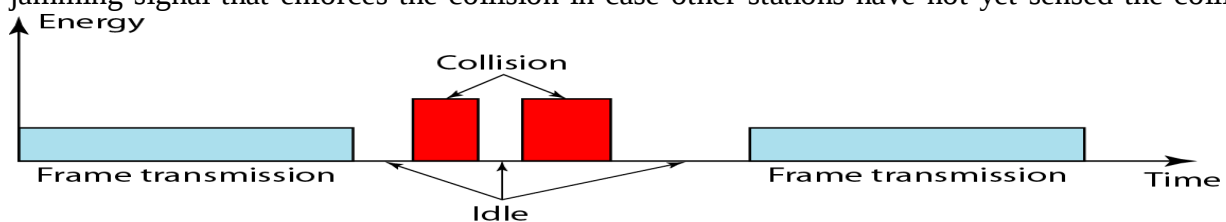
The CSMA method does not specify the procedure following a collision. Carrier sense multiple access with collision detection (CSMA/CD) augments the algorithm to handle the collision. In this method, a station monitors the medium after it sends a frame to see if the transmission was successful. If so, the station is finished. If, however, there is a collision, the frame is sent again. Although each station continues to send bits in the frame until it detects the collision, we show what happens as the first bits collide.



At time t_1 , station A has executed its persistence procedure and starts sending the bits of its frame. At time t_2 , station C has not yet sensed the first bit sent by A. Station C executes its persistence procedure and starts sending the bits in its frame, which propagate both to the left and to the right. The collision occurs sometime after time t_2 . Station C detects a collision at time t_3 when it receives the first bit of A's frame. Station C immediately (or after a short time, but we assume immediately) aborts transmission. Station A detects collision at time t_4 when it receives the first bit of C's frame; it also immediately aborts transmission.



In *CSMA/CD*, transmission and collision detection is a continuous process. We do not send the entire frame and then look for a collision. The station transmits and receives continuously and simultaneously (using two different ports). We use a loop to show that transmission is a continuous process. We constantly monitor in order to detect one of two conditions: either transmission is finished or a collision is detected. Either event stops transmission. When we come out of the loop, if a collision has not been detected, it means that transmission is complete; the entire frame is transmitted. Otherwise, a collision has occurred. The third difference is the sending of a short jamming signal that enforces the collision in case other stations have not yet sensed the collision.



Energy Level: The level of energy in a channel can have three values: zero, normal, and abnormal.

- At the zero level, the channel is idle.
- At the normal level, a station has successfully captured the channel and is sending its frame.
- At the abnormal level, there is a collision and the level of the energy is twice the normal level.

A station that has a frame to send or is sending a frame needs to monitor the energy level to determine if the channel is idle, busy, or in collision mode.

Carrier Sense Multiple Access with Collision Avoidance (CSMA/CA)

The basic idea behind *CSMA/CD* is that a station needs to be able to receive while transmitting to detect a collision. When there is no collision, the station receives one signal: its own signal. When there is a collision, the station receives two signals: its own signal and the signal transmitted by a

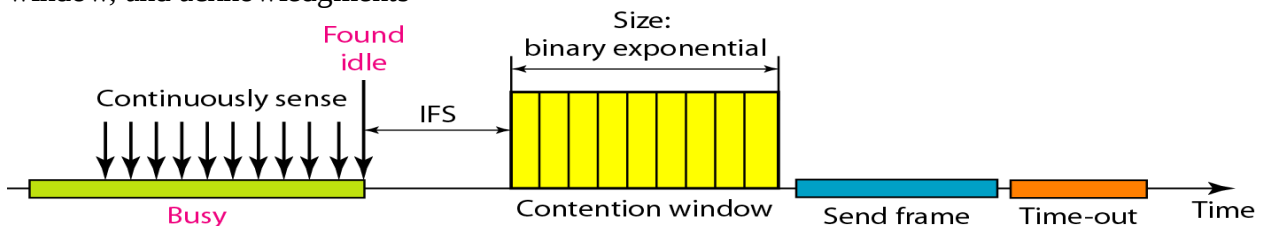
second station.

To distinguish between these two cases, the received signals in these two cases must be significantly different. In other words, the signal from the second station needs to add a significant amount of energy to the one created by the first station.

In a wired network, the received signal has almost the same energy as the sent signal because either the length of the cable is short or there are repeaters that amplify the energy between the sender and the receiver. This means that in a collision, the detected energy almost doubles.

However, in a wireless network, much of the sent energy is lost in transmission. The received signal has very little energy. Therefore, a collision may add only 5 to 10 percent additional energy. This is not useful for effective collision detection.

We need to avoid collisions on wireless networks because they cannot be detected. Carrier sense multiple access with collision avoidance (CSMA/CA) was invented for this network. Collisions are avoided through the use of CSMA/CA's three strategies: the interframe space, the contention window, and acknowledgments



Interframe Space (IFS)

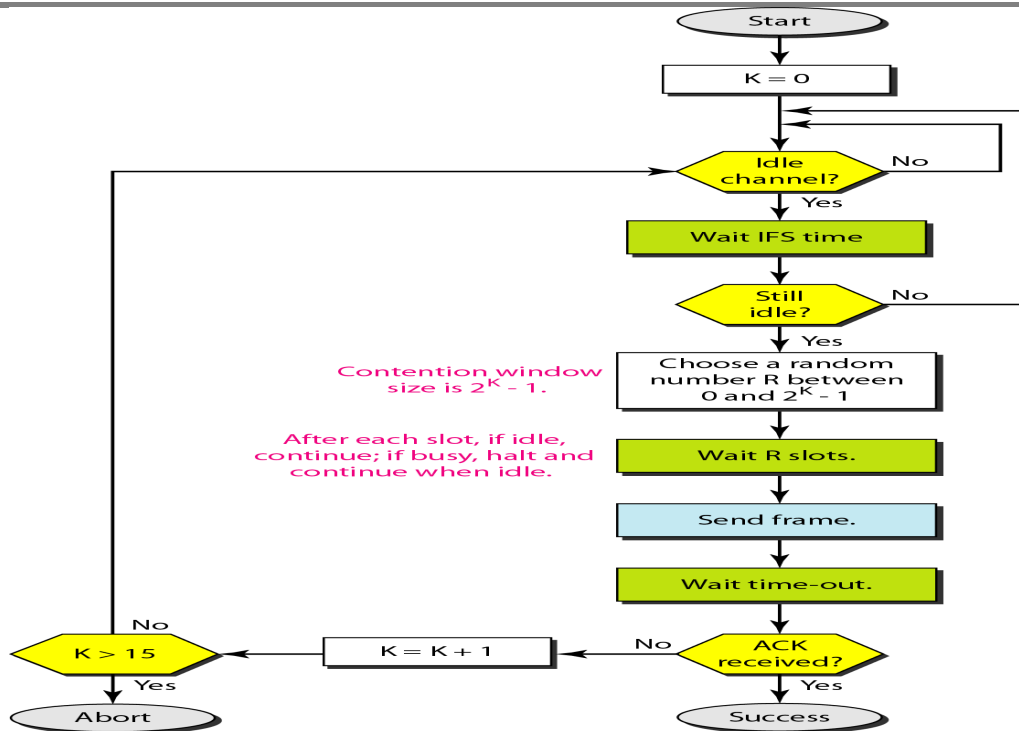
First, collisions are avoided by deferring transmission even if the channel is found idle. When an idle channel is found, the station does not send immediately. It waits for a period of time called the interframe space or IFS. Even though the channel may appear idle when it is sensed, a distant station may have already started transmitting. The distant station's signal has not yet reached this station. The IFS time allows the front of the transmitted signal by the distant station to reach this station. If after the IFS time the channel is still idle, the station can send, but it still needs to wait a time equal to the contention time. The IFS can also be used to define the priority of a station or a frame.

Contention Window

The contention window is an amount of time divided into slots. A station that is ready to send chooses a random number of slots as its wait time. The number of slots in the window changes according to the binary exponential back-off strategy. This means that it is set to one slot the first time and then doubles each time the station cannot detect an idle channel after the IFS time. One interesting point about the contention window is that the station needs to sense the channel after each time slot. However, if the station finds the channel busy, it does not restart the process; it just stops the timer and restarts it when the channel is sensed as idle. This gives priority to the station with the longest waiting time.

Acknowledgment

With all these precautions, there still may be a collision resulting in destroyed data. In addition, the data may be corrupted during the transmission. The positive acknowledgment and the time-out timer can help guarantee that the receiver has received the frame.



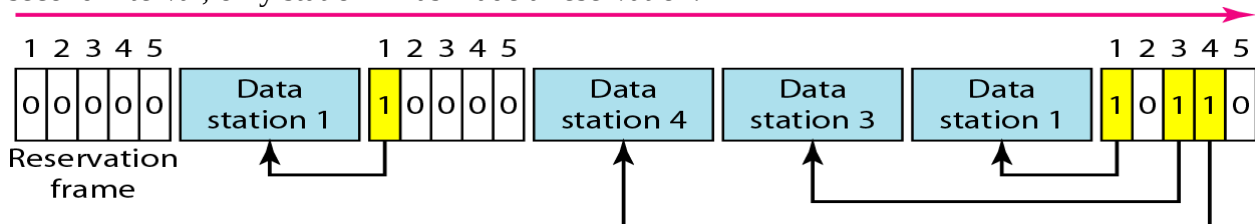
CONTROLLED ACCESS

In controlled access, the stations consult one another to find which station has the right to send. A station cannot send unless it has been authorized by other stations.

Reservation

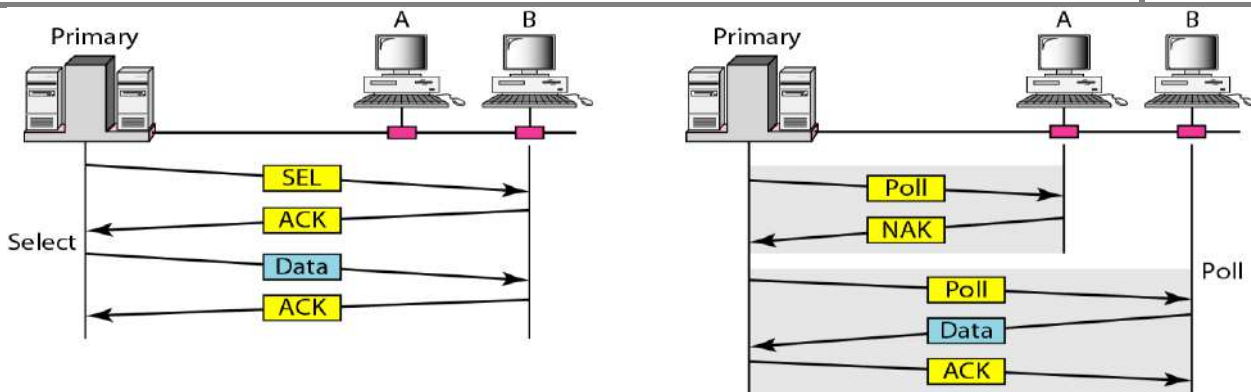
In the reservation method, a station needs to make a reservation before sending data. Time is divided into intervals. In each interval, a reservation frame precedes the data frames sent in that interval.

If there are N stations in the system, there are exactly N reservation minislots in the reservation frame. Each minislot belongs to a station. When a station needs to send a data frame, it makes a reservation in its own minislot. The stations that have made reservations can send their data frames after the reservation frame. Figure shows a situation with five stations and a five-minislot reservation frame. In the first interval, only stations 1, 3, and 4 have made reservations. In the second interval, only station 1 has made a reservation.



Polling

Polling works with topologies in which one device is designated as a primary station and the other devices are secondary stations. All data exchanges must be made through the primary device even when the ultimate destination is a secondary device. The primary device controls the link; the secondary devices follow its instructions. It is up to the primary device to determine which device is allowed to use the channel at a given time. The primary device, therefore, is always the initiator of a session



If the primary wants to receive data, it asks the secondaries if they have anything to send; this is called poll function. If the primary wants to send data, it tells the secondary to get ready to receive; this is called select function.

The *select* function is used whenever the primary device has something to send. Remember that the primary controls the link. If the primary is neither sending nor receiving data, it knows the link is available.

If it has something to send, the primary device sends it. What it does not know, however, is whether the target device is prepared to receive. So the primary must alert the secondary to the upcoming transmission and wait for an acknowledgment of the secondary's ready status. Before sending data, the primary creates and transmits a select (SEL) frame, one field of which includes the address of the intended secondary.

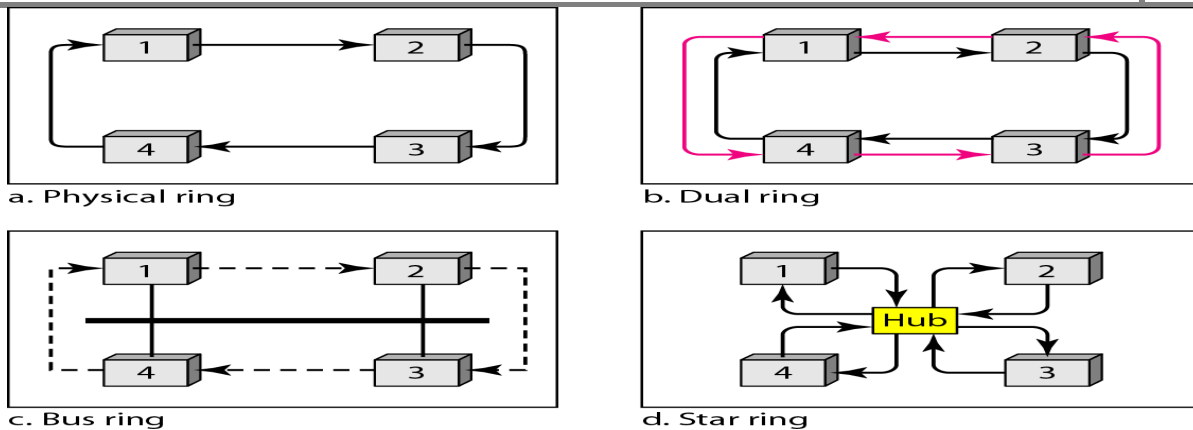
The *poll* function is used by the primary device to solicit transmissions from the secondary devices. When the primary is ready to receive data, it must ask (poll) each device in turn if it has anything to send. When the first secondary is approached, it responds either with a NAK frame if it has nothing to send or with data (in the form of a data frame) if it does. If the response is negative (a NAK frame), then the primary polls the next secondary in the same manner until it finds one with data to send. When the response is positive (a data frame), the primary reads the frame and returns an acknowledgment (ACK frame), verifying its receipt.

Token Passing

In the token-passing method, the stations in a network are organized in a logical ring. In other words, for each station, there is a *predecessor* and a *successor*. The predecessor is the station which is logically before the station in the ring; the successor is the station which is after the station in the ring. The current station is the one that is accessing the channel now. The right to this access has been passed from the predecessor to the current station. The right will be passed to the successor when the current station has no more data to send. But how is the right to access the channel passed from one station to another? In this method, a special packet called a token circulates through the ring. The possession of the token gives the station the right to access the channel and send its data. When a station has some data to send, it waits until it receives the token from its predecessor. It then holds the token and sends its data.

When the station has no more data to send, it releases the token, passing it to the next logical station in the ring. The station cannot send data until it receives the token again in the next round. In this process, when a station receives the token and has no data to send, it just passes the data to the next station.

Token management is needed for this access method. Stations must be limited in the time they can have possession of the token. The token must be monitored to ensure it has not been lost or destroyed. For example, if a station that is holding the token fails, the token will disappear from the network. Another function of token management is to assign priorities to the stations and to the types of data being transmitted. And finally, token management is needed to make low-priority stations release the token to highpriority stations.

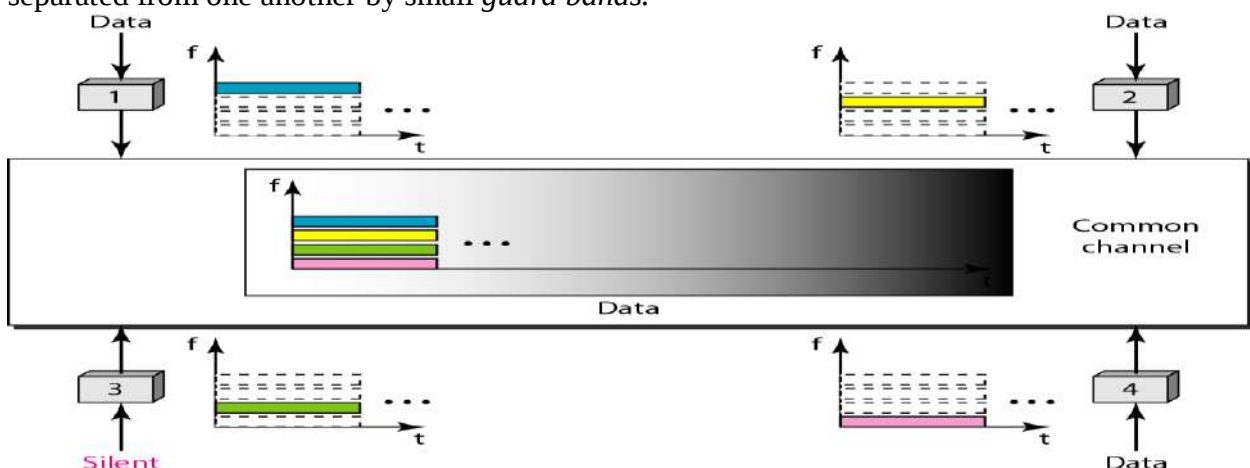


CHANNELIZATION

Channelization is a multiple-access method in which the available bandwidth of a link is shared in time, frequency, or through code, between different stations. There are three channelization protocols: FDMA, TDMA, and CDMA.

Frequency-Division Multiple Access (FDMA)

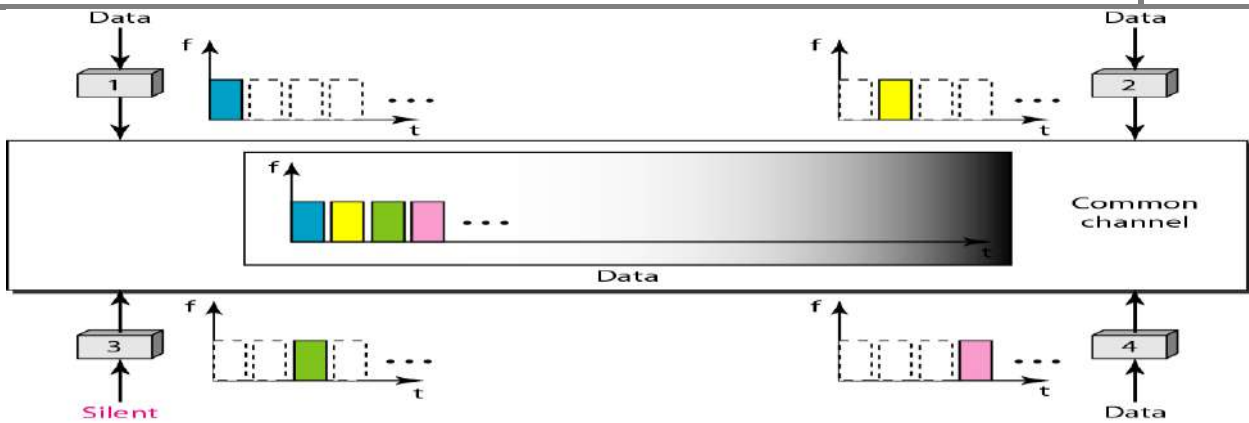
In frequency-division multiple access (FDMA), the available bandwidth is divided into frequency bands. Each station is allocated a band to send its data. In other words, each band is reserved for a specific station, and it belongs to the station all the time. Each station also uses a bandpass filter to confine the transmitter frequencies. To prevent station interferences, the allocated bands are separated from one another by small *guard bands*.



FDMA is an access method in the data link layer. The data link layer in each station tells its physical layer to make a bandpass signal from the data passed to it. The signal must be created in the allocated band. There is no physical multiplexer at the physical layer. The signals created at each station are automatically bandpass-filtered. They are mixed when they are sent to the common channel.

Time-Division Multiple Access (TDMA)

In time-division multiple access (TDMA), the stations share the bandwidth of the channel in time. Each station is allocated a time slot during which it can send data. Each station transmits its data in its assigned time slot. Figure shows the idea behind TDMA.



The main problem with TDMA lies in achieving synchronization between the different stations. Each station needs to know the beginning of its slot and the location of its slot. This may be difficult because of propagation delays introduced in the system if the stations are spread over a large area. To compensate for the delays *guard times* are inserted. Synchronization is normally accomplished by having some synchronization bits (normally referred to as preamble bits) at the beginning of each slot.

Code-Division Multiple Access (CDMA)

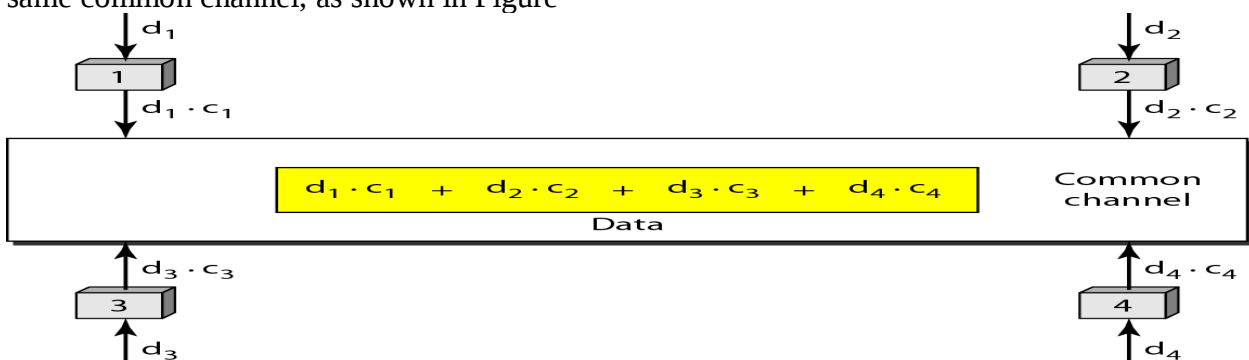
Code-division multiple access (CDMA) was conceived several decades ago. Recent advances in electronic technology have finally made its implementation possible. CDMA differs from FDMA because only one channel occupies the entire bandwidth of the link. It differs from TDMA because all stations can send data simultaneously; there is no timesharing.

CDMA is based on coding theory. Each station is assigned a code, which is a sequence of numbers called chips. Whenever a station needs to send data it sends an encoded data bit.

Let us assume we have four stations 1, 2, 3, and 4 connected to the same channel. The data from station 1 are d_1 , from station 2 are d_2 , and so on. The code assigned to the first station is c_1 , to the second is c_2 , and so on. We assume that the assigned codes have two properties.

1. If we multiply each code by another, we get 0.
2. If we multiply each code by itself, we get 4 (the number of stations).

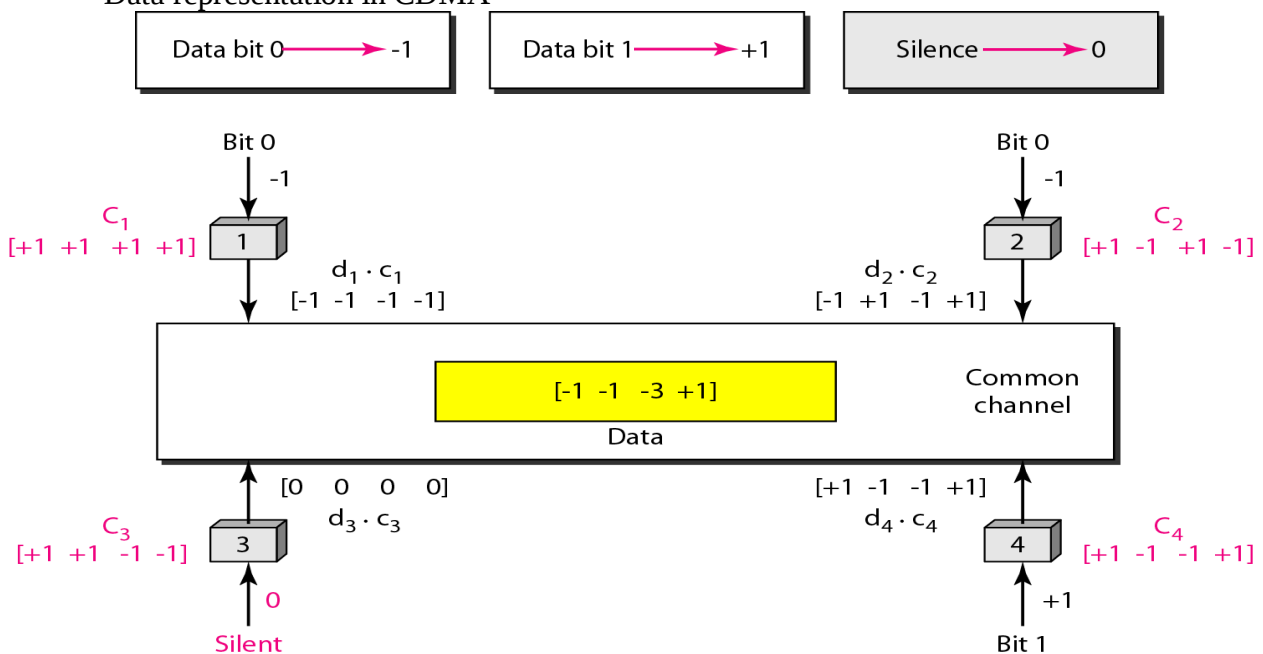
With these two properties in mind, let us see how the above four stations can send data using the same common channel, as shown in Figure



- One channel carries all transmissions simultaneously
- Two properties: If we multiply each code by another, we get 0. If we multiply each code by itself, we get 4
- Data = $(d_1 \cdot c_1 + d_2 \cdot c_2 + d_3 \cdot c_3 + d_4 \cdot c_4) \cdot c_1$
 $= d_1 \cdot c_1 \cdot c_1 + d_2 \cdot c_2 \cdot c_1 + d_3 \cdot c_3 \cdot c_1 + d_4 \cdot c_4 \cdot c_1 = 4 \cdot d_1$
- Sequence of numbers called **chips**



- Orthogonal sequences have the following properties:
 - Each sequence is made of N elements, where N is the number of stations
 - If we multiply a sequence by a number, every element in the sequence is multiplied by that element (**scalar multiplication**)
 - If we multiply two equal sequence, element by element, and add the results, we get N (**inner product**)
 - If we multiply two different sequence, element by element, and add the results, we get 0
 - Adding two sequence means adding the corresponding elements. The result is another sequence
- Data representation in CDMA



POSSIBLE QUESTIONS**5X8=40 Marks**

1. What is block coding? Explain with examples.
2. Describe the random-access protocols with appropriate sketches.
3. Elucidate the principles of cyclic codes with examples
4. Describe the selective repeat ARQ protocol for a noisy channel.
5. Elucidate the Stop-and-wait and Go-Back-N protocols with appropriate sketches.
6. List out the Channelization Protocols and explain them.
7. What are the Controlled Access Protocols? Explain with neat diagrams
8. Illustrate the principle behind the ALOHA Protocols with neat diagrams.
9. Write a note on: (i) CRC (ii) Checksum (iii) Parity code
10. Write a brief note on: (i) Framing (ii) Piggybacking



KARPAGAM ACADEMY OF HIGHER EDUCATION
(Deemed to be University Established Under Section 3 of UGC A
Pollachi Main Road, Eachanari Post, Coimbatore – 641 021.
DEPARTMENT OF COMPUTER SCIENCE
DATA COMMUNICATION NETWORKS (15CSU502)
ONLINE EXAM QUESTION BANK

UNIT-3

Questions	opt1	opt2	opt3
Transmission errors are usually detected at the.....layer of OSI model	physical	datalink	network
Transmission errors are usually corrected at the.....layer of OSI model	network	transport	
Datalink layer imposes amechanism to avoid overwhelming the receiver	flow control	error control	access control
Error control mechanism of datalink layer is achieved through aadded to the end of frame.	header	trailer	adress
The datalink layer is responsible for moving.....from one hop to next	packets	frames	signals
In a single_bit error,how many bits in a data unit are changed	one	two	four
In a burst error,how many bits in a data unit are changed	less than 2	2 or more than 2	2
The length of the burst error is measured from	first bit to last bit	first corrupted bit to last corrupted bit	two
Single bit error will least occur in.....data transmissions	serial	parallel	synchronous
To detect errors or correct errors,we need to send with data	address	frames	extra bits
Which of the following best describes a single bit error	a single bit is inverted	a single bit is inverted per data unit	a single bit is inverted per transmission
In block coding,we divide our message intp blocks,each of k bits,called	dataword	codeword	integers

In block coding,the length of the block is	k	r	k+r
Block coding can detect onlyerror	single	burst	multiple
We needredundant bits for error correction than for error detection	less	more	equal
The corresponding codeword for the dataword 01 is.....	011	000	101
The hamming distance can easily be found if we apply the operation	XOR	OR	AND
The hamming distance is the smallest hamming distance between all possible pairs in a set of words	minimum	maximum	equal
The hamming distance d(000,111) is	1	0	2
To guarantee correction of upto t errors in all cases,the minimum hamming distance in a block code must be	$d(\min)=2t+1$	$d(\min)=2t-1$	$d(\min)=2t$
To guarantee correction of upto s errors in all cases,the minimum hamming distance in a block code must be	$d(\min)=s-1$	$d(\min)=s+1$	$d(\min)=s$
A simple_parity check code is a single bit error detecting code in which n=..... with $d(\min)=2$	K	$K*1$	K-1
The codeword corresponding to the dataword 1111 is	11110	11111	11101
A simple_parity check code can detect an Number of errors	odd	even	prime
The hamming code is a method of	error detection	error correcton	error encapsulation
To make the hamming code respond to a burst error of size N,we need to make codewords of our frame	N+1	N-1	N
CRC is used in network such as	WAN	LAN and WAN	LAN

In CRC there is no error if the remainder at the receiver is.....	equal to the remainder at the sender	all 0's	non zero
At the CRC checker,.....means that the data unit is damaged.	string of 0's	string of 1's	a string of alternating 1's and 0's
..... Is a regulation of data transmission so that the receiver buffer do not become overwhelmed	flow control	error control	access control
.....in the datalink layer separates a message from one source to a destination or from other message to other destinations	packets	address	framing
.....is the process of adding 1 extra byte whenever there is a flag or escape character in text	byte stuffing	redundancy	bit stuffing
.....is the process of adding 1 extra 0 whenever five consecutive 1's follows a 0 in the data.	byte stuffing	redundancy	bit stuffing
.....in the data link layer is based on automatic repeat request, which is the retransmission of data	error control	flow control	access control
At any time an error is detected in an exchange specified frames are retransmitted and process is called....	ARQ	ACK	NAK
The datalink layer at the sender side gets data from its.....layer	network	physical	application
ARQ stands for	acknowledge repeat request	automatic repeat request	automatic repeat quantisation
Which of the following is a data link layer function	line discipline	error control	flow control

In protocols the flow and error control information such as ACK and NAK is included in the data frames in a technique called.....	stop and wait	go_back	A and B
In stop and wait ARQ ,the sequence of numbers is based on.....	modulo-2-arithmetic	modulo-12-arithmetic	modulo-N-arithmetic
Error correction inis done by keeping a copy of the send frames and retransmitting of the frame when time expires	stop and wait ARQ	ARQ	ACK
In the Go_Back N protocol,the sequence numbers are modulo.....	2^m	2^{m-1}	2^{m+1}
In sliding window ,the range which is the concern of the sender is called.....	send sliding window	receive sliding window	piggybacking
Piggybacking is used to improve the efficiency of theprotocols.	bidirectional	unidirectional	multidirectional
The send window can slideslots when a valid acknowledgment arrive	one or more	one	two
The upper sublayer that is responsible for flow and error control is called.....control	logical	media access	logical and physical
The MAC(media access control)sublayer co-ordinates the datalink task within a specified.....	LAN	MAN	WAN
The lower sublayer that is responsible for multiple access resolution is calledcontrol	Logical	media access	logical and physical
In the sliding window method or flow control several frame can be beat a time	transit	received	logical and physical
The sliding window of the sender expands to thewhen acknowledgement are received	left	middle	right

Error detecting codes requirenuber of redundant bits.	less	equal	more
The datalink layer transforms thea raw transmission facility to a reliable link and is responsible for node_to_node delivery.	datalink	physical	network
Datalink layer divided into functionality oriented sublayer.	one	zero	two
The send window in Go_Back N maximum size can be	2^m	2^{m+1}	2
In stop and wait ARQ and Go_Back_N ARQ,the size of the send window is.....	0	3	1
The relationship between m and n in hamming code is	$n=2m-1$	$n=m$	$n=m-1$
A simple parity_check code is a single_bit error detecting code in which $n=k+1$ with $d_{\min}$	3	1	0
.....mechanism of datalink layer is achieved through added to the trailer added to the end of frame.	ARQ	ARC	Error control
In,we divide our message into blocks	convolution coding	block coding	linear coding
Thelayer at the sender site gets data from its network layer.	physical	datalink	application
In theprotocol,the sequence numbers are modulo 2^m	Go_Back N	Simplest	Stop and wait

ct 1956)
INDIA

opt4	Answer
transport	physical
physical	datalink
none of the above	flow control
frames	trailer
message	frames
five	one
3	2 or more than 2
three	first corrupted bit to last corrupted bit
asynchronous	serial
packets	extra bits
any of the above	a single bit is inverted per data unit
decimal	dataword

k-r	k+r
type	single
less than or equal to	more
110	011
NAND	XOR
not equal	minimum
3	2
$d(\min)=t+1$	$d(\min)=2t+1$
$d(\min)=0$	$d(\min)=s+1$
K+1	K+1
11011	11110
non- prime	odd
error detection and correction	error correcton
0	N
MAN	LAN and WAN

the quotient at the sender	all 0's
a non-zero remainder	a non-zero remainder
connection control	flow control
switching	framing
character stuffing	byte stuffing
character stuffing	bit stuffing
connection control	error control
SEL	ARQ
transport	network
automatic retransmission request	automatic repeat quantisation
all the above	all the above

piggybacking	piggybacking
modulo-1- arithmetic	modulo-2-arithmetic
NAQ	stop and wait ARQ
2	2m
sliding	send sliding window
omnidirectiona 1	bidirectional
two or more	one or more
physical	logical
LAN and MAN	LAN
physical	media access
physical	transit
top	right

less than or equal to	more
transport	physical
three	two
2^{m-1}	2^{m-1}
2	1
$n=2m+1$	$n=2m-1$
2	2
Flow control	Error control
A and C	block coding
transport	datalink
ARQ	Go_Back N

UNIT IV

SYLLABUS

UNIT-IV

Network Layer: IPv4 addresses – Internetworking – IPv4 – Delivery and Forwarding – Unicast Routing Protocols.

Transport Layer: Process to Process Delivery – User Datagram Protocol – Transmission Control Protocol.

TEXT BOOK

1. Behrouz A. Forouzan. 2006. Data Communication and Networking, 4th Edition, McGraw Hill, New Delhi.

UNIT IV

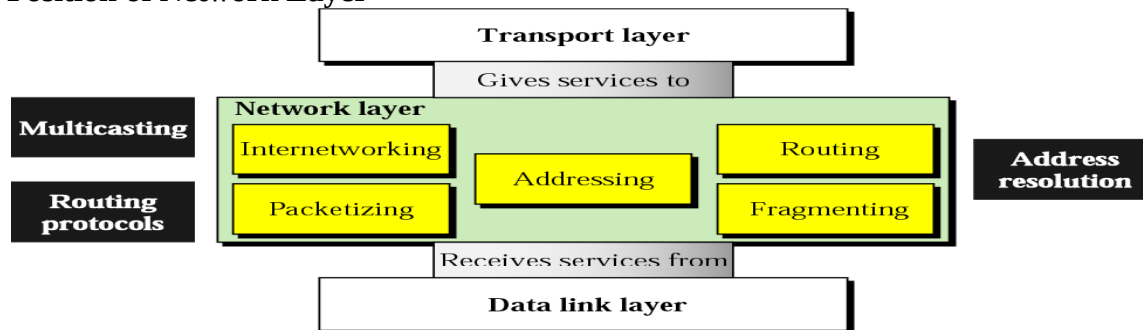
NETWORK LAYER

Network Layer

Network layer functions

- Deliver packets from sending to receiving hosts
- network layer protocols in *every* host, router three important functions:
- *path determination*: route taken by packets from source to dest. *Routing algorithms*
- *forwarding*: move packets from router's input to appropriate router output
- *call setup*: some network architectures require router call setup along path before data flows

Position of Network Layer



Internetworking

— Logically connecting heterogeneous networks to look like single network to upper transport and application layers.

• Addressing

— Each device (a computer or a router) over the Internet must have unique and universally accepted address.

• Routing

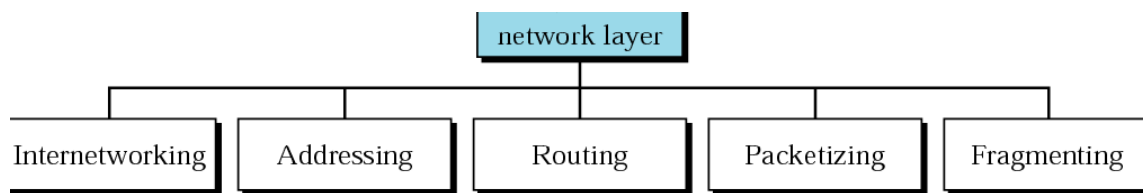
— Packet cannot choose its route to the destination. The routers connecting LANs and WANs make this decision.

• Packetizing

— The network layer encapsulates datagram/segments received from upper layers and makes packets out of them.

• Fragmenting

- Each router de-encapsulates the IP datagram from the received frame, processes it and encapsulates it into another frame.



IPv4 ADDRESSES

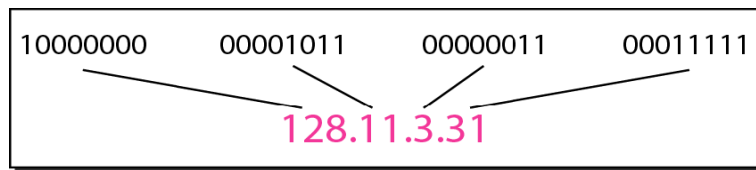
An IPv4 address is a 32-bit address that uniquely and universally defines the connection of a device (for example, a computer or a router) to the Internet.

Address Space

- An IPv4 address is 32 bits long.

Note

- The IPv4 addresses are unique and universal.
- Dotted-decimal notation and binary notation for an IPv4 address
 - Binary Notation
 - Dotted-Decimal Notation
- Identifier used in network layer to identify each device connected to the Internet
- 32-bit binary address that uniquely and universally defines the connection of a host or a router to the Internet.
- In Internet, no two devices can have the same IP
- For readability, we divide the IP address into 4 bytes.
- Dotted-decimal notation: Each byte is separated by dots.



Example

Change the following IPv4 addresses from binary notation to dotted-decimal notation.

- a. 10000001 00001011 00001011 11101111
- b. 11000001 10000011 00011011 11111111

Solution

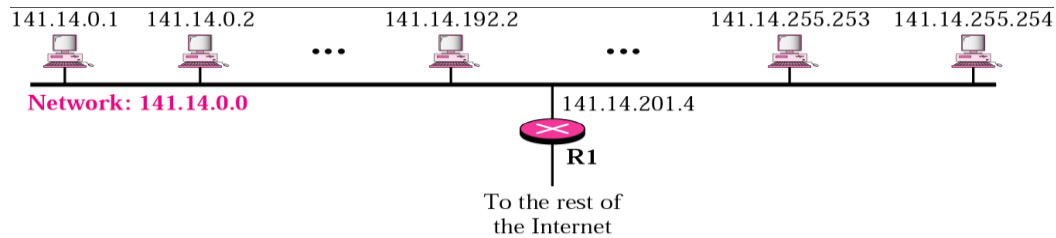
We replace each group of 8 bits with its equivalent decimal number and add dots for separation.

- a. 129.11.11.239
- b. 193.131.27.255

Levels of hierarchy

- Levels of Hierarchy
 - To reach a host on the Internet, we must first reach the network by using the first portion of the address (netid)

- Then we must reach the host itself by using the second portion (hostid)
- IP addresses are designed with two levels of hierarchy.



Delivery and Forwarding

Subnetting

- Sub-netting

—We can divide a network into sub-networks while making the world knows only the main network.

—In sub-netting, a network is divided into several smaller groups with each sub-network (or subnet) having its own sub-network address.

Mask

A router routes the packet based on network address and subnetwork address.

- A router inside a network routes based on subnetwork address but a router outside a network routes based on network address.
- Router uses the 32-bit mask to identify the network address.
- Routers outside an organization use a default mask; the routers inside an organization use a subnet mask

Classless Addressing

- A range of addresses meant a block of addresses in class A, B, or C.
- What about a small business that needed only 16 addresses? Or a household that needed only two addresses?
 - ISPs provide IP; people connect via dial-up modem, DSL, or cable modem to the ISP.
 - Variable-length blocks: No class boundaries.
 - Mask: Provide a block, it is given the first address and mask.
 - Subnetting
 - Classless Inter Domain Routing (CIDR)

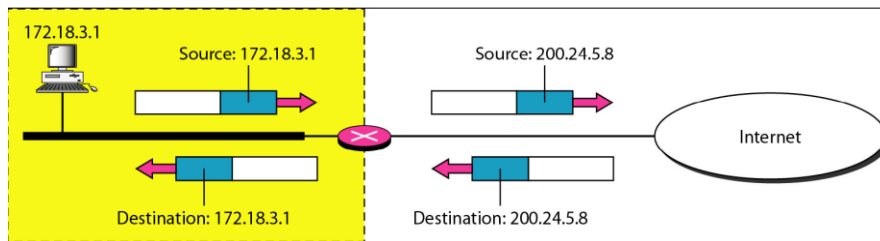
Dynamic Address Configuration

- Each computer has IP address, subnet mask, IP address of a router, IP address of a name server; This information is usually stored in a configuration file and accessed by the computer during the bootstrap (boot)process.
 - Dynamic Host Configuration Protocol (DHCP) is a protocol designed to provide the information dynamically (based on demand).
 - DHCP is a client-server program.
 - When a DHCP client requests a temporary IP address, the DHCP server goes to the pool of available (unused) IP addresses and assigns an IP address for a negotiable period of time.
- A Network Address Translation (NAT) implementation

- NAT enables a user to have a large set of addresses internally and one address, or a small set of addresses, externally.

Address Translation

- All the outgoing packets go through the NAT router, which replaces the source address in the packet with the global NAT address.
- All incoming packets also pass through the NAT router, which replaces the destination address in the packet (the NAT router global address) with the appropriate private address.



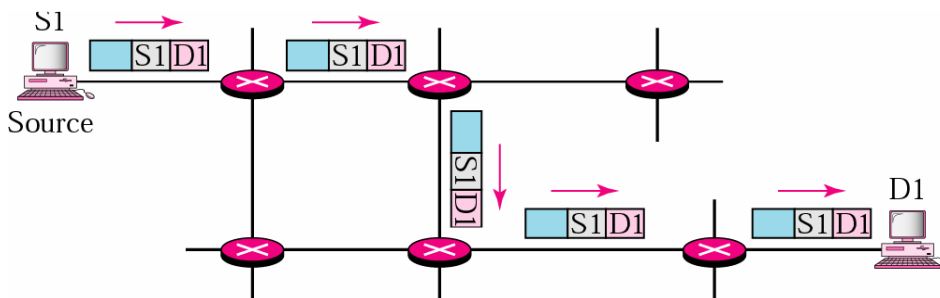
Routing Protocols

There are two types of routing protocols they are:

1 Unicast routing

2 Multicast routing

UNICAST ROUTING



- In unicast routing, there is only one source and only one destination.
- When a router receives a packet, it forwards the packet through only one of its ports (the one belonging to the optimum path) as defined in routing table. It discards the packet, if there is no route.

Metric of different protocols

- Metric is the cost assigned for passing through a network.
 - The total metric of a particular router is equal to the sum of the metrics of networks that comprise the route.
 - A router chooses the route with smallest metric.
- RIP (Routing Information Protocol): Cost of passing each network is same; it is one hop count.
 - If a packet passes through 10 networks to reach the destination, the total cost is 10 hop counts.
- OSPF (Open Shortest Path First): Administrator can assign cost for passing a network based on type of service required.

- OSPF allows each router to have more than one routing table based on required type of service.
- Maximum throughput, minimum delay
- BGP (Border Gateway Protocol): Criterion is the policy, which is set by the administrator.

Interior and Exterior routing

- Autonomous System: Group of networks and routers under the authority of a single administration.
- Routers inside an autonomous system is referred to as interior routing.
- Routing between autonomous systems is referred to as exterior routing.

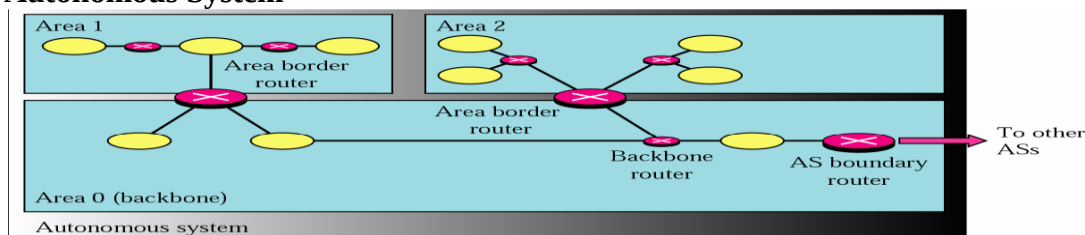
Routing Information Protocol (RIP)

- RIP is based on Distance vector routing.
- Distance vector routing
 - Sharing knowledge about the entire autonomous system: Each router periodically shares its knowledge about the entire autonomous system with its neighbours.
 - Sharing only with neighbours through all its interfaces.
 - Sharing at regular intervals: 30 seconds.
- Routing table
 - Has one entry for each destination network of which the router is aware.
 - Each entry has destination network address, the shortest distance to reach the destination in hop count, and next router to which the packet should be delivered to reach its final destination.
 - Hop count is the number of networks that a packet encounters to reach its final destination.

OSPF

- Open Shortest Path First
- Special routers called autonomous system boundary routers are responsible for dissipating information about other autonomous systems into the current system.
- OSPF divides an autonomous system into areas.

Autonomous System



Areas in an Autonomous System

- Area is a collection of networks, hosts, and routers all contained within an autonomous system.
- Routers inside an area flood the area with routing information.
- Area border routers: Summarize the information about the area and send it to other routers.
- Backbone area [Primary area]: All the areas inside an autonomous system must be connected to the backbone. Routers in this area are called as backbone routers. This area identification number is 0.

Link state routing table for router A

Network	Cost	Next Router	Other Information
N1	5	C	
N2	7	D	
N3	10	B	
N4	11	D	
N5	15	C	

BGP

- Border Gateway Protocol
- Inter-autonomous system routing protocol.
- BGP is based on a routing method called path vector routing.

Why D.V. and L.S. are not good enough?

- Distance Vector routing
 - Sometimes we don't want the route with smallest hop count as the preferred route [like, avoiding non-secure routes].
 - D.V routing information provides only the hop count and not the path that leads to that destination.
 - A router that receives a distance vector advertisement packet may be fooled if the shortest path is actually calculated through the receiving router itself.
- Link State routing
 - Internet is too big for this routing method
 - To use link state routing for the whole internet would require each router to have a huge link state database.
 - It would also take a long time for each router to calculate its routing table using the Dijkstra algorithm
- Path Vector routing
 - Each entry in the routing table contains the destination network, the next router, and the path to reach the destination.
 - The path is usually defined as an ordered list of autonomous systems that a packet should travel through to reach the destination.

Transport Layer

The transport layer is responsible for the delivery of a message from one process to another.

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

A connectionless transport layer treats each segment as an independent packet and delivers it to the transport layer at the destination machine. A connection-oriented transport layer makes a connection with the transport layer at the destination machine first before delivering the packets. After all the data is transferred, the connection is terminated. In the transport layer, a message is normally divided into transmittable segments. A connectionless protocol, such as UDP, treats

each segment separately. A connection-oriented protocol, such as TCP and SCTP, creates a relationship between the segments using sequence numbers

- understand principles behind transport layer services:
 - multiplexing/demultiplexing
 - reliable data transfer
 - flow control
 - congestion control
- learn about transport layer protocols in the Internet:
 - UDP: connectionless transport
 - TCP: connection-oriented transport
 - TCP congestion control

Transport-layer services

- Multiplexing and demultiplexing
- Connectionless transport: UDP
- Principles of reliable data transfer
- Principles of reliable data transfer
 - segment structure
 - reliable data transfer
 - flow control
 - connection management
- Principles of congestion control
- TCP congestion control

Transport services and protocols

- provide *logical communication* between app processes running on different hosts
- transport protocols run in end systems
 - send side: breaks app messages into segments, passes to network layer
 - rcv side: reassembles segments into messages, passes to app layer

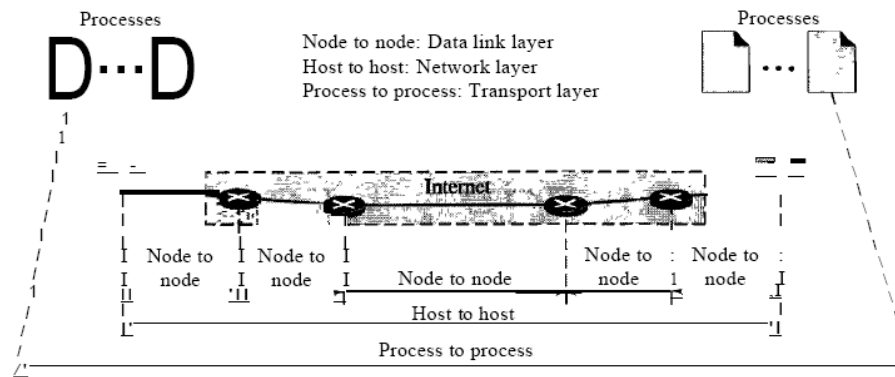
Transport vs. network layer

- *network layer*: logical communication between hosts
 - Point-to-point
- *transport layer*: logical communication between processes
 - relies on and enhances, network layer services also called “End-to-End”

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-the delivery of a packet, part of a message, from one process to another. Two processes communicate in a client/server relationship.



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, we 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. The destination port number is needed for delivery; the source port number is needed for the reply.

In the Internet model, the 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.

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 passes the packet to the network layer.

Connectionless Versus Connection-Oriented Service

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

Connectionless Service

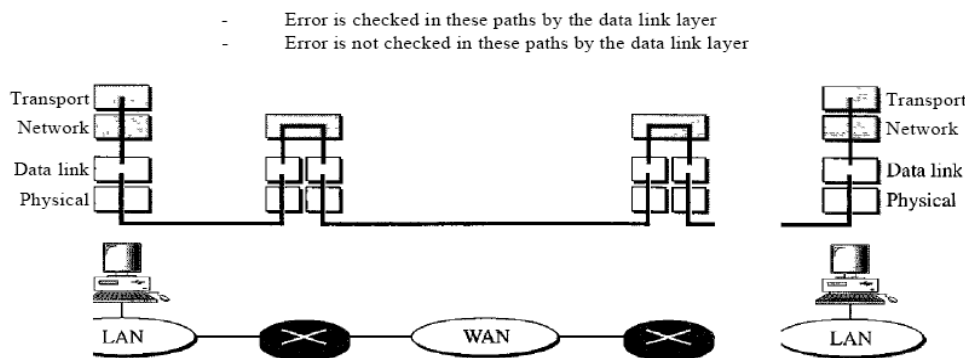
In a connectionless service, the packets are sent from one party to another with no need for connection establishment or connection release. The packets are not numbered; they may be delayed or lost or may arrive out of sequence. There is no acknowledgment either. We will see shortly that one of the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



USER DATAGRAM PROTOCOL (UDP)

The User Datagram Protocol (UDP) is called a connectionless, unreliable transport protocol. It does not add anything to the services of IP except to provide process-to-process communication instead of host-to-host communication. Also, it performs very limited error checking.

Well-Known Ports for UDP

Table 23.1 shows some well-known port numbers used by UDP. Some port numbers can be used by both UDP and TCP.

<i>Port</i>	<i>Protocol</i>	<i>Description</i>
7	Echo	Echoes a received datagram back to the sender
9	Discard	Discards any datagram that is received
11	Users	Active users

<i>Port</i>	<i>Protocol</i>	<i>Description</i>
13	Daytime	Returns the date and the time
17	Quote	Returns a quote of the day
19	Chargen	Returns a string of characters
53	Nameserver	Domain Name Service
67	BOOTPs	Server port to download bootstrap information
68	BOOTPc	Client port to download bootstrap information
69	TFTP	Trivial File Transfer Protocol
III	RPC	Remote Procedure Call
123	NTP	Network Time Protocol
161	SNMP	Simple Network Management Protocol
162	SNMP	Simple Network Management Protocol (trap)

User Datagram

UDP packets, called user datagrams, have a fixed-size header of 8 bytes. Figure below shows the format of a user datagram.

The fields are as follows:

- o Source port number. This is the port number used by the process running on the source host. It is 16 bits long, which means that the port number can range from 0 to 65,535. If the source host is the client (a client sending a request), the port number, in most cases, is an ephemeral port number requested by the process and chosen by the UDP software running on the source host. If the source host is the server (a server sending a response), the port number, in most cases, is a well-known port number.

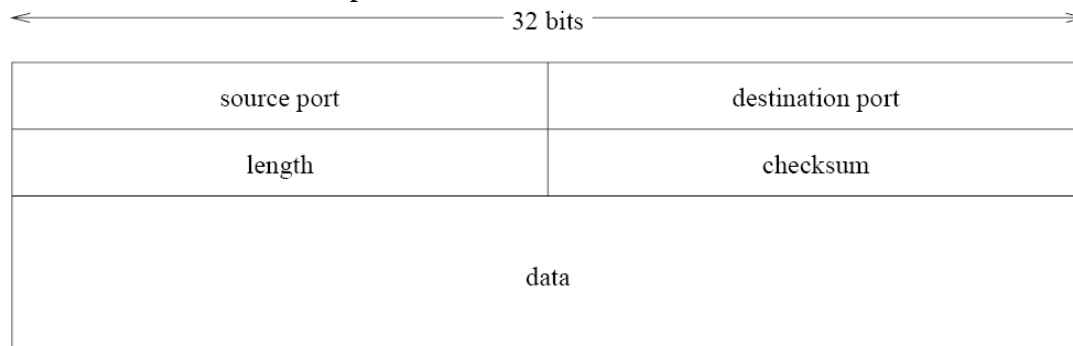


Figure 15.2. UDP header format

Destination port number. This is the port number used by the process running on the destination host. It is also 16 bits long. If the destination host is the server (a client sending a request), the port number, in most cases, is a well-known port number. If the destination host is the client (a server sending a response), the port number, in most cases, is an ephemeral port number. In this case, the server copies the ephemeral port number it has received in the request packet.

Length. This is a 16-bit field that defines the total length of the user datagram, header plus data. The 16 bits can define a total length of 0 to 65,535 bytes. However, the total length needs to be much less because a UDP user datagram is stored in an IP datagram with a total length of 65,535 bytes. The length field in a UDP user datagram is actually not necessary. A user datagram is encapsulated in an IP datagram. There is a field in the IP datagram that defines the total length. There is another field in the IP datagram that defines the length of the header. So if we subtract the value of the second field from the first, we can deduce the length of a UDP datagram that is encapsulated in an IP datagram.

UDP length = IP length - IP header's length

However, the designers of the UDP protocol felt that it was more efficient for the destination UDP to calculate the length of the data from the information provided in the UDP user datagram rather than ask the IP software to supply this information.

We should remember that when the IP software delivers the UDP user datagram to the UDP layer, it has already dropped the IP header.

o Checksum. This field is used to detect errors over the entire user datagram (header plus data). The checksum is discussed next. Checksum. We have also shown how to calculate the checksum for the IP and ICMP packet. We now show how this is done for UDP.

UDP Operation

UDP uses concepts common to the transport layer. These concepts will be discussed here briefly, and then expanded in the next section on the TCP protocol.

Connectionless Services

As mentioned previously, UDP provides a connectionless service. This means that each user datagram sent by UDP is an independent datagram. There is no relationship between the different user datagrams even if they are coming from the same source process and going to the same destination program. The user datagrams are not numbered. Also, there is no connection establishment and no connection termination, as is the case for TCP. This means that each user datagram can travel on a different path. One of the ramifications of being connectionless is that the process that uses UDP cannot send a stream of data to UDP and expect UDP to chop them into different related user datagrams. Instead each request must be small enough to fit into one user datagram. Only those processes sending short messages should use UDP.

Flow and Error Control

UDP is a very simple, unreliable transport protocol. There is no flow control and hence no window mechanism. The receiver may overflow with incoming messages. There is no error control mechanism in UDP except for the checksum. This means that the sender does not know if a message has been lost or duplicated. When the receiver detects an error through the checksum, the user datagram is silently discarded. The lack of flow control and error control means that the process using UDP should provide these mechanisms.

Encapsulation and Decapsulation

To send a message from one process to another, the UDP protocol encapsulates and decapsulates messages in an IP datagram.

Transmission Control Protocol (TCP)

Transmission Control Protocol (TCP). TCP, like UDP, is a process-to-process (program-to-program) protocol. TCP, therefore, like UDP, uses port numbers. Unlike UDP, TCP is a connection oriented protocol; it creates a virtual connection between two TCPs to send data. In addition, TCP uses flow and error control mechanisms at the transport level.

In brief, TCP is called a *connection-oriented, reliable* transport protocol. It adds connection-oriented and reliability features to the services of IP.

TCP Services

Before we discuss TCP in detail, let us explain the services offered by TCP to the processes at the application layer.

Process-to-Process Communication

Like UDP, TCP provides process-to-process communication using port numbers. able below lists some well-known port numbers used by TCP. The Transmission Control Protocol (TCP) is one of the main transport layer protocols used with IP. It is a connection oriented protocol based on the connectionless IP protocol. Because it is the lowest layer which has end-to-end communication, it needs to handle things such as lost packets. In this respect it is similar to the data-link layer which must handle errors on an individual link.

Port	Protocol	Description
7	Echo	Echoes a received datagram back to the sender
9	Discard	Discards any datagram that is received
11	Users	Active users
13	Daytime	Returns the date and the time
17	Quote	Returns a quote of the day
19	Chargen	Returns a string of characters
20	FIP, Data	File Transfer Protocol (data connection)
21	FIP, Control	File Transfer Protocol (control connection)
23	TELNET	Tenninal Network
25	SMTP	Simple Mail Transfer Protocol
53	DNS	Domain Name Server
67	BOOTP	Bootstrap Protocol
79	Finger	Finger
80	HTTP	Hypertext Transfer Protocol
111	RPC	Remote Procedure Call

Stream Delivery Service

TCP, unlike UDP, is a stream-oriented protocol. In UDP, a process (an application program) sends messages, with predefined boundaries, to UDP for delivery. UDP adds its own header to each of these messages and delivers them to IP for transmission. Each message from the process is called a user datagram and becomes, eventually, one IP datagram. Neither IP nor UDP recognizes any relationship between the datagrams. TCP, on the other hand, allows the

sending process to deliver data as a stream of bytes and allows the receiving process to obtain data as a stream of bytes. TCP creates an environment in which the two processes seem to be connected by an imaginary "tube" that carries their data across the Internet.



Full-Duplex Communication

TCP offers full-duplex service, in which data can flow in both directions at the same time. Each TCP then has a sending and receiving buffer, and segments move in both directions.

Connection-Oriented Service

TCP, unlike UDP, is a connection-oriented protocol. When a process at site A wants to send and receive data from another process at site B, the following occurs:

1. The two TCPs establish a connection between them.
2. Data are exchanged in both directions.
3. The connection is terminated.

Note that this is a virtual connection, not a physical connection. The TCP segment is encapsulated in an IP datagram and can be sent out of order, or lost, or corrupted, and then resent. Each may use a different path to reach the destination. There is no physical connection.

TCP creates a stream-oriented environment in which it accepts the responsibility of delivering the bytes in order to the other site. The situation is similar to creating a bridge that spans multiple islands and passing all the bytes from one island to another in one single connection.

Reliable Service

TCP is a reliable transport protocol. It uses an acknowledgment mechanism to check the safe and sound arrival of data. We will discuss this feature further in the section on error control.

TCP Features

Numbering System

Although the TCP software keeps track of the segments being transmitted or received, there is no field for a segment number value in the segment header. Instead, there are two fields called the sequence number and the acknowledgment number. These two fields refer to the byte number and not the segment number. Byte Number TCP numbers all data bytes that are transmitted in a connection. Numbering is independent in each direction. When TCP receives bytes of data from a process, it stores them in the sending buffer and numbers them. The numbering does not necessarily start from 0. Instead, TCP generates a random number between 0 and $2^{32} - 1$ for the number of the first byte. For example, if the random number happens to be 1057 and the total data to be sent are 6000 bytes, the bytes are numbered from 1057 to 7056. We will see that byte numbering is used for flow and error control. The bytes of data being transferred in each connection are numbered by TCP. The numbering starts with a randomly generated number. Sequence Number After the bytes have been numbered, TCP assigns a sequence number to each

segment that is being sent. The sequence number for each segment is the number of the first byte carried in that segment.

Flow Control

TCP, unlike UDP, provides *flow control*. The receiver of the data controls the amount of data that are to be sent by the sender. This is done to prevent the receiver from being overwhelmed with data. The numbering system allows TCP to use a byte-oriented flow control.

Error Control

To provide reliable service, TCP implements an error control mechanism. Although error control considers a segment as the unit of data for error detection (loss or corrupted segments), error control is byte-oriented.

Congestion Control

TCP, unlike UDP, takes into account congestion in the network. The amount of data sent by a sender is not only controlled by the receiver (flow control), but is also determined by the level of congestion in the network.

Segment

Before we discuss TCP in greater detail, let us discuss the TCP packets themselves. A packet in TCP is called a segment.

The format of a TCP segment header is shown in figure below.

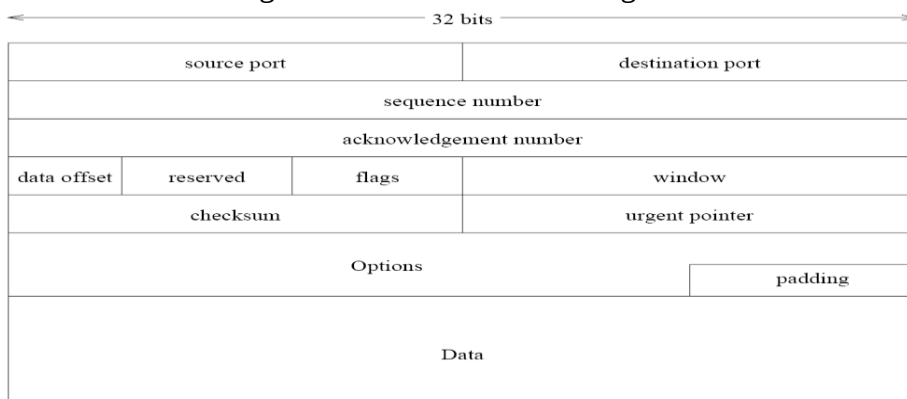


Figure 15.1. TCP header format

source port

All of the address fields in the lower layer protocols are only concerned with getting the packet to the correct host. Often, however, we want to have multiple connections between two hosts. The source port is simply the number of the outgoing connection from the source host.

destination port

Similarly, this is the number of the incoming connection on the destination host. There must be a program on the destination host which has somehow told the networking system on that host that it will accept packets destined for this port number. Standard system services such as SMTP, NNTP and NTP (all described in later chapters) have well known standard port numbers. Thus to connect to the SMTP port on a particular host (in order to transmit some email) a TCP connection would be set up with the correct destination port number (25). The source port number does not matter except that it should be unique on the sending machine so that replies can be received correctly.

sequence number

This is the sequence number of this packet. It differs from the usual data-link layer sequence number in that it is in fact the sequence number of the first byte of information and is incremented by the number of bytes in this packet for the next message. In other words, it counts the number of bytes transmitted rather than the number of packets.

acknowledgement number .This is the sequence number of the last byte being acknowledged. This is a piggy-backed acknowledgement. This field is the offset in the packet of the beginning of the data field. (In other words it is the length of the header.)

flags

This field contains several flags relating to the transfer of the packet. We will not consider them further here.

window

This field is used in conjunction with the acknowledgement number field. TCP uses a sliding window protocol with a variable window size (often depending on the amount of buffer space available. This field contains the number of bytes which the host is willing to accept from the remote host.

POSSIBLE QUESTIONS

UNIT- 4

Descriptive Type Questions:

- 1) What is NAT? How can NAT help in address
- 2) Differentiate Intra- and Interdomain routing.
- 3) Explain client/server paradigm
- 4) Describe delivery with a diagram
- 5) Explain intra –and inter domain routing
- 6) Briefly discuss about internetworking
- 7) Define address space in IPV4
- 8) What is the need for network layer?



KARPAGAM ACADEMY OF HIGHER EDUCATION
Deemed to be University Established Under Section 3 of UGC Act 1956)

Ilachi Main Road, Eachanari Post, Coimbatore – 641 021. INDIA

DEPARTMENT OF COMPUTER SCIENCE

DATA COMMUNICATION NETWORKS (15CSU502)

ONLINE EXAM QUESTION BANK

UNIT-4

Question	Option 1	Option 2	Option 3	Option 4
Internet at the network layer is a _____ network.	packet-switched	.LAN	connection	connetionless
Internet has chosen the datagram approach to _____ in network layer	routers	packets	switching	protocol
Internet is made of so many _____ networks.	homogenous	hetrogenous	MAN	multipoint
Communication at network layer in the internet is _____.	connectionless	point-to-point	connection oriented	packet-switched
What is the abbreviation for IPV4 _____.	Inter Protocol Versus 4	Inter Position Version 4	Internet protocol version 4	Internet Position Versus 4
IPV4 provides the term 'best-effort' means that _____.	no error control	error control	error detection	datagram
Packets in the IPV4 layer are called _____.	frames	datagroup	switching	datagrams
A datagram is a variable length packet consisting of _____ parts.	one	six	two	three
The total length field defines the total length of the datagram including _____.	footer	header	flags	frames
Abbravation for MTU _____.	Minimum Transfer Unit	Maximum Transfer Unit	Maximum Travel Unit	Minimum Travel Unit
_____ in the IPV4 packet covers only header,not the data.	Check subtract	Check sum	options	Check product
Options can be used for network testings and _____.	checking	packets	types	debugging
A no-operation option is a _____ byte used as a filler between option.	three	six	one	four
_____ can only used as the last option.	end-of-option	first-of-option	options	no options

Record route can list up to_____ router address.	fifteen	sixty	nine	ten
_____route has less rigid.	loose source	strict source	no route	record
_____is expressed in millisecond,from midnight.	time stand	time stamp	time shot	time start
In packet format,the extension headers and data from the upper layer conains upto_____bytes of information.	65033	65535	65536	65035
Base header with_____ fields.	eight	ten	five	six
The 4bit field defines the_____ number of the IP.	versus	header	.footer	version
Delivery has_____types.	3	4	5	2
Source and destination of the packet are located on the same physical network called_____.	Indirect delivery	Inward delivery	Direct delivery	Outward delivery
One technique to reduce the content of a routing table is_____.	before-hop	next-hop	first hop	last hop
The Routing table holds only the address of the next hop_____.	next hop	route method	network method	host method
A second technique to reduce the routing table_____.	next hop	default	forward	network specific
All hosts connected to the same network as one single entity_____.	route	next-hop	host specific	network specific
In classless addressing,atleast_____col umns in a routing table.	5	6	3	4
In an address aggregation,the network for each organization is_____.	independent	dependent	network specific	host specific
The routing table can be either_____.	static	static and dynamic	static or dynamic	dynamic
A static routing table can be used in a_____internet.	big	small	multi	LAN
Dynamic routing protocols are_____.	RIP	OSPF	BGP	RIP OSPF and BGP

The one of the flag is not present,the router is down_____.	G	U	H	D
Flag D means_____.	added by direction	added	added by redirection	subtracted by direction
Routing inside an autonomous system_____	Intra	Inter	Inside	Outside
Abbreviation for BGP_____.	Border Gateway Process	Bit Gateway Process	Border Gateway Protocol	Byte Gateway Protocol
A node sends its routing table,at every_____in a periodic update.	33s	30s	31s	35s
_____algorithm creates a shortest path tree from a graph.	data	dakstra	define	dijkstra
An area is a collection of_____.	networks	hosts	route	networks, hosts and route
_____link is a network and is connected to only one router.	stub	point-to-point	transient	route
Multicasting of the relationship is_____.	one-to-one	many-to-one	one-to-many	many-to-many
_____layer is responsible for process-to-process delivery.	transport	physical	application	network
Internet has decided to use universal port numbers for severs called_____.	well-unknown port	well-known port	well-known protocol	well-unknown process
IANA has divided the port numbers into_____ranges.	six	four	five	three
_____a connection,is first established between the sender and receiver.	connection-oriented	connectionless	token	dialog
UDP is called_____.	connection-oriented	check point	token	connetionless
UDP length = IP length - _____.	IP length	IP breadth	IP header's length	IP header's breadth
UDP is a suitable transport protocol for_____.	unicasting	multicasting	nocasting	bicasting
TCP groups a number of bytes together into a packet called_____.	segment	encapsulation	datagram	data binding

The acknowledgement number is _____.	natural	whole	integers	cumulative
_____ flag is used to terminate the connection.	TER	FIN	URG	PSH
_____ protocol is used to remote procedure call.	DNS	PRC	RPC	RPCC
An ACK segment,if carrying _____ data consumes no sequence number.	no	2	3	5
In TCP,one end can stop sending data while still receiving data is _____.	full-close	full-open	half-close	half- open
The value of RTO is dynamic in TCP and is updated based on _____ segment.	RTO	RTT	ACK	ARQ

The block size of class C address is 255
 In classless addressing , the first address is 256
 _____ algorithm creates a shortest path data 128
 The last address in the block is 127
 An IPv4 address is 32 bits 128
 Addresses in class C are used by direction Bit Gateway Protocol D

ANSWER
packet-switched
protocol
heterogeneous
connectionless
Internet protocol version 4
no error control
datagrams
two
header
Maximum Transfer Unit
Checksum
debugging
one
end-of-option

nine
loose source
time stamp
65535
eight
version
2
Direct delivery
next-hop
route method
network specific
host specific
4
independent
static or dynamic
small
RIP OSPF and BGP

U
added by redirection
Inside
Border Gateway Protocol
30s
dijkstra
networks, hosts and route
stub
one-to-many
transport
well-known port
three
connection-oriented
connetionless
IP header's length
unicasting
segment

cumulative
FIN
RPC
no
half-close
RTT

255
first
dijkstra

1
32

D

UNIT V

SYLLABUS

UNIT-V

Transport Layer: Data Traffic – Congestion Control.

Application Layer: Electronic Mail – File Traffic – WWW and HTTP – Symmetric Key and Asymmetric Key Cryptography – Security Services – Message Integrity – Message Authentication – Digital Signature.

TEXT BOOK

1. Behrouz A. Forouzan. 2006. Data Communication and Networking, 4th Edition, McGraw Hill, New Delhi.

UNIT V

TRANSPORT LAYER

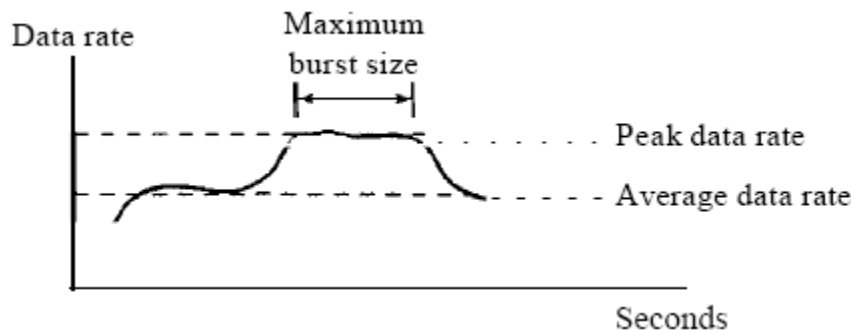
Transport Layer

DATA TRAFFIC

The main focus of congestion control and quality of service is data traffic. In congestion control we try to avoid traffic congestion. In quality of service, we try to create an appropriate environment for the traffic.

Traffic Descriptor

Traffic descriptors are qualitative values that represent a data flow. Figure shows a traffic flow with some of these 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 value 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 just 1 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 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 above Figure .

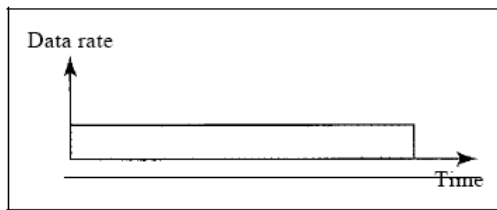
Constant Bit Rate

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.

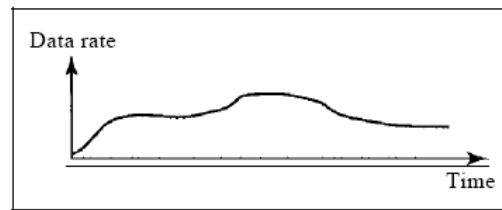
The maximum burst size is not applicable. This type of traffic is very easy for a network to handle since it is predictable. The network knows in advance how much bandwidth to allocate for this type of flow.

Variable Bit Rate

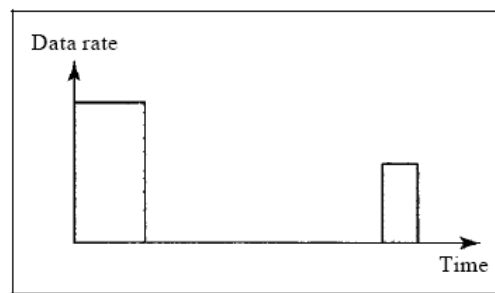
In the variable-bit-rate (VBR) category, the rate of the data flow changes in time, with the changes smooth instead of sudden and sharp. In this type of flow, the average data rate and the peak data rate are different. The maximum burst size is usually a small value.



a. Constant bit rate



b. Variable bit rate



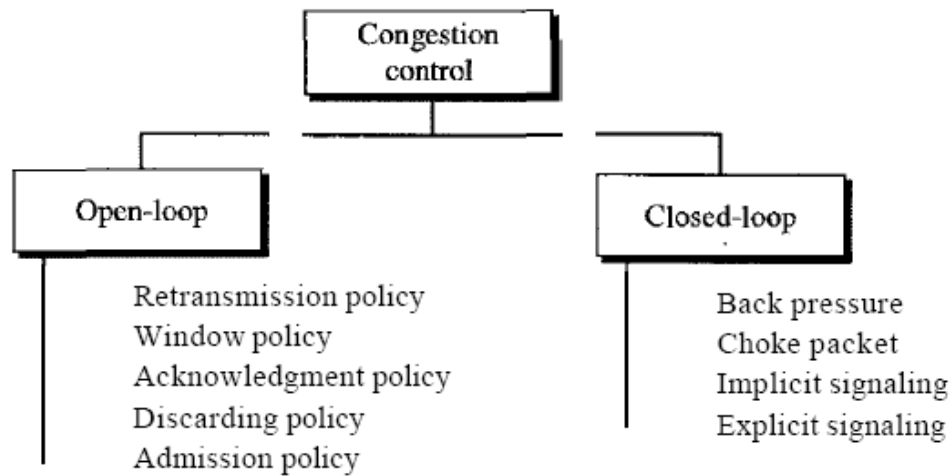
c. Bursty

Bursty

In the **bursty data** category, the data rate changes suddenly in a very short time. It may jump from zero, for example, to 1 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 bit 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.

CONGESTION CONTROL

Congestion control refers to techniques and mechanisms that can either prevent congestion, before it happens, or remove congestion, after it has happened.



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.

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.

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.

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.

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.

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 virtualcircuit connection if there is congestion in the network or if there is a possibility of future congestion.

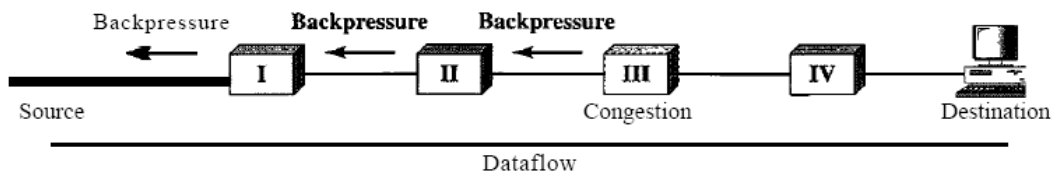
Closed-Loop Congestion Control

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.

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.



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 informs 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.

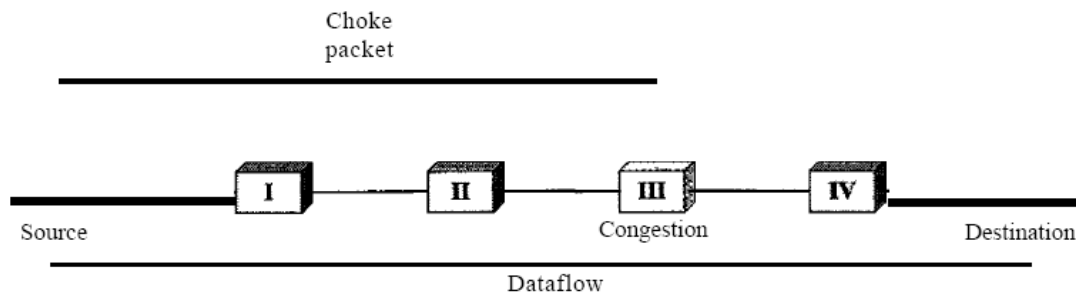
None of the virtual-circuit networks we studied in this book use backpressure. It was, however, implemented in the first virtual-circuit network, X.25. The technique cannot be implemented in a datagram network because in this type of network, a node (router) does not have the slightest knowledge of the upstream router.

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 traveled 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 sourcehost, using a source quench ICMP message. The warning message goes directly to the source station; the intermediate routers, and does not take any action.

This figure shows the idea of a choke packet.



Implicit Signaling

In implicit signaling, there is no communication between the congested node or nodes and the source. The source guesses that there is a 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.

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.

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. **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.

Application Layer

Electronic Mail

The most heavily used application in virtually any distributed system is electronic mail. The Simple Mail Transfer Protocol (SMTP) has always been the workhorse of the TCP/IP suite. However, SMTP has traditionally been limited to the delivery of simple text messages. In recent years, there has been a demand for the capability to deliver mail containing various types of data, including voice, images, and video clips. To satisfy this requirement, a new electronic mail standard, which builds on SMTP, has been defined: the Multi-Purpose Internet Mail Extension (MIME). In this section, we first examine SMTP, and then look at MIME.

SMTP

SMTP is the standard protocol for transferring mail between hosts in the TCP/IP suite; it is defined in RFC 821. Although messages transferred by SMTP usually follow the format defined in RFC 822, described later, SMTP is not concerned with the format or content of messages themselves, with two exceptions. This concept is often expressed by saying that SMTP uses information written on the *envelope* of the mail (message header), but does not look at the contents (message body) of the envelope. The two exceptions:

1. SMTP standardizes the message character set as 7-bit ASCII.
2. SMTP adds log information to the start of the delivered message that indicates the path the message took.
 - RFC 821
 - not concerned with format of messages or data
 - covered in RFC 822 (see later)
 - SMTP uses info written on envelope of mail
 - message header
 - does not look at contents
 - message body
 - except:
 - standardize message character set to 7 bit ASCII
 - add log info to start of message

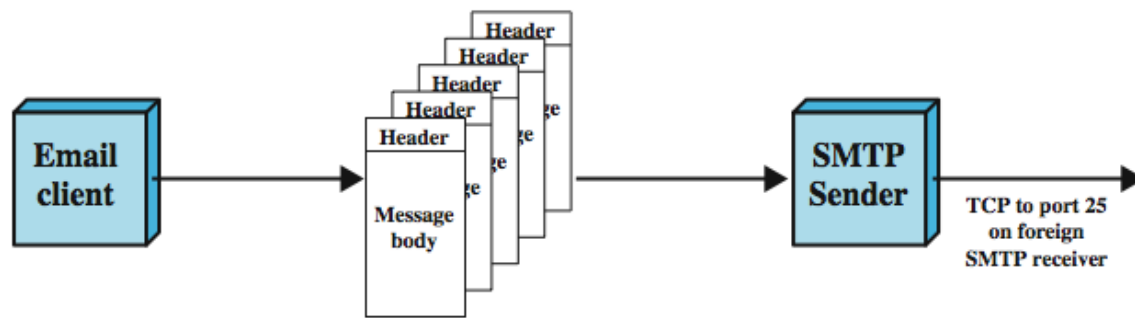
Basic Operation

To begin, mail is created by a user agent program in response to user input. Each created message consists of a header that includes the recipient's e-mail address and other information, and a body containing the message to be sent. These messages are then queued in some fashion and provided as input to an SMTP Sender program, which is typically an always-present server program on the host.

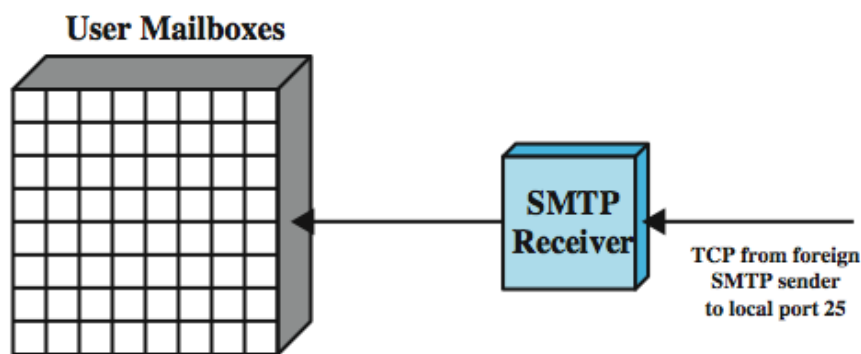
- email message is created by user agent program (mail client), and consists of:
 - header with recipient's address and other info
 - body containing user data
- messages queued and sent as input to SMTP sender program
 - typically a server process (daemon on UNIX)

SMTP Mail Flow

Figure below illustrates the overall flow of mail in a typical system. Although much of this activity is outside the scope of SMTP, the figure illustrates the context within which SMTP typically operates.



(a) Outgoing Mail



(b) Incoming Mail

Mail Message Contents

Although the structure of the outgoing mail queue will differ depending on the host's operating system, each queued message conceptually has two parts:

1. The message text, consisting of the RFC 822 header (constitutes the message envelope and includes an indication of the intended recipient or recipients), and the body of the message, composed by the user.
2. A list of mail destinations.

The list of mail destinations for the message is derived by the user agent from the RFC 822 message header. In some cases, the destination or destinations are literally specified in the message header. In other cases, the user agent may need to expand mailing list names, remove duplicates, and replace mnemonic names with actual mailbox names. If any blind carbon copies (BCCs) are indicated, the user agent needs to prepare messages that conform to this requirement. The basic idea is that the multiple formats and styles preferred by humans in the user interface are replaced by a standardized list suitable for the SMTP send program.

- each queued message has two parts
 - message text
 - RFC 822 header with envelope and list of recipients
 - message body, composed by user
 - list of mail destinations
 - derived by user agent from header
 - may be listed in header
 - may require expansion of mailing lists

- may need replacement of mnemonic names with mailbox names
- if BCCs indicated, user agent needs to prepare correct message format

Conversation

It is important to note that the SMTP protocol is limited to the conversation that takes place between the SMTP sender and the SMTP receiver. SMTP's main function is the transfer of messages, although there are some ancillary functions dealing with mail destination verification and handling. The rest of the mail-handling apparatus depicted in Figure below is beyond the scope of SMTP and may differ from one system to another.

SMTP System Overview

The operation of SMTP consists of a series of commands and responses exchanged between the SMTP sender and receiver. The initiative is with the SMTP sender, who establishes the TCP connection. Once the connection is established, the SMTP sender sends commands over the connection to the receiver. Each command generates exactly one reply from the SMTP receiver.

SMTP Commands

Each command consists of a single line of text, beginning with a four-letter command code followed in some cases by an argument field. Most replies are a single-line, although multiple-line replies are possible. The table lists first those commands that all receivers must be able to recognize. The other commands are optional and may be ignored by the receiver.

Connection Setup

Basic SMTP operation occurs in three phases: connection setup, exchange of one or more command-response pairs, and connection termination.

In the connection setup phase, an SMTP sender will attempt to set up a TCP connection with a target host when it has one or more mail messages to deliver to that host. The sequence is quite simple:

1. The sender opens a TCP connection with the receiver.
2. Once the connection is established, the receiver identifies itself with "220 Service Ready".
3. The sender identifies itself with the HELO command.
4. The receiver accepts the sender's identification with "250 OK".

If the mail service on the destination is unavailable, the destination host returns a "421 Service Not Available" reply in step 2 and the process is terminated.

Mail Transfer

Once a connection has been established, the SMTP sender may send one or more messages to the SMTP receiver. There are three logical phases to the transfer of a message:

1. A MAIL command identifies the originator of the message. The MAIL command gives the reverse path, which can be used to report errors. If the receiver is prepared to accept messages from this originator, it returns a "250 OK" reply. Otherwise the receiver returns a reply indicating failure to execute the command (codes 451, 452, 552) or an error in the command (codes 421, 500, 501).
2. One or more RCPT commands identify the recipients for this message. The RCPT command identifies an individual recipient of the mail data; multiple recipients are specified by multiple use of this command. A separate reply is returned for each RCPT command. The receiver can accept the destination with a 250 reply; or return an appropriate fail/error response.
3. A DATA command transfers the message text. The advantage of using a separate RCPT phase is that the sender will not send the message until it is assured that the receiver is prepared to

receive the message for at least one recipient, thereby avoiding the overhead of sending an entire message only to learn that the destination is unknown. Once the SMTP receiver has agreed to receive the mail message for at least one recipient, the SMTP sender uses the DATA command to initiate the transfer of the message. The end of the message is indicated by a line containing only a period.

Multipurpose Internet Mail Extension (MIME)

MIME is an extension to the RFC 822 framework that is intended to address some of the problems and limitations of the use of SMTP and RFC 822 for electronic mail, in a manner that is compatible with existing RFC 822 implementations. [RODR02] lists the following limitations of the SMTP/822 scheme:

1. SMTP cannot transmit executable files or other binary objects. A number of schemes are in use for converting binary files into a text form that can be used by SMTP mail systems, including the popular UNIX UUencode/UUdecode scheme. However, none of these is a standard or even a de facto standard.
2. SMTP cannot transmit text data that includes national language characters because these are represented by 8-bit codes with values of 128 decimal or higher, and SMTP is limited to 7-bit ASCII.
3. SMTP servers may reject mail messages over a certain size.
4. SMTP gateways that translate between the character codes ASCII and EBCDIC do not use a consistent set of mappings, resulting in translation problems.
5. SMTP gateways to X.400 electronic mail networks cannot handle nontextual data included in X.400 messages.
6. Some SMTP implementations do not adhere completely to the SMTP standards defined in RFC 821.

The MIME specification includes the following elements:

1. Five new message header fields are defined, which may be included in an RFC 822 header. These fields provide information about the body of the message. The five header fields defined in MIME are:

- **MIME-Version:** Must have the parameter value 1.0. This field indicates that the message conforms to the RFCs.
- **Content-Type:** Describes the data contained in the body with sufficient detail that the receiving user agent can pick an appropriate agent or mechanism to present the data to the user or otherwise deal with the data in an appropriate manner.
- **Content-Transfer-Encoding:** Indicates the type of transformation that has been used to represent the body of the message in a way that is acceptable for mail transport.
- **Content-ID:** Used to uniquely identify MIME entities in multiple contexts.
- **Content-Description:** A plaintext description of the object with the body; this is useful when the object is not displayable

2. A number of content formats are defined, thus standardizing representations that support multimedia electronic mail.

3. Transfer encodings are defined that enable the conversion of any content format into a form that is protected from alteration by the mail system.

Content Type/Subtype

- Text body - in given character set
- Multipart - body contains multiple parts

- Message
- Image
- Video
- Audio
- Application

FILE TRANSFER

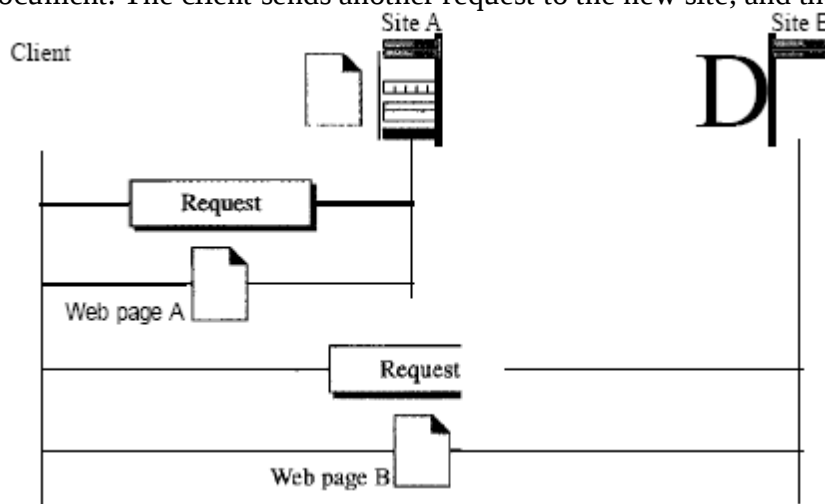
Transferring files from one computer to another is one of the most common tasks expected from a networking or internetworking environment. As a matter of fact, the greatest volume of data exchange in the Internet today is due to file transfer. File transfer Protocol (*FTP*).

WWW and HTTP

The **World Wide Web** (WWW) is a repository of information linked together from points all over the world. The WWW has a unique combination of flexibility, portability, and user-friendly features that distinguish it from other services provided by the Internet. The WWW project was initiated by CERN (European Laboratory for Particle Physics) to create a system to handle distributed resources necessary for scientific research.

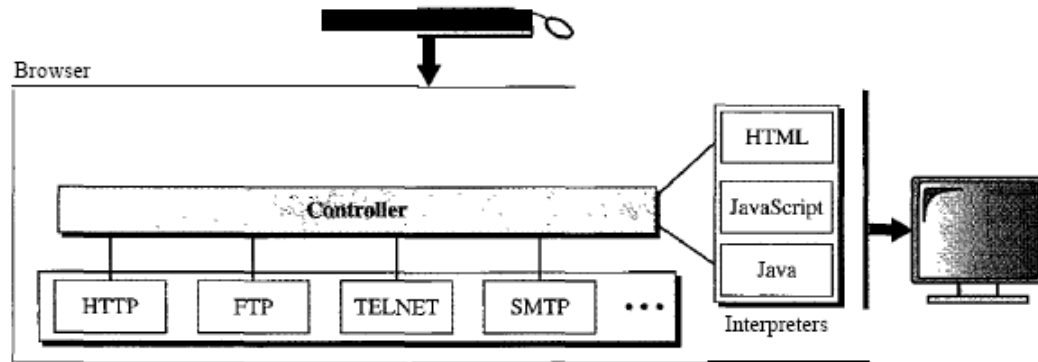
ARCHITECTURE

The WWW today is a distributed client/server service, in which a client using a browser can access a service using a server. However, the service provided is distributed over many locations called *sites*. Each site holds one or more documents, referred to as *Web pages*. Each Web page can contain a link to other pages in the same site or at other sites. The pages can be retrieved and viewed by using browsers. Let us go through the scenario shown in Figure below. The client needs to see some information that it knows belongs to site A. It sends a request through its browser, a program that is designed to fetch Web documents. The request, among other information, includes the address of the site and the Web page, called the URL, which we will discuss shortly. The server at site A finds the document and sends it to the client. When the user views the document, she finds some references to other documents, including a Web page at site B. The reference has the URL for the new site. The user is also interested in seeing this document. The client sends another request to the new site, and the new page is retrieved.



Client (Browser)

A variety of vendors offer commercial browsers that interpret and display a Web document, and all use nearly the same architecture. Each browser usually consists of three parts: a controller, client protocol, and interpreters. The controller receives input from the keyboard or the mouse and uses the client programs to access the document. After the document has been accessed, the controller uses one of the interpreters to display the document on the screen. The client protocol can be one of the protocols described previously such as FfP or HTIP (described later in the chapter). The interpreter can be HTML, Java, or JavaScript, depending on the type of document.



Server

The Web page is stored at the server. Each time a client request arrives, the corresponding document is sent to the client. To improve efficiency, servers normally store requested files in a cache in memory; memory is faster to access than disk. A server can also become more efficient through multithreading or multiprocessing. In this case, a server can answer more than one request at a time.

Uniform Resource Locator

A client that wants to access a Web page needs the address. To facilitate the access of documents distributed throughout the world, HTTP uses locators. The uniform resource locator (URL) is a standard for specifying any kind of information on the Internet. The URL defines four things: protocol, host computer, port, and path. The *protocol* is the client/server program used to retrieve the document. Many different protocols can retrieve a document; among them are FTP or HTTP. The most common today is HTTP. The host is the computer on which the information is located, although the name of the computer can be an alias. Web pages are usually stored in computers, and computers are given alias names that usually begin with the characters "www". This is not mandatory, however, as the host can be any name given to the computer that hosts the Web page. The URL can optionally contain the port number of the server. If the *port* is included, it is inserted between the host and the path, and it is separated from the host by a colon. Path is the pathname of the file where the information is located. Note that the path can itself contain slashes that, in the UNIX operating system, separate the directories from the subdirectories and files.

Cookies

The World Wide Web was originally designed as a stateless entity. A client sends a request; a server responds. Their relationship is over. The original design of WWW, retrieving publicly

available documents, exactly fits this purpose. Today the Web has other functions; some are listed here.

1. Some websites need to allow access to registered clients only.
2. Websites are being used as electronic stores that allow users to browse through the store, select wanted items, put them in an electronic cart, and pay at the end with a credit card.
3. Some websites are used as portals: the user selects the Web pages he wants to see.
4. Some websites are just advertising.

For these purposes, the cookie mechanism was devised. We discussed the use of cookies at the transport layer in Chapter 23; we now discuss their use in Web pages.

Creation and Storage of Cookies

The creation and storage of cookies depend on the implementation; however, the principle is the same

1. When a server receives a request from a client, it stores information about the client in a file or a string. The information may include the domain name of the client, the contents of the cookie (information the server has gathered about the client such as name, registration number, and so on), a timestamp, and other information depending on the implementation.
2. The server includes the cookie in the response that it sends to the client.
3. When the client receives the response, the browser stores the cookie in the cookie directory, which is sorted by the domain server name.

Using Cookies

When a client sends a request to a server, the browser looks in the cookie directory to see if it can find a cookie sent by that server. If found, the cookie is included in the request. When the server receives the request, it knows that this is an old client, not a new one. Note that the contents of the cookie are never read by the browser or disclosed to the user. It is a cookie *made* by the server and *eaten* by the server. Now let us see how a cookie is used for the four previously mentioned purposes:

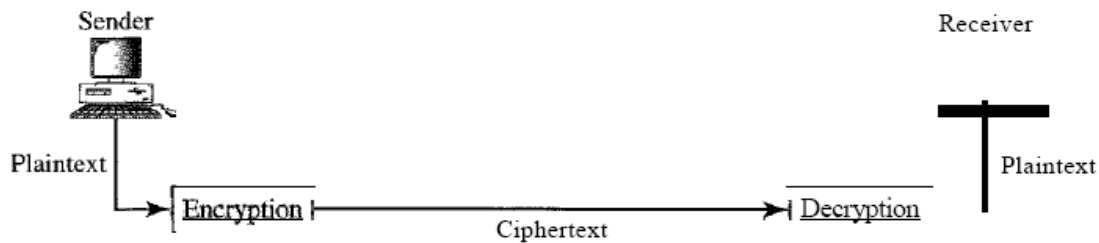
1. The site that restricts access to registered clients only sends a cookie to the client when the client registers for the first time. For any repeated access, only those clients that send the appropriate cookie are allowed.
2. An electronic store (e-commerce) can use a cookie for its client shoppers. When a client selects an item and inserts it into a cart, a cookie that contains information about the item, such as its number and unit price, is sent to the browser. If the client selects a second item, the cookie is updated with the new selection information. And so on. When the client finishes shopping and wants to check out, the last cookie is retrieved and the total charge is calculated.
3. A Web portal uses the cookie in a similar way. When a user selects her favorite pages, a cookie is made and sent. If the site is accessed again, the cookie is sent to the server to show what the client is looking for.
4. A cookie is also used by advertising agencies. An advertising agency can place banner ads on some main website that is often visited by users. The advertising agency supplies only a URL that gives the banner address instead of the banner itself.

When a user visits the main website and clicks on the icon of an advertised corporation, a request is sent to the advertising agency. The advertising agency sends the banner, a GIF file, for example, but it also includes a cookie with the ID of the user. Any future use of the banners adds to the database that profiles the Web behavior of the user. The advertising agency has compiled the interests of the user and can sell this information to other parties. This use of cookies has

made them very controversial. Hopefully, some new regulations will be devised to preserve the privacy of users.

Cryptography

Cryptography, a word with Greek origins, means "secret writing." However, we use the term to refer to the science and art of transforming messages to make them secure and immune to attacks.



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 ciphertext 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 a secure communication. On the contrary, one cipher can serve millions of communicating pairs.

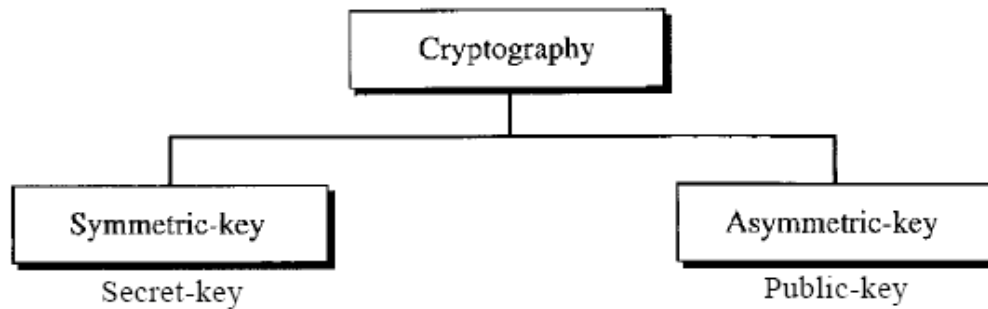
Key

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 the 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.

Alice, Bob, and Eve In cryptography, it is customary to use three characters in an information exchange scenario; we use Alice, Bob, and Eve. Alice is the person who needs to send secure data. Bob is the recipient of the data. Eve is the person who somehow disturbs the communication between Alice and Bob by intercepting messages to uncover the data or by sending her own disguised messages. These three names represent computers or processes that actually send or receive data, or intercept or change data.

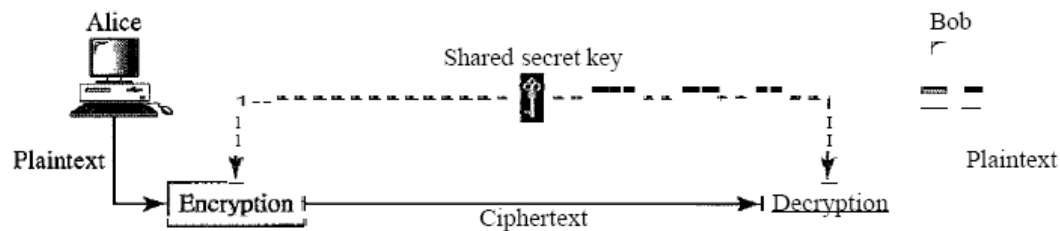
Two Categories

We can divide all the cryptography algorithms (ciphers) into two groups: symmetrickey (also called secret-key) cryptography algorithms and asymmetric (also called public-key) cryptography algorithms.



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.

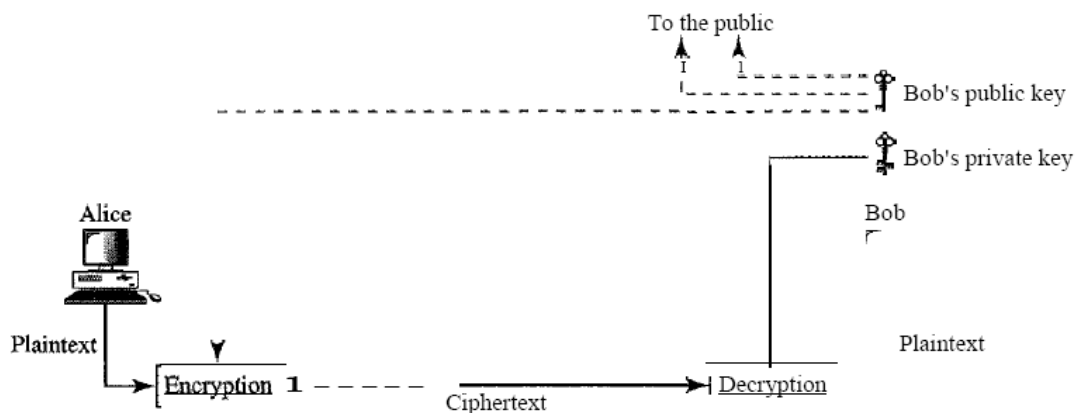


In symmetric-key cryptography, the same key is used by the sender (for encryption) and the receiver (for decryption). The key is shared.

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 Figure below, imagine Alice wants to send a message to Bob. Alice uses the public key



to encrypt the message. When the message is received by Bob, the private key is used to decrypt the message.

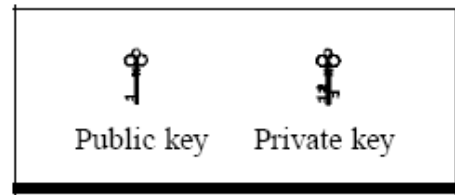
Three Types of Keys

The reader may have noticed that we are dealing with three types of keys in cryptography: the secret key, the public key, and the private key. The first, the secret key, is the shared key used in symmetric-key cryptography. The second and the third are the public and private keys used in

asymmetric-key cryptography. We will use three different icons for these keys throughout the book to distinguish one from the others



Symmetric-key cryptography

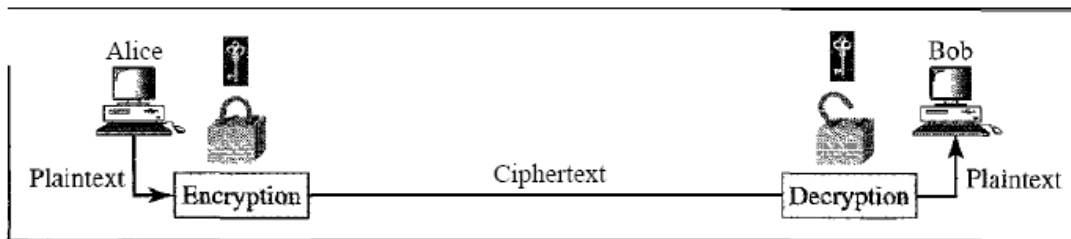


Asymmetric-key cryptography

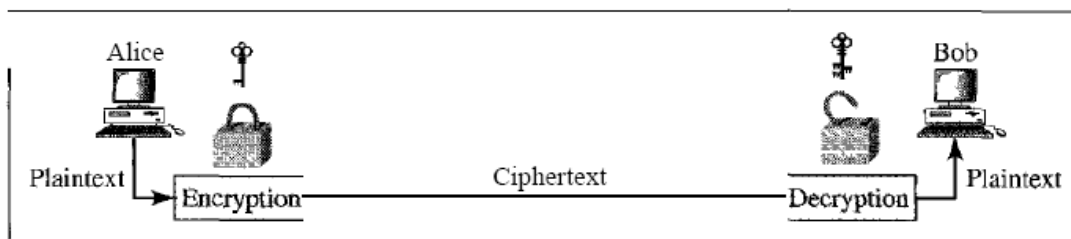
Comparison

Let us compare symmetric-key and asymmetric-key cryptography. Encryption can be thought of as electronic locking; decryption as electronic unlocking. The sender puts the message in a box and locks the box by using a key; the receiver unlocks the box with a key and takes out the message. The difference lies in the mechanism of the locking and unlocking and the type of keys used.

In symmetric-key cryptography, the same key locks and unlocks the box. In asymmetric-key cryptography, one key locks the box, but another key is needed to unlock it.



a. Symmetric-key cryptography



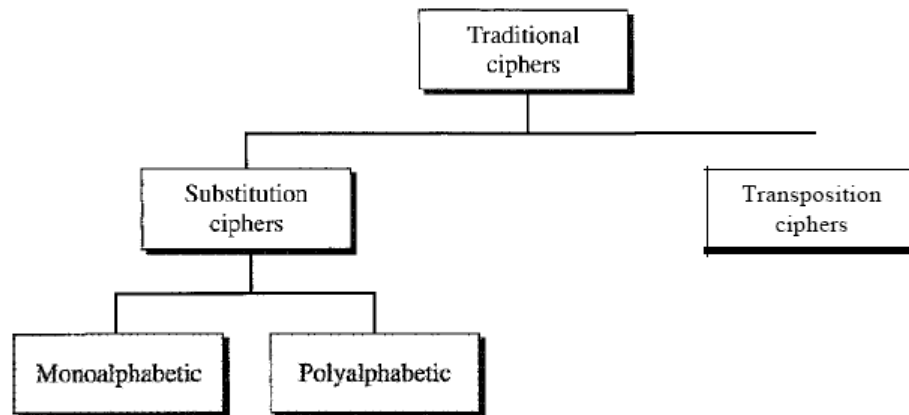
b. Asymmetric-key cryptography

SYMMETRIC-KEY CRYPTOGRAPHY

Symmetric-key cryptography started thousands of years ago when people needed to exchange secrets (for example, in a war).

Traditional Ciphers

We can divide traditional symmetric-key ciphers into two broad categories: substitution ciphers and transposition ciphers



Substitution Cipher

A substitution cipher substitutes one symbol with another. If the symbols in the plaintext are alphabetic characters, we replace one character with another. For example, we can replace character A with D, and character T with Z. If the symbols are digits (0 to 9), we can replace 3 with 7, and 2 with 6. Substitution ciphers can be categorized as either monoalphabetic or polyalphabetic ciphers.

A substitution cipher replaces one symbol with another

In a monoalphabetic cipher, a character (or a symbol) in the plaintext is always changed to the same character (or symbol) in the ciphertext regardless of its position in the text. For example, if the algorithm says that character A in the plaintext is changed to character D, every character A is changed to character D. In other words, the relationship between characters in the plaintext and the ciphertext is a one-to-one relationship.

In a polyalphabetic cipher, each occurrence of a character can have a different substitute. The relationship between a character in the plaintext to a character in the ciphertext is a one-to-many relationship. For example, character A could be changed to D in the beginning of the text, but it could be changed to N at the middle. It is obvious that if the relationship between plaintext characters and ciphertext characters is one-to-many, the key must tell us which of the many possible characters can be chosen for encryption. To achieve this goal, we need to divide the text into groups of characters and use a set of keys. For example, we can divide the text "THISISANEASYTASK" into groups of 3 characters and then apply the encryption using a set of 3 keys. We then repeat the procedure for the next 3 characters.

Transposition Ciphers

In a transposition cipher, there is no substitution of characters; instead, their locations change. A character in the first position of the plaintext may appear in the tenth position of the ciphertext. A character in the eighth position may appear in the first position. In other words, a transposition cipher reorders the symbols in a block of symbols

A transposition cipher reorders (permutes) symbols in a block of symbols.

Key In a transposition cipher, the key is a mapping between the position of the symbols in the plaintext and cipher text. For example, the following shows the key using a block of four characters:

Plaintext:

Cipher text:

2 4 1 3
1 2 3 4

In encryption, we move the character at position 2 to position 1, the character at position 4 to position 2, and so on. In decryption, we do the reverse. Note that, to be more effective, the key should be long, which means encryption and decryption of long blocks of data. Figure shows encryption and decryption for our four-character block using the above key. The figure shows that the encryption and decryption use the same key. The encryption applies it from downward while decryption applies it upward.

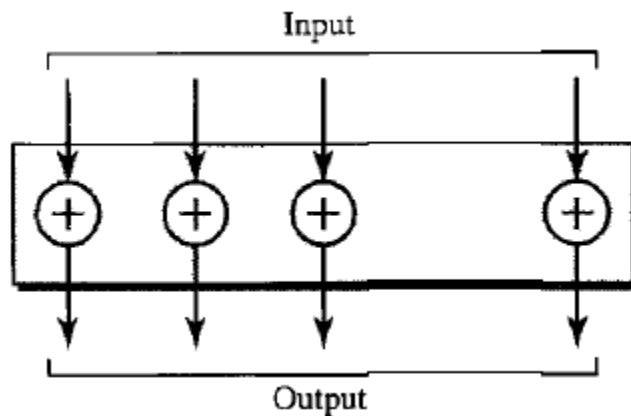
Simple Modern Ciphers

The traditional ciphers we have studied so far are character-oriented. With the advent of the computer, ciphers need to be bit-oriented. This is so because the information to be encrypted is not just text; it can also consist of numbers, graphics, audio, and video data.

It is convenient to convert these types of data into a stream of bits, encrypt the stream, and then send the encrypted stream. In addition, when text is treated at the bit level, each character is replaced by 8 (or 16) bits, which means the number of symbols becomes 8 (or 16). Mingling and mangling bits provides more security than mingling and mangling characters. Modern ciphers use a different strategy than the traditional ones. A modern symmetric cipher is a combination of simple ciphers. In other words, a modern cipher uses several simple ciphers to achieve its goal.

XOR Cipher

Modern ciphers today are normally made of a set of **simple ciphers**, which are simple predefined functions in mathematics or computer science. The first one discussed here is called the **XOR cipher** because it uses the exclusive-or operation as defined in computer science.

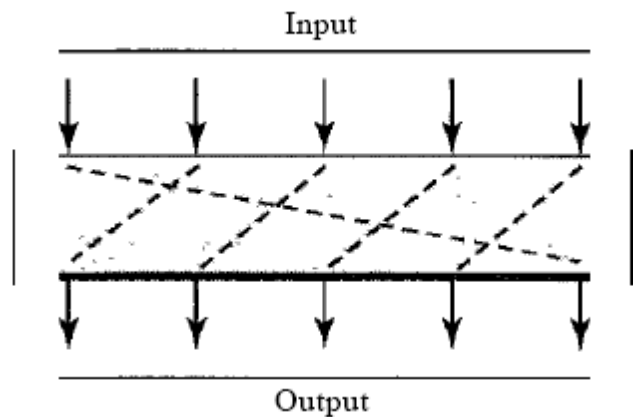


An XOR operation needs two data inputs plaintext, as the first and a key as the second. In other words, one of the inputs is the block to be the encrypted, the other input is a key; the result is the encrypted block. Note that in an XOR cipher, the size of the key, the plaintext, and the ciphertext are all the same. XOR ciphers have a very interesting property: the encryption and decryption are the same.

Rotation Cipher

Another common cipher is the rotation cipher, in which the input bits are rotated to the left or right. The rotation cipher can be keyed or keyless. In keyed rotation, the value of the key defines the number of rotations; in keyless rotation the number of rotations is fixed.

Figure below shows an example of a rotation cipher. Note that the rotation cipher can be considered a special case of the transpositional cipher using bits instead of characters.



The rotation cipher has an interesting property. If the length of the original stream is N , after N rotations, we get the original input stream. This means that it is useless to apply more than $N - 1$ rotations. In other words, the number of rotations must be between 1 and $N - 1$.

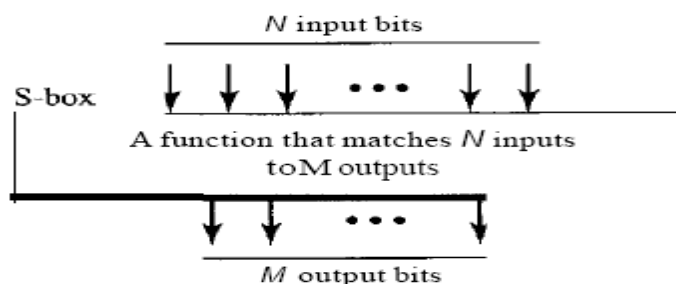
The decryption algorithm for the rotation cipher uses the same key and the opposite rotation direction. If we use a right rotation in the encryption, we use a left rotation in decryption and vice versa.

Substitution Cipher: S-box

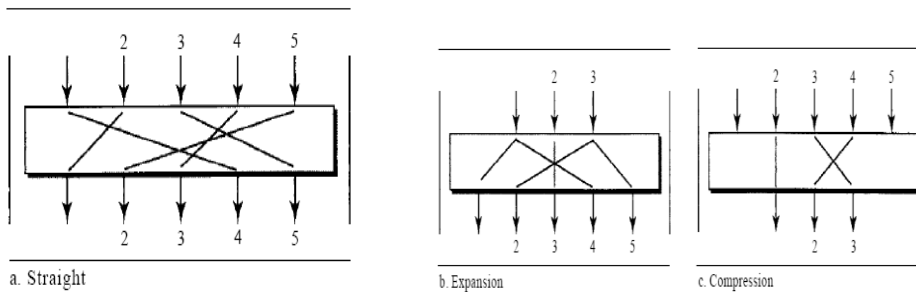
An S-box (substitution box) parallels the traditional substitution cipher for characters. The input to an S-box is a stream of bits with length N ; the result is another stream of bits with length M . And N and M are not necessarily the same. Figure shows an S-box. The S-box is normally keyless and is used as an intermediate stage of encryption or decryption. The function that matches the input to the output may be defined mathematically or by a table.

Transposition Cipher: P-box

A P-box (permutation box) for bits parallels the traditional transposition cipher for characters. It performs a transposition at the bit level; it transposes bits. It can be implemented



in software or hardware, but hardware is faster. P-boxes, like S-boxes, are normally keyless. We can have three types of permutations in P-boxes: the **straight permutation**, **expansion permutation**, and **compression permutation** as shown in Figure



A straight permutation cipher or a straight P-box has the same number of inputs as outputs. In other words, if the number of inputs is N , the number of outputs is also N . In an expansion permutation cipher, the number of output ports is greater than the number of input ports. In a compression permutation cipher, the number of output ports is less than the number of input ports.

Modern Round Ciphers

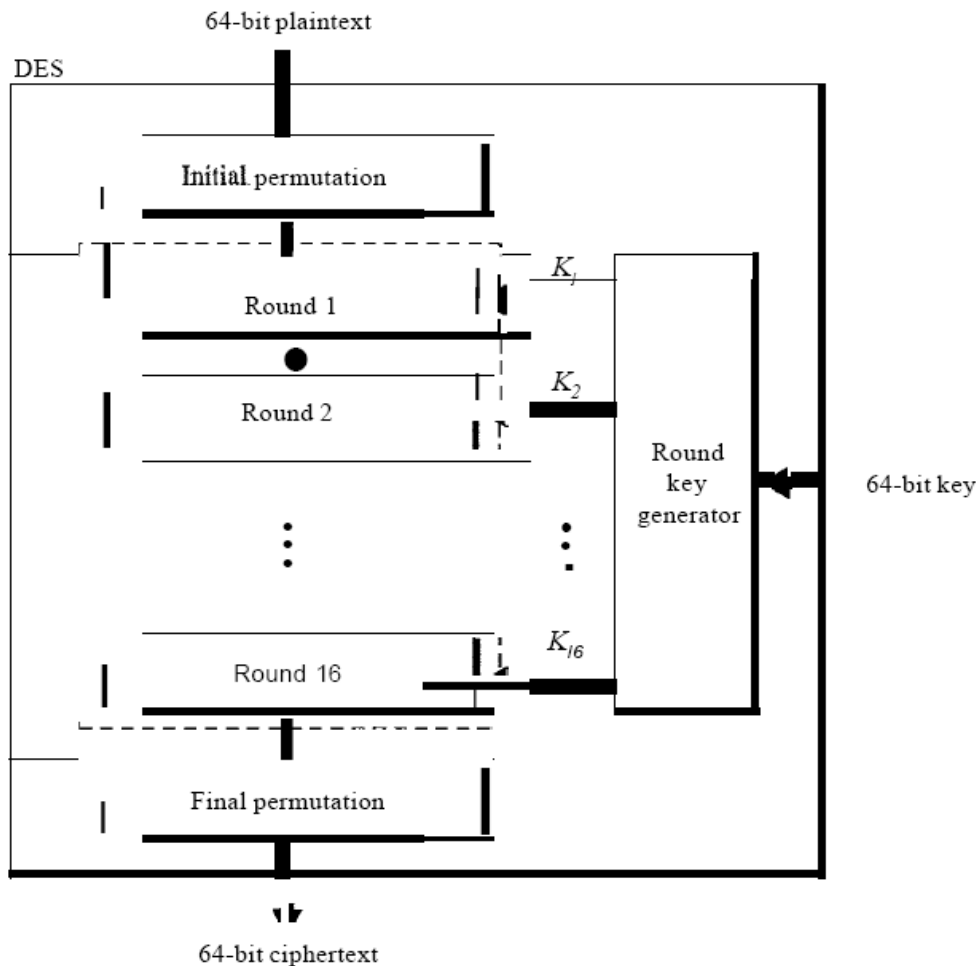
The ciphers of today are called **round ciphers** because they involve multiple **rounds**, where each round is a complex cipher made up of the simple ciphers that we previously described. The key used in each round is a subset or variation of the general key called the round key. If the cipher has N rounds, a key generator produces N keys, $K_1, K_2, \dots, K_N$, where K_1 is used in round 1, K_2 in round 2, and so on.

In this section, we introduce two modern symmetric-key ciphers: DES and AES.

These ciphers are referred to as block ciphers because they divide the plaintext into blocks and use the same key to encrypt and decrypt the blocks. DES has been the de facto standard until recently. AES is the formal standard now.

Data Encryption Standard (DES)

One example of a complex block cipher is the Data Encryption Standard (DES). DES was designed by IBM and adopted by the U.S. government as the standard encryption method for nonmilitary and nonclassified use. The algorithm encrypts a 64-bit plaintext block using a 64-bit key, as shown in Figure

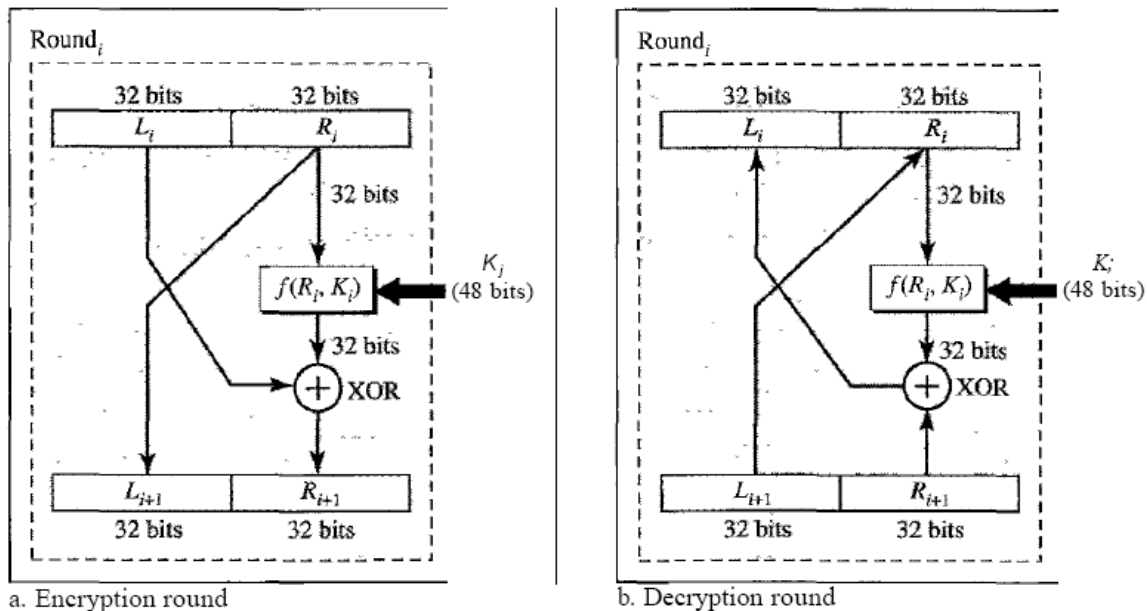


DES has two transposition blocks (P-boxes) and 16 complex round ciphers (they are repeated). Although the 16 iteration round ciphers are conceptually the same, each uses a different key derived from the original key.

The initial and final permutations are keyless straight permutations that are the inverse of each other. The permutation takes a 64-bit input and permutes them according to predefined values. Each round of DES is a complex round cipher, as shown in Figure below. Note that the structure of the encryption round ciphers is different from that of the decryption one.

DES Function

The heart of DES is the **DES function**. The DES function applies a 48-bit key to the rightmost 32 bits R_i to produce a 32-bit output. This function is made up of four operations: an XOR, an expansion permutation, a group of S-boxes, and a straight permutation

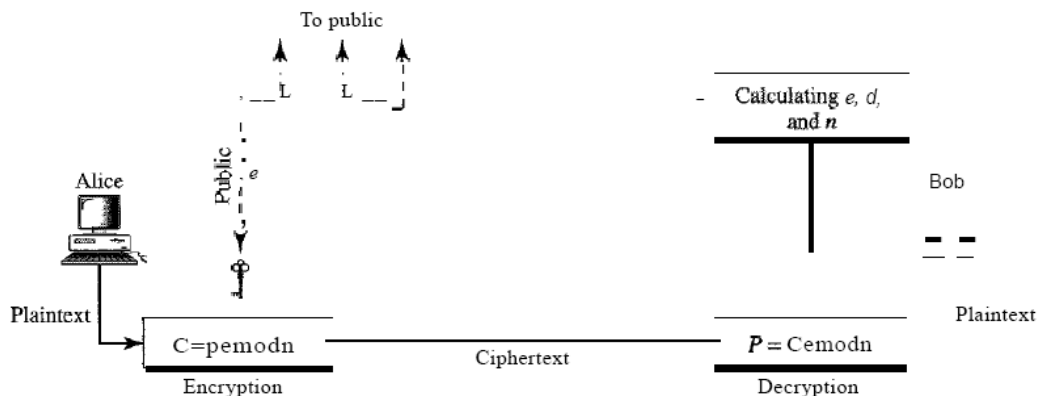


ASYMMETRIC-KEY CRYPTOGRAPHY

An asymmetric-key (or public-key) cipher uses two keys: one private and one public. We discuss two algorithms: RSA and Diffie-Hellman.

RSA

The most common public key algorithm is RSA, named for its inventors Rivest, Shamir, and Adleman (RSA). It uses two numbers, e and d , as the public and private keys.



The two keys, e and d , have a special relationship to each other, a discussion of this relationship is beyond the scope of this book. We just show how to calculate the keys without proof.

Selecting Keys

Bob use the following steps to select the private and public keys:

1. Bob chooses two very large prime numbers p and q . Remember that a prime number is one that can be divided evenly only by 1 and itself.
2. Bob multiplies the above two primes to find n , the modulus for encryption and decryption. In other words, $n ::= p \times q$.
3. Bob calculates another number $\phi ::= (p-1) \times (q-1)$.
4. Bob chooses a random integer e . He then calculates d so that $d \times e ::= 1 \text{ mod } \phi$.
5. Bob announces e and n to the public; he keeps ϕ and d secret.

Encryption

Anyone who needs to send a message to Bob can use n and e . For example, if Alice needs to send a message to Bob, she can change the message, usually a short one, to an integer. This is the plaintext. She then calculates the ciphertext, using e and n .

$$C = \text{pt}^e \pmod{n}$$

Alice sends C , the ciphertext, to Bob.

Decryption

Bob keeps p and d private. When he receives the ciphertext, he uses his private key d to decrypt the message:

$$P = C^d \pmod{n}$$

Restriction

For RSA to work, the value of P must be less than the value of n . If P is a large number, the plaintext needs to be divided into blocks to make P less than n .

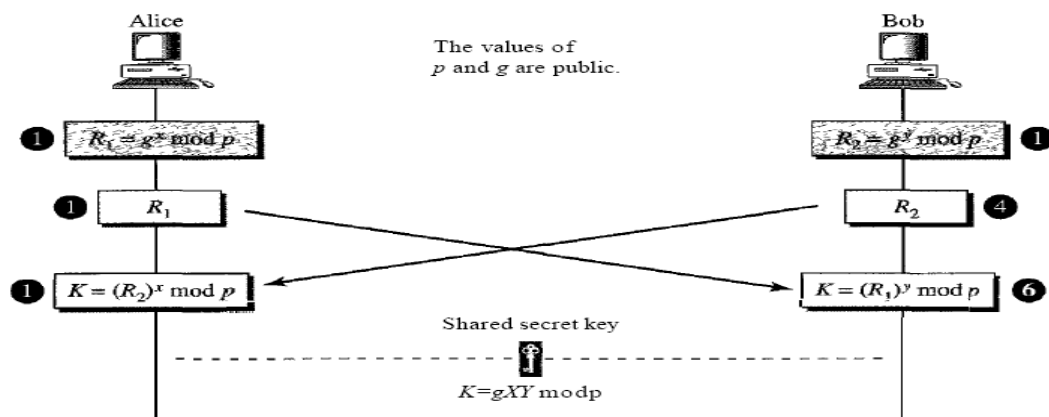
Diffie-Hellman

RSA is a public-key cryptosystem that is often used to encrypt and decrypt symmetric keys. Diffie-Hellman, on the other hand, was originally designed for key exchange. In the

Diffie-Hellman cryptosystem, two parties create a symmetric session key to exchange data without having to remember or store the key for future use. They do not have to meet to agree on the key; it can be done through the Internet. Let us see how the protocol works when Alice and Bob need a symmetric key to communicate. Before establishing a symmetric key, the two parties need to choose two numbers p and g . The first number, p , is a large prime number on the order of 300 decimal digits (1024 bits). The second number is a random number. These two numbers need not be confidential. They can be sent through the Internet; they can be public.

Procedure

Figure below shows the procedure. The steps are as follows:



Step 1: Alice chooses a large random number x and calculates $R_1 = g^x \pmod{p}$.

o Step 2: Bob chooses another large random number y and calculates $R_2 = g^y \pmod{p}$.

o Step 3: Alice sends R_1 to Bob. Note that Alice does not send the value of x ; she sends only R_1 .

o Step 4: Bob sends R_2 to Alice. Again, note that Bob does not send the value of y , he sends only R_2 .

o Step 5: Alice calculates $K = (R_2)^x \pmod{p}$.

o Step 6: Bob also calculates $K = (R_1)^y \pmod{p}$.

The symmetric key for the session is K .

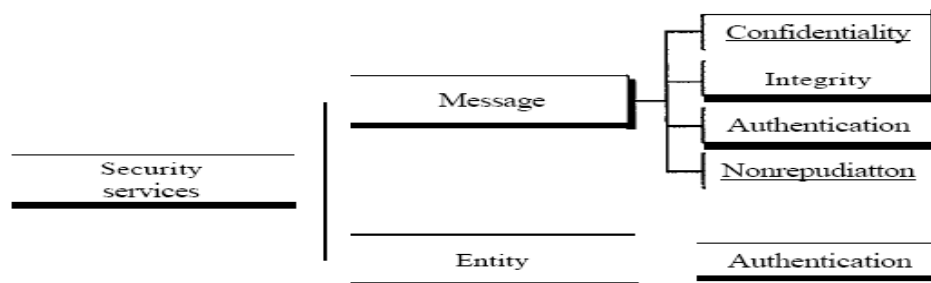
$(f \bmod p)Y \bmod p = (gY \bmod p)X \bmod p = fY \bmod p$

Bob has calculated $K = (R^{-1} \bmod p)(f \bmod p)Y \bmod p = lfY \bmod p$. Alice has calculated $K = (R^2)X \bmod p = (gY \bmod p)X \bmod p = lfY \bmod p$. Both have reached the same value without Bob knowing the value of x and without Alice knowing the value of y .

SECURITY SERVICES

Security Requirements

- **confidentiality** - protect data content/access. Requires that data only be accessible by authorized parties. This type of access includes printing, displaying, and other forms of disclosure, including simply revealing the existence of an object.
- **integrity** - protect data accuracy. Requires that only authorized parties can modify data. Modification includes writing, changing, changing status, deleting, and creating.
- **availability** - ensure timely service. Requires that data are available to authorized parties.
- **authenticity** - protect data origin. Requires that a host or service be able to verify the identity of a user.



Passive Attacks

A useful means of classifying security attacks (RFC 2828) is in terms of *passive attacks* and *active attacks*. A passive attack attempts to learn or make use of information from the system but does not affect system resources. An active attack attempts to alter system resources or affect their operation.

- Passive attacks are in the nature of eavesdropping on, or monitoring of, transmissions. The goal of the opponent is to obtain information that is being transmitted. Two types of passive attacks are release of message contents (which may contain sensitive or confidential information), and traffic analysis which is subtler. Even with encryption protecting message contents, an opponent might still be able to observe the pattern of these messages. The opponent could determine the location and identity of communicating hosts and could observe the frequency and length of messages being exchanged. This information might be useful in guessing the nature of the communication that was taking place.
- Passive attacks are very difficult to detect because they do not involve any alteration of the data. Typically, the message traffic is sent and received in an apparently normal fashion and neither the sender nor receiver is aware that a third party has read the messages or observed the traffic pattern. However, it is feasible to prevent the success of

these attacks, usually by means of encryption. Thus, the emphasis in dealing with passive attacks is on prevention rather than detection.

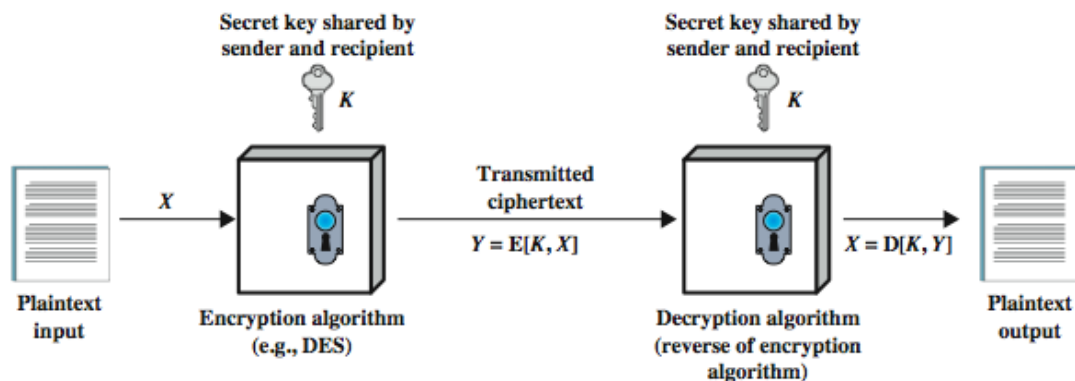
Active Attacks

Active attacks involve some modification of the data stream or the creation of a false stream and can be subdivided into four categories: masquerade, replay, modification of messages, and denial of service.

- Masquerade attacks take place when one entity pretends to be a different entity. A masquerade attack usually includes some other forms of active attack.
- Replay involves the passive capture of a data unit and its subsequent retransmission to produce an unauthorized effect.
- Modification of messages simply means that some portion of a legitimate message is altered, or that messages are delayed or reordered, to produce an unauthorized effect.
- The denial of service prevents or inhibits the normal use or management of communications facilities. This attack may have a specific target. Another form of service denial is the disruption of an entire network or a server.

Active attacks present the opposite characteristics of passive attacks. Whereas passive attacks are difficult to detect, measures are available to prevent their success. On the other hand, it is quite difficult to prevent active attacks absolutely, because to do so would require physical protection of all communications facilities and paths at all times. Instead, the goal is to detect them and to recover from any disruption or delays caused by them. Because the detection has a deterrent effect, it may also contribute to prevention.

Symmetric Encryption



Symmetric encryption, also referred to as conventional encryption or single-key encryption, was the only type of encryption in use prior to the introduction of public-key encryption in the late 1970s. It remains by far the more widely used of the two types of encryption. A symmetric encryption scheme has five ingredients, as shown in Stallings DCC8e Figure above:

Plaintext: This is the original message or data that is fed into the algorithm as input.

- **Encryption algorithm:** The encryption algorithm performs various substitutions and transformations on the plaintext.
- **Secret key:** The secret key is also input to the encryption algorithm. The exact substitutions and transformations performed by the algorithm depend on the key.

- **Ciphertext:** This is the scrambled message produced as output. It depends on the plaintext and the secret key. For a given message, two different keys will produce two different ciphertexts.
- **Decryption algorithm:** This is essentially the encryption algorithm run in reverse. It takes the ciphertext and the secret key and produces the original plaintext.

Data Encryption Standard

- DES has been the dominant encryption algorithm since its introduction in 1977. However, because DES uses only a 56-bit key, it was only a matter of time before computer processing speed made DES obsolete. In 1998, the Electronic Frontier Foundation (EFF) announced that it had broken a DES challenge using a special-purpose "DES cracker" machine that was built for less than \$250,000. The attack took less than three days. The EFF has published a detailed description of the machine, enabling others to build their own cracker.

Triple DEA

The life of DES was extended by the use of triple DES (3DES), which involves repeating the basic DES algorithm three times, using either two or three unique keys, for a key size of 112 or 168 bits. The principal drawback of 3DES is that the algorithm is relatively sluggish in software. A secondary drawback is that both DES and 3DES use a 64-bit block size. For reasons of both efficiency and security, a larger block size is desirable.

Advanced Encryption Standard

- NIST issued call for proposals for an Advanced Encryption Standard (AES) in 1997
 - security strength equal to or better than 3DES
 - significantly improved efficiency
 - symmetric block cipher with block length 128 bits
 - key lengths 128, 192, and 256 bits
 - evaluation include security, computational efficiency, memory requirements, hardware and software suitability, and flexibility
 - AES issued as FIPS 197 in 2001

AES Description

- **assume key length 128 bits**
- **input a 128-bit block (square matrix of bytes)**
 - copied into state array, modified at each stage
 - after final stage, state copied to output
- **128-bit key (square matrix of bytes)**
 - expanded into array of 44 32-bit key schedule words
- **byte ordering by column**
- also used for data integrity

Secure Hash Algorithm

The Secure Hash Algorithm (SHA) was developed by NIST and published as a federal information processing standard (FIPS 180) in 1993; a revised version was issued as FIPS 180-1 in 1995 and is generally referred to as SHA-1. SHA-1 produces a hash value of 160 bits. In 2002, NIST produced a new version of the standard, FIPS 180-2, that defined three new versions of SHA, with hash value lengths of 256, 384, and 512 bits, known as SHA-256, SHA-384, and

SHA-512. These new versions have the same underlying structure and use the same types of modular arithmetic and logical binary operations as SHA-1. In 2005, NIST announced the intention to phase out approval of SHA-1 and move to a reliance on the other SHA versions by 2010. Shortly thereafter, a research team described an attack in which two separate messages could be found that deliver the same SHA-1 hash using 2^{69} operations, far fewer than the 2^{80} operations previously thought needed to find a collision with an SHA-1 hash. This result should hasten the transition to the other versions of SHA. The SHA-512 algorithm takes as input a message with a maximum length of less than 2^{128} bits and produces as output a 512-bit message digest. The input is processed in 1024-bit blocks.

- **Secure Hash Algorithm (SHA)**
 - **SHA defined in FIPS 180 (1993), 160-bit hash**
 - **SHA-1 defined in FIPS 180-1 (1995)**
 - **SHA-256, SHA-384, SHA-512 defined in FIPS 180-2 (2002), 256/384/512-bit hashes**
- **SHA-1 being phased out, attack known**
- **SHA-512 processes input message**
 - **with total size less than 2^{128} bits**
 - **in 1024 bit blocks**
 - **to produce a 512-bit digest**

Digital Signatures

Public-key encryption can be used in another way, as illustrated in Stallings DCC8e. Suppose that Bob wants to send a message to Alice and, although it is not important that the message be kept secret, he wants Alice to be certain that the message is indeed from him. In this case Bob uses his own private key to encrypt the message. When Alice receives the ciphertext, she finds that she can decrypt it with Bob's public key, thus proving that the message must have been encrypted by Bob. No one else has Bob's private key and therefore no one else could have created a ciphertext that could be decrypted with Bob's public key. Therefore, the entire encrypted message serves as a digital signature. In addition, it is impossible to alter the message without access to Bob's private key, so the message is authenticated both in terms of source and in terms of data integrity.

RSA Algorithm

One of the first public-key schemes was developed in 1977 by Ron Rivest, Adi Shamir, and Len Adleman at MIT and first published in 1978. The RSA scheme has since that time reigned supreme as the only widely accepted and implemented approach to public-key encryption. RSA is a block cipher in which the plaintext and ciphertext are integers between 0 and $n - 1$ for some n . Encryption and decryption are of the following form, for some plaintext block M and ciphertext block C :

$$C = M^e \bmod n$$

$$M = C^d \bmod n = (M^e)^d \bmod n = M^{ed} \bmod n$$

Both sender and receiver must know the values of n and e , and only the receiver knows the value of d . This is a public-key encryption algorithm with a public key of $PU = \{e, n\}$ and a private key of $PR = \{d, n\}$. For this algorithm to be satisfactory for public-key encryption, some

requirements must be met. The first two are easy. The third can be met for large values of e and n .

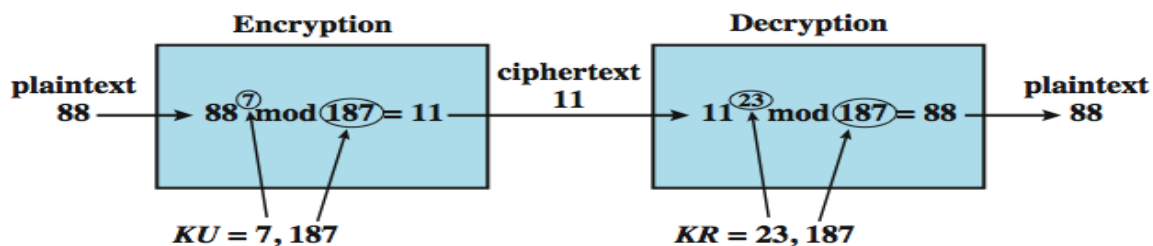
Stallings DCC8e summarizes the RSA algorithm. Begin by selecting two prime numbers, p and q and calculating their product n , which is the modulus for encryption and decryption. Next, we need the quantity $\phi(n)$, referred to as the **Euler** totient of n , which is the number of positive integers less than n and relatively prime to n . Then select an integer e that is relatively prime to $\phi(n)$ [i.e., the greatest common divisor of e and $\phi(n)$ is 1]. Finally, calculate d such that $de \bmod \phi(n) = 1$. It can be shown that d and e have the desired properties.

Key Generation	
Select p, q	p and q both prime
Calculate $n = p \times q$	
Calculate $\phi(n) = (p - 1)(q - 1)$	
Select integer e	$\gcd(\phi(n), e) = 1; 1 < e < \phi(n)$
Calculate d	$d = e^{-1} \bmod \phi(n)$
Public key	$KU = \{e, n\}$
Private key	$KR = \{d, n\}$

Encryption	
Plaintext:	$M < n$
Ciphertext:	$C = M^e \bmod n$

Decryption	
Ciphertext:	C
Plaintext:	$M = C^d \bmod n$

RSA Example



An example is shown in Stallings DCC8e Figure above. For this example, the keys were generated as follows:

1. Select two prime numbers, $p = 17$ and $q = 11$.
2. Calculate $n = pq = 17 \times 11 = 187$.
3. Calculate $\phi(n) = (p - 1)(q - 1) = 16 \times 10 = 160$.
4. Select e such that e is relatively prime to $\phi(n) = 160$ and less than $\phi(n)$; we choose $e = 7$.
5. Determine d such that $de \bmod 160 = 1$ and $d < 160$. The correct value is $d = 23$, because $23 \times 7 = 161 = 10 \times 160 + 1$.

The resulting keys are public key $PU = \{7, 187\}$ and private key $PR = \{23, 187\}$. The example shows the use of these keys for a plaintext input of $M = 88$. For encryption, we need to calculate $C = 88^7 \bmod 187$. For decryption, we calculate $M = 11^{23} \bmod 187$.

RSA Security

There are two possible approaches to defeating the RSA algorithm. The first is the brute force approach: try all possible private keys. Thus, the larger the number of bits in e and d , the more secure the algorithm. However, because the calculations involved, both in key generation and in encryption/decryption, are complex, the larger the size of the key, the slower the system will run. Most discussions of the cryptanalysis of RSA have focused on the task of factoring n into its two prime factors. For a large n with large prime factors, factoring is a hard problem, but not as hard as it used to be. A striking illustration of this is the following. In 1977, the three inventors of RSA dared *Scientific American* readers to decode a cipher they printed in Martin Gardner's "Mathematical Games" column. They offered a \$100 reward for the return of a plaintext sentence, an event they predicted might not occur for some 40 quadrillion years. In April of 1994, a group working over the Internet and using over 1600 computers claimed the prize after only eight months of work. This challenge used a public-key size (length of n) of 129 decimal digits, or around 428 bits. This result does not invalidate the use of RSA; it simply means that larger key sizes must be used. Currently, a 1024-bit key size (about 300 decimal digits) is considered strong enough for virtually all applications.

- **brute force search of all keys**
 - **given size of parameters is infeasible**
 - **but larger keys do slow calculations**
- **factor n to recover p & q**
 - **a hard problem**
 - **well known 129 digit challenge broken in 1994**
 - **key size of 1024-bits (300 digits) currently secure for most apps**

MESSAGE INTEGRITY

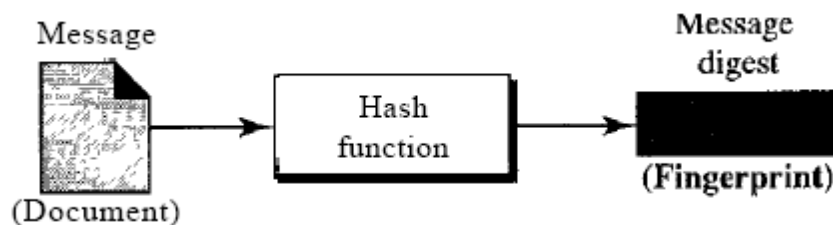
- Encryption and decryption provide secrecy, or confidentiality, but not **integrity**. However, on occasion we may not even need secrecy, but instead must have integrity. For example, Alice may write a will to distribute her estate upon her death. The will does not need to be encrypted. After her death, anyone can examine the will. The integrity of the will, however, needs to be preserved. Alice does not want the contents of the will to be changed. As another example, suppose Alice sends a message instructing her banker,
 - Bob, to pay Eve for consulting work. The message does not need to be hidden from Eve
 - because she already knows she is to be paid. However, the message does need to be safe
 - from any tampering, especially by Eve.

Message and Message Digest

The electronic equivalent of the document and fingerprint pair is the message and message digest pair: To preserve the integrity of a message, the message is passed through an algorithm called a hash function. The hash function creates a compressed image of the message that can be used as a fingerprint. Figure below shows the message, hash function, and the message digest.

Difference

The two pairs document/fingerprint and message/message digest are similar, with some differences. The document and fingerprint are physically linked together; also, neither needs to be kept secret. The message and message digest can be unlinked (or sent) separately and, most importantly, the message digest needs to be kept secret. The message digest is either kept secret in a safe place or encrypted if we need to send it through a communications channel.



Creating and Checking the Digest

The message digest is created at the sender site and is sent with the message to the receiver. To check the integrity of a message, or document, the receiver creates the hash function again and compares the new message digest with the one received. If both are the same, the receiver is sure that the original message has not been changed. Of course, we are assuming that the digest has been sent secretly.

Hash Function Criteria

To be eligible for a hash, a function needs to meet three criteria: one-wayness, resistance to weak collision, and resistance to strong collision

One-wayness

A hash function must have **one-wayness**; a message digest is created by a one-way hashing function. We must not be able to recreate the message from the digest. Sometimes it is difficult to make a hash function 100 percent one-way; the criteria state that it must be extremely difficult or impossible to create the message if the message digest is given. This is similar to the document/fingerprint case. No one can make a document from a fingerprint

Message Authentication

Encryption protects against passive attack (eavesdropping). A different requirement is to protect against active attack (falsification of data and transactions). Protection against such attacks is known as message authentication which allows communicating parties to verify that received messages are authentic. The two important aspects are to verify that the contents of the message have not been altered and that the source is authentic. We may also wish to verify a message's timeliness (it has not been artificially delayed and replayed) and sequence relative to other messages flowing between two parties.

- protection against active attacks with
 - falsification of data
 - falsification of source
- authentication allows receiver to verify that message is authentic
 - has not been altered
 - is from claimed/authentic source
 - timeliness

Authentication Using Symmetric Encryption

It is possible to perform authentication simply by the use of symmetric encryption. If we assume that only the sender and receiver share a key (which is as it should be), then only the genuine sender would be able successfully to encrypt a message for the other participant. Furthermore, if the message includes an error-detection code and a sequence number, the receiver is assured that no alterations have been made and that sequencing is proper. If the message also includes a timestamp, the receiver is assured that the message has not been delayed beyond that normally expected for network transit.

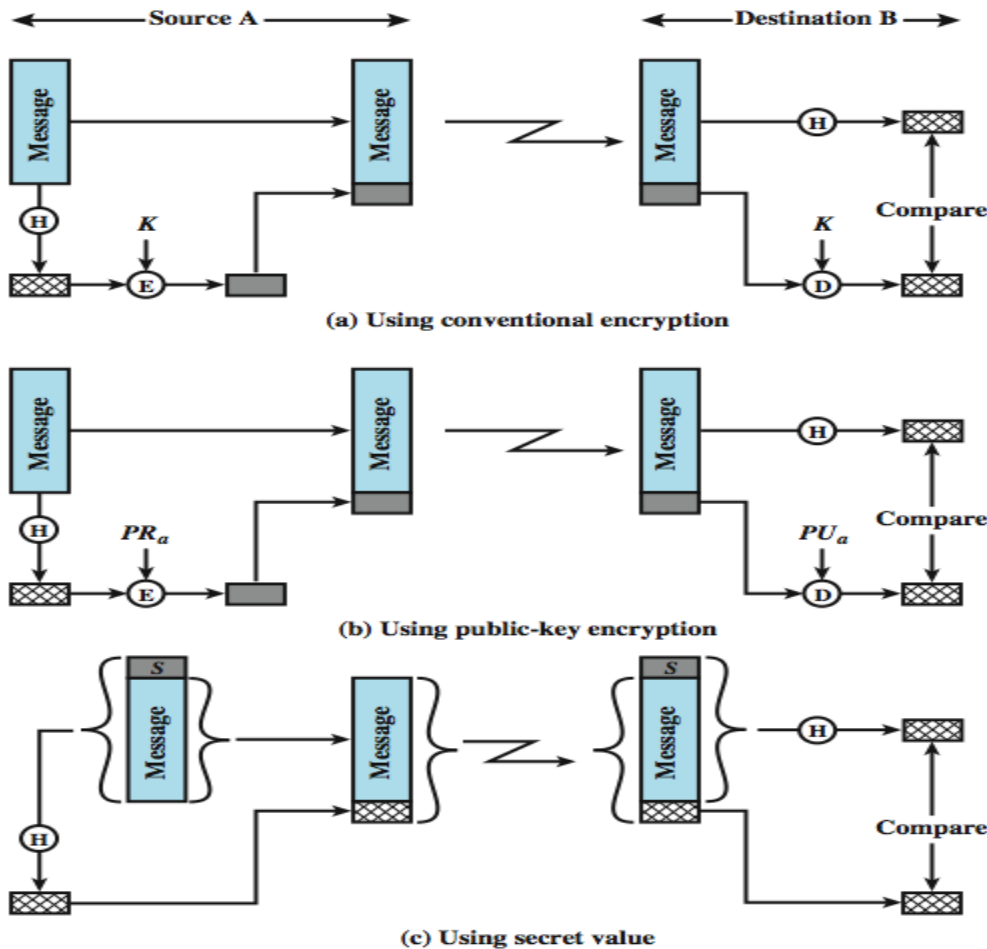
Authentication Without Encryption

- authentication tag generated and appended to each message
- message not encrypted
- useful when don't want encryption because:
 - messages broadcast to multiple destinations
 - have one destination responsible for authentication
 - one side heavily loaded
 - encryption adds to workload
 - can authenticate random messages
 - programs authenticated without encryption can be executed without decoding

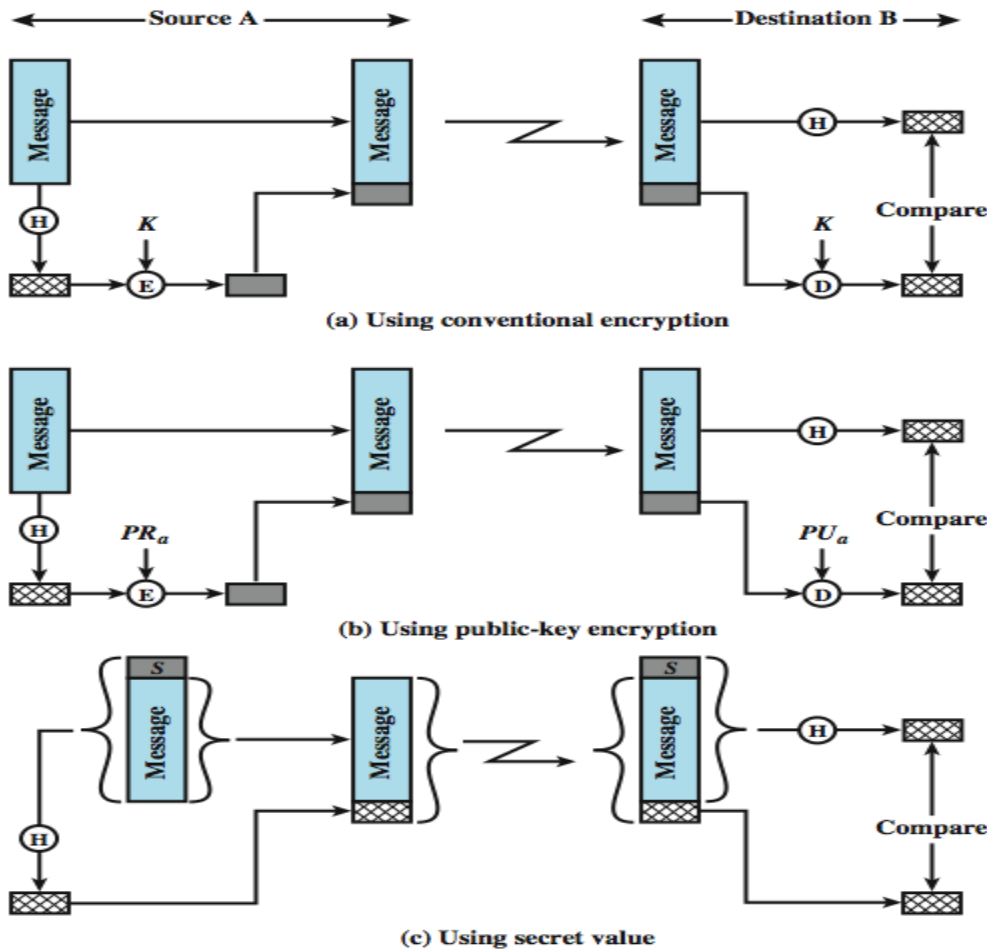
Message Authentication Code

- generate authentication code based on shared key and message
- common key shared between A and B
- if only sender and receiver know key and code matches:
 - receiver assured message has not altered
 - receiver assured message is from alleged sender
 - if message has sequence number, receiver assured of proper sequence
- can use various algorithms, eg. DES

One authentication technique involves the use of a secret key to generate a small block of data, known as a message authentication code, that is appended to the message. This technique assumes that two communicating parties, say A and B, share a common secret key K_{AB} . If we assume that only the receiver and the sender know the identity of the secret key, and if the received code matches the calculated code, then



- 1. The receiver is assured that the message has not been altered. If an attacker alters the message but does not alter the code, then the receiver's calculation of the code will differ from the received code. Because the attacker is assumed not to know the secret key, the attacker cannot alter the code to correspond to the alterations in the message.
- 2. The receiver is assured that the message is from the alleged sender. Because no one else knows the secret key, no one else could prepare a message with a proper code.
- 3. If the message includes a sequence number (such as is used with X.25, HDLC, and TCP), then the receiver can be assured of the proper sequence, because an attacker cannot successfully alter the sequence number.
- A number of algorithms could be used to generate the code. The National Bureau of Standards, in its publication *DES Modes of Operation*, recommends the use of DES. DES is used to generate an encrypted version of the message, and the last number of bits of ciphertext are used as the code. A 16- or 32-bit code is typical. The process just described is similar to encryption, but is less vulnerable.



POSSIBLE QUESTIONS

UNIT- 5

Descriptive Type Questions:

1. Describe congestion control
2. What is URL and describe its components
3. Explain in detail about Web documents categories.
4. Write a short note about data traffic
5. Describe FTP
6. Explain in detail about electronic mail process with neat diagram
7. Describe client server architecture
8. Write the short note on active documents



KARPAGAM ACADEMY OF HIGHER EDUCATION
Deemed to be University Established Under Section 3 of UGC Act 1956)

Eachi Main Road, Eachanari Post, Coimbatore – 641 021. INDIA

DEPARTMENT OF COMPUTER SCIENCE

DATA COMMUNICATION NETWORKS (15CSU502)

ONLINE EXAM QUESTION BANK

UNIT-5

QUESTION	OPTION 1	OPTION 2	OPTION 3	OPTION 4
----- are qualitative values that represent a flow data	data traffic	data descriptor	data traffic and data descriptor	traffic data
the----- define the maximum data rate of the traffic	peak data rate	maximum burst size	bandwidth	effective bandwidth
the----- define the maximum length of time the traffic is generated in the peak rate	effective bandwidth	constant rate	peak data rate	maximum burst size
the----- is the bandwidth that the network needs to allocate for the flow of traffic	effective bandwidth	peak data rate	maximum burst size	data descriptor
a constant-bit-rate is also called as -----	fixed rate	nonfixed rate	definite rate	indefinite rate
in -----the rate of data flow changes in time, with the change smooth instead of sudden and sharp	constant-bit-rate	variable-bit-rate	both a & b	bit rate
in the -----the data rate changes suddenly in a very short times	variable-bit-rate	constant-bit-rate	bursty data	constant-bit-rate
congestion control is divided into-----types	1	2	3	4
a -----is a mechanism that can prevent before and after it happens	open-loop	closed-loop	congestion control	Congestion avoidance
in----- control policies are applied to prevent congestion before it happens	open-loop congestion	closed-loop congestion	Congestion avoidance	congestion control

if the sender feels that a sent packets is lost ,the packet needs to ----- -----	transmission	delete	retransmission	open
the type of ----- at the sender may also affect congestion	closed-loop	window	control	discarding
the ----- policy imposed by the receiver may also effect	acknowledgment	discarding	admission	window
a good ----- policy by the routers may prevent congestion and the same time may not harm the integrity network	admission	window	discarding	acknowledgment
an----- policy , which is a quality of service mechanism	acknowledgment	window	daiscarding	admission
a -----is mechanism try to alleviate congestion after it happens	open-loop	closed-loop	congestion control	Congestion avoidance
the technique od----- refer to congestion control mechanism in which a congestion node stops receiving data from the immediate upstream nodes	choke packet	control	window	backpressure
in----- is anodeto node congestion control that start with a node and propagates	backpressure	choke packet	none	window
a----- is apacket sent by anode to the source to inform it of congestion	control	choke packet	admission	backpressure
in -----there is no communication between the congested nodes and source	explot signaling	left side	implicit signaling	right side
lack of reliablity means losing a -----	packet	control	data flow	admission

a ----- in a file transfer or E-mail is less important	jitter	delay	reliability	speed
a ----- in the variation in delay for packets belonging to the same flow	jitter	reliability	delay	speed
different application need different -----	maximum burst size	effective bandwidth	bandwidth	peak data rate
packets from different flows arrive at ----- ---	scheduling	fifo	bandwidth	switch
a good ----- technique treats the different flows in a pair in appropriate manner	bandwidth	scheduling	admission	window
several scheduling are designed to improve ----- ----	quality of service	quality of data	quality of control	quality of data flow
in ----- queuing , packets wait in a buffer(queue) until the node is ready to process them	lifo	linked	fifo	circular
in ----- queuing , packets are first assigned to a priority class	fifo	lifo	circular	priority
in priority the packets in a----- priority queue are processed first	lowest	highest	medium	topest
a better scheduling method is----- queuing, in this ,the packets are still assigned to different classes	weighted fair	priority	both a & b	fair queuing
the----- is a mechanism to control the amount and rate of traffic sent to the network	priority	data descriptor	traffic shaping	weighted fair
the ----- does not credit an idle host	token bucket	leak bucket	full bucket	empty bucket

a----- algorithm shapes bursty traffic into fixed-rate traffic by averaging the data rate	leak bucket	token bucket	empty bucket	bus
the ----- bucket allows the bursty traffic at a regulated maximum rate	empty bucket	leak bucket	token bucket	star
the----- can be combined to credit an idle host and at the same time regulate the traffic	leak bucket	token bucket	empty bucket	leak and token bucket
_____ allows us to send message include text,audio and video .	mail	internet	E-mail	WWW
the _____ client established a connection with MTA server on the system	MTA	alice	UA	system server
the first component of an electronic mail system is the _____	alice	server	user agent	services provider
_____ is the example of user agents are mail,pine,and elm	user agent	command driven	GUI-based	E-mail
_____ define the names of a special files	local part	domain name	mime	local and domain
the second part of address is _____	system server	internet	domain name	local part
MIME is _____	multiple internet mail extensions	multipurpose interface mail extensions	multipurpose internet mail exchange	multipurpose internet mail extensions
_____ has delete and keep mode	pop	pop2	pop3	pop1
_____ is the mechanism provided by TCP/IP for copying a file from one host to another	FTP	MIME	UA	pop3
_____ is the default format for transferring text files	image	ASCII	data structure	record structure
_____ is the default format for transferring binary files	image	data structure	record structure	ASCII

in the _____ formate, the file is a continuous stream of byte	file structure	record structure	data structure	image
the service provider is distrubuted over many location called _____	internet	sites	www	http
theweb page store at the _____	hard disk	disk	client	server
the _____ is the computer on which the information is located	path	sites	host	cookies
_____ is the pathname of the file where the information is located	host	path	server	sites
_____ is language for creating web pages	HTML	C	C++	java
a _____ is created by a web server whenever a browser request the document	common gate way	dynamic document	script	static script
_____ is the protocol used mainly to access data on the world wide web	communicatio	network	WWW	HTTP
a _____ replaces one symbol with another	substitution cipher	monoalphabeti c cipheth	traditional cipher	ceasar cipher
a _____ reorders (permutes) symbols in a block of symbols	traditional cipher	substitution cipher	transposition cipher	monoalphabetic cipher
the _____ is sometimes referred to as the caesar cipher	monoalphabet ic cipheth	shift cipher	traditional cipher	transposition cipher
a _____ is a techique that emplys the morden block ciphers such as DES and AES	modes	modes of operation	operating server	OS
the most common public key algorithm is _____	RSA	RSSA	RSS	ARQ

_____ means that the data can must arrive at the receiver axactly as they were sent	integrity	message integrity	authentication	message authentication
_____ is the service beyond message integrity	message authentication	message integrity	integrity	authentication
a digital signature needs a _____ system	private-key	primary-key	public-key	secondary- key
a digital signature today provides _____	message authentication	integrity	message integrity	authentication
a _____ keys between two parties is used only once	session	primary-key	private-key	public-key

ANSWER
data descriptor
peak data rate
maximum burst size
effective bandwidth
fixed rate
variable-bit-rate
bursty data
2
congestion control
open-loop congestion

retransmission

window

acknowledgment

discarding

admission

closed-loop

backpressure

backpressure

choke packet

implicit signaling

packet

delay
jitter
bandwidth
switch
scheduling
quality of service
fifo
priority
highest
weighted fair
traffic shaping
leak bucket

leak bucket
token bucket
both a & b
E-mail
MTA
user agent
command driven
local part
domain name
multipurpose internet mail extensions
pop3
FTP
ASCII
image

file structure
sites
server
host
path
HTML
dynamic document
HTTP
substitution cipher
transposition cipher
shift cipher
modes of operation
RSA

message integrity
message authentication
public-key
message integrity
session

Register Number _____
[15CSU502]

KARPAGAM ACADEMY OF HIGHER EDUCATION
KARPAGAM UNIVERSITY
(Deemed to be University Established under Section 3 of UGC Act 1956)

BSc Computer Science
FIRST INTERNAL EXAMINATION- JULY 2017
DATA COMMUNICATION NETWORKS

CLASS: III BSc (CS) – A & B

Time: 2 hours

DATE & SESSION: 18 .07. 2017, N

Maximum: 50 marks

PART – A (20 x 1 = 20 Marks)

Answer ALL the Questions

1. The system must deliver data to the correct destination is called_____
a. Accuracy b. jitter c. delivery d. timeliness
2. In_____, the communication is unidirectional.
a. duplex mode b. full duplex mode c. half duplex mode d. simplex mode
3. One long cable acts as a _____to link all the devices in a network.
a. Bus b. mesh c. hub d. backbone
4. MAN stands for _____
a. Metropolitan area network b. Metropolitan area network
c. Metropolitan area network d. Macro area network
5. In physical layer we can transfer data into _____
a. Frame b. packet c. bit d. byte
6. _____is a type of transmission impairment in which an outside source such as crosstalk corrupts a signal
a. Attenuation b. distortion c. noise d. decibel
7. FTP_____
a. file transmit protocol b. file transmission protocol
c. file transfer protocol d. flip transfer protocol
8. A multipoint is also called as_____
a. multi line b. multi drop c. multi level d. single level
9. A_____ is the set of rules.
a. Protocols b. transmission medium c. networks d. ip
10. The _____is the number of bits sent in a second.
a. Bit length b. bandpass c. bandwidth d. bit rate
11. _____standards are often established originally by manufactures.
a. de jure b. de facto c. de fact d. semantics
12. The_____layer is responsible for providing services to the user.
a. Presentation b. data link c. application d. network

13. RARP is _____
- Reverse address resolution protocol
 - reverse address revolutionized protocol
 - reverse address result protocol
 - reverse address research protocol
14. As frequency increases, the period _____
- Decreases
 - increases
 - remains the same
 - doubles
15. A _____ signal is a composite analog signal with an infinite bandwidth
- simple
 - composite
 - digital
 - Analog
16. A _____ connection provides a dedicated link between two devices.
- Point-to-point
 - multi-point
 - mesh
 - physical
17. The _____ layer is responsible for process to process delivery.
- physical
 - presentation
 - networks
 - transport
18. Which multiplexing technique transmits analog signals
- FDM
 - TDM
 - WDM
 - TDM and WDM
19. A _____ is a set of devices connected by communication links.
- Protocols
 - networks
 - computer
 - printer
20. The _____ layer is responsible for movements of bits from one hop to next.
- data link
 - physical
 - transport
 - session

PART – B (3 x 10 = 30 Marks)

Answer ALL the Questions

21. (a). Elucidate the layered architecture of OSI reference model with a neat diagram.
[OR]
(b). Write a note on the transmission impairments
22. (a). Describe the architecture of LAN, MAN, WAN and their differences with neat sketches.
[OR]
(b). List out the types of addresses used in TCP/IP protocol suite and explain the association with the layers.
23. (a). Explain the various network topologies with a neat diagram
[OR]
(b). Illustrate the working of frequency division multiplexing with a neat diagram.

Register Number _____
[15CSU502]

KARPAGAM ACADEMY OF HIGHER EDUCATION
KARPAGAM UNIVERSITY
(Deemed to be University Established under Section 3 of UGC Act 1956)

BSc Computer Science
FIRST INTERNAL EXAMINATION- JULY 2017
DATA COMMUNICATION NETWORKS
ANSWER KEY

CLASS: III BSc (CS) – A & B

DATE & SESSION: 18 .07. 2017, N

Time: 2 hours

Maximum: 50 marks

PART – A (20 x 1 = 20 Marks)

Answer ALL the Questions

1. The system must deliver data to the correct destination is called_____
a. **Accuracy** b. jitter c. delivery d. timeliness
2. In_____, the communication is unidirectional.
a. duplex mode b. full duplex mode c. half duplex mode d. **simplex mode**
3. One long cable acts as a _____to link all the devices in a network.
a. Bus b. mesh c. hub d. **backbone**
4. MAN stands for _____
a. Metropolitan area network b. **Metropolitan area network**
c. Metropolitical area network d. Macro area network
5. In physical layer we can transfer data into _____
a. Frame b. packet c. **bit** d. byte
6. _____is a type of transmission impairment in which an outside source such as crosstalk corrupts a signal
a. Attenuation b. distortion c. **noise** d. decibel
7. FTP_____
a. file transmit protocol b. **file transmission protocol**
c. file transfer protocol d. flip transfer protocol
8. A multipoint is also called as_____
a. multi line b. **multi drop** c. multi level d. single level
9. A_____ is the set of rules.
a. **Protocols** b. transmission medium c. networks d. ip
10. The _____is the number of bits sent in a second.
a. Bit length b. bandpass c. bandwidth d. **bit rate**
11. _____standards are often established originally by manufactures.
a. de jure b. **de facto** c. de fact d. semantics
12. The_____layer is responsible for providing services to the user.

- a. Presentation b. data link c. **application** d. network

13. RARP is _____

- a. **Reverse Address Resolution Protocol**
b. Reverse Address Revolutionized Protocol
c. Reverse Address Result Protocol
d. Reverse Address Research Protocol

14. As frequency increases, the period _____

- a. **Decreases** b. increases c. remains the same d. doubles

15. A _____ signal is a composite analog signal with an infinite bandwidth

- a. simple b. composite c. **digital** d. Analog

16. A _____ connection provides a dedicated link between two devices.

- a. **Point-to-point** b. multi-point c. mesh d. physical

17. The _____ layer is responsible for process to process delivery.

- a. physical b. presentation c. networks d. **transport**

18. Which multiplexing technique transmits analog signals

- a. **FDM** b. TDM c. WDM d. TDM and WDM

19. A _____ is a set of devices connected by communication links.

- a. Protocols b. **networks** c. computer d. printer

20. The _____ layer is responsible for movements of bits from one hop to next.

- a. **data link** b. physical
c. transport d. session

PART – B (3 x 10 = 30 Marks)

Answer ALL the Questions

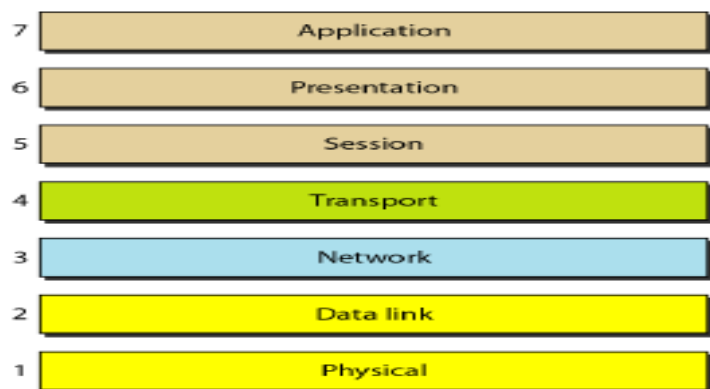
21 (a). Elucidate the layered architecture of OSI reference model with a neat diagram.

THE OSI MODEL

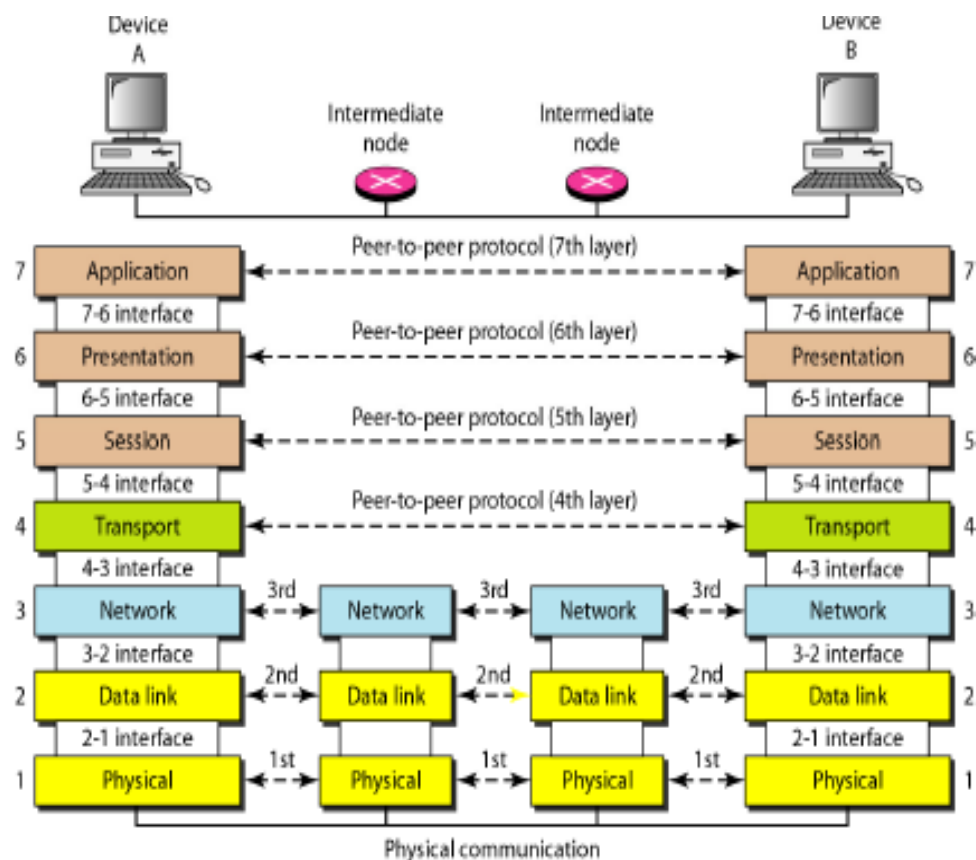
Established in 1947, the International Standards Organization (ISO) is a multinational body dedicated to world wide agreement on international standards.

An ISO standard that covers all aspects of network communications is the Open Systems Interconnection (OSI) model. It was first introduced in the late 1970s.

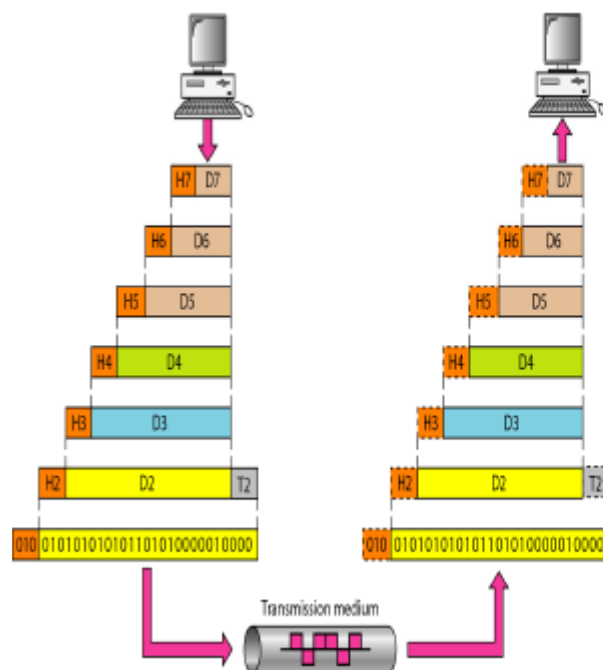
Seven layers of the OSI model



Peer-to-peer processes



Encapsulation

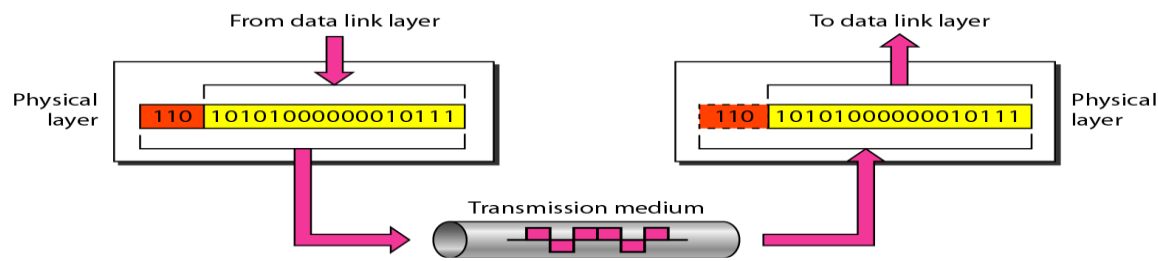


- The process starts at layer 7 (application layer), then moves from layer to layer in descending, sequential order.
- At each layer, a **header** is added to the data unit.
- At layer 2, a **trailer** is added as well.
- When the formatted data unit passes through the physical layer (layer 1) it is changed into an electromagnetic or optical signal and transported along a physical link
- At the destination the reverse process is performed

THE OSI MODEL AND LAYERS

In this section we briefly describe the functions of each layer in the OSI model.

Physical Layer



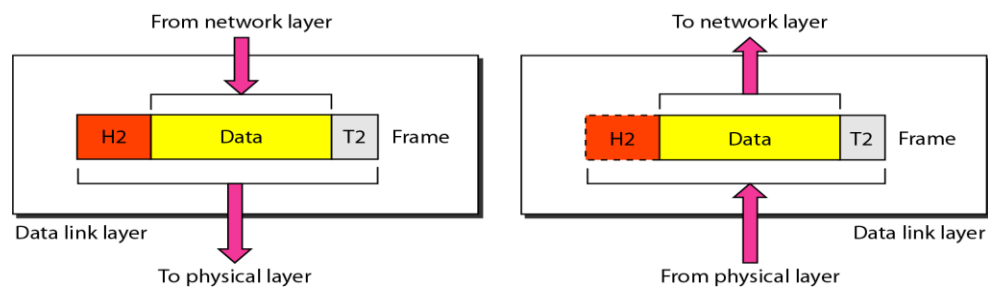
- The physical layer is responsible for movements of individual bits from one hop (node) to the next
- Mechanical and electrical specification, the procedures and functions

Duties:

- Physical characteristics of interfaces and media
- Representation of bits
- Data rate
- Synchronization of bits
- Line configuration
- Physical topology
- Transmission mode

Data link layer

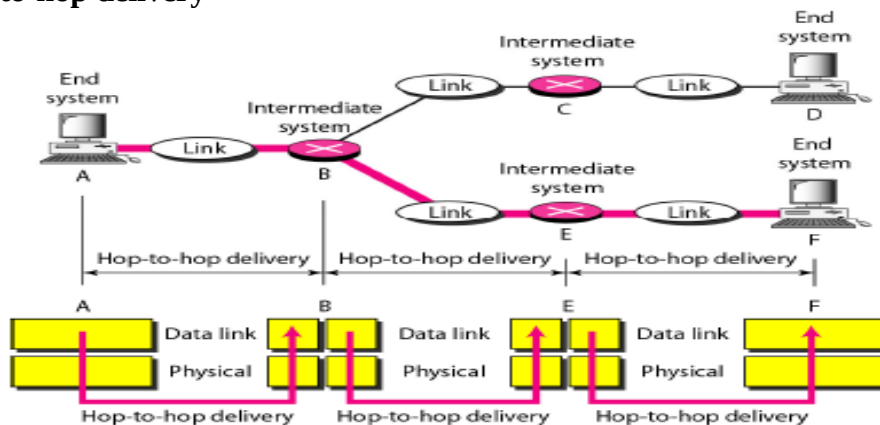
- The data link layer is responsible for moving frames from one hop (node) to the next
- Transform the physical layer to a reliable (error-free) link



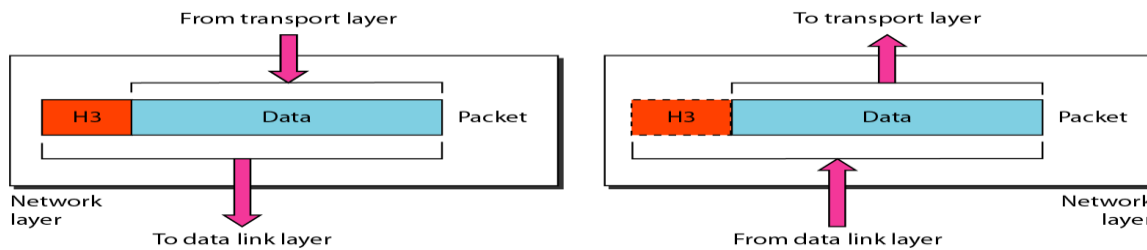
Duties:

- Framing
- Physical addressing
- Flow control
- Error control
- Access control

Hop-to-hop delivery



Network layer

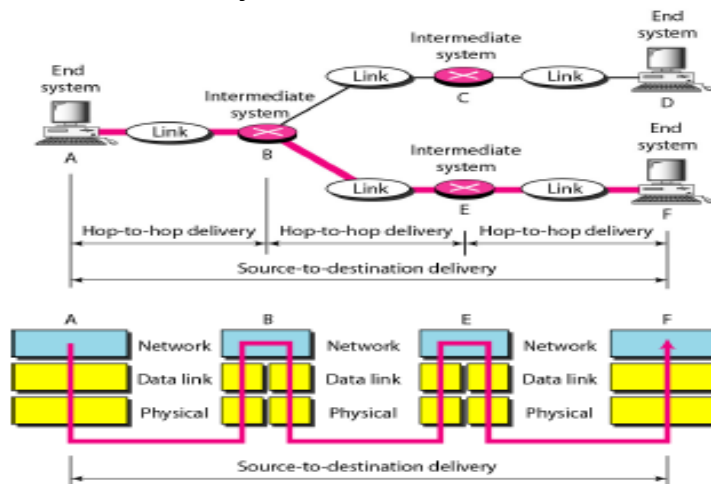


The network layer is responsible for the delivery of individual packets from the source host to the destination host.

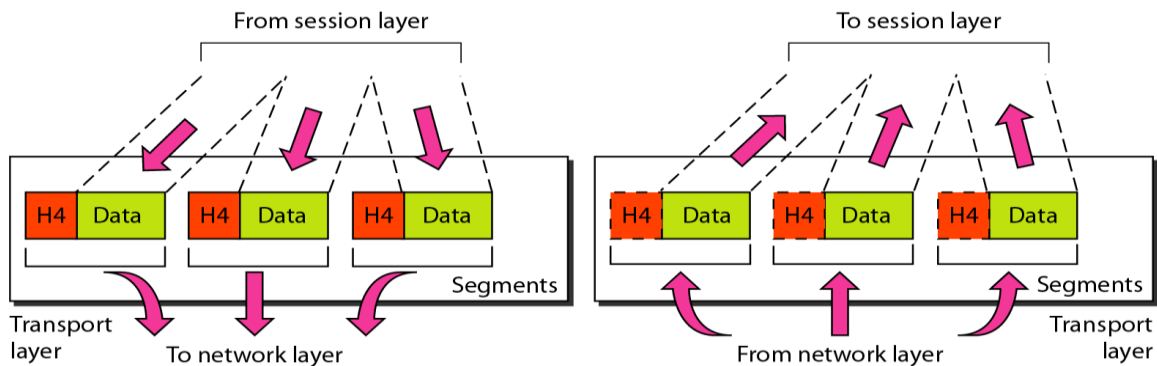
Duties:

- Logical addressing
- Routing

Source-to-destination delivery



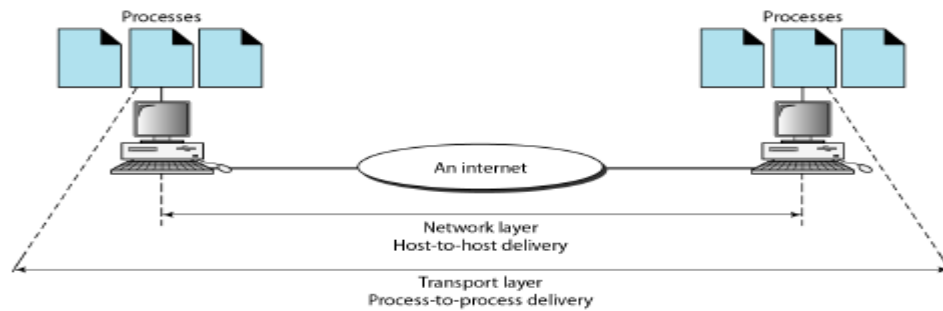
Transport layer



The transport layer is responsible for the delivery of a message from one process to another.

Duties:

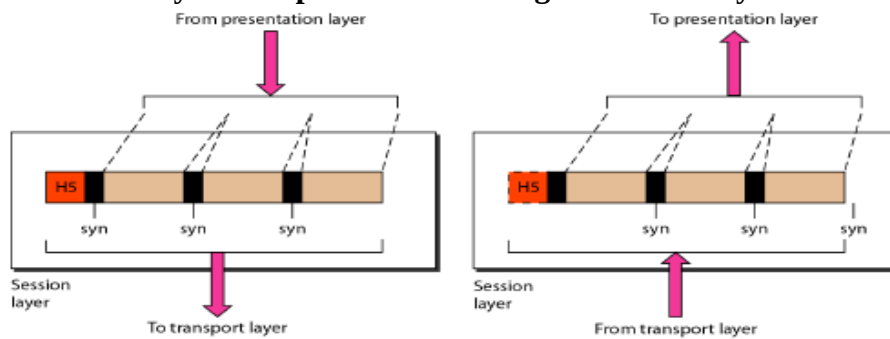
- Service-point (port) addressing
- Segmentation and reassembly
- Connection control
- Flow control
- Error control



**Reliable
process-to-
process delivery**

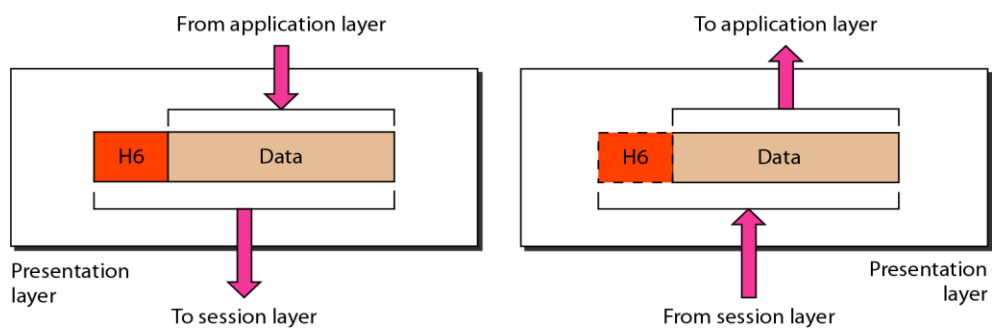
Session layer

The session layer is responsible for dialog control and synchronization.

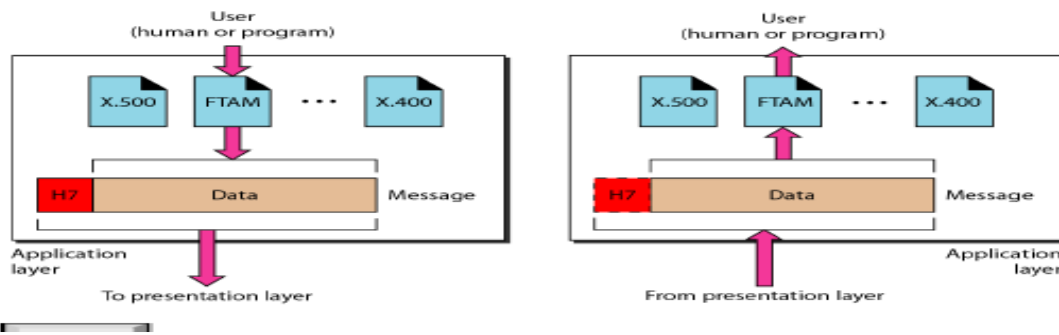


Presentation layer

The presentation layer is responsible for translation, compression, and encryption.



Application layer



The application layer is responsible for providing services to the user.

Services:

- Network virtual terminal
- Mail services
- File transfer, access, and management
- Directory services

[OR]

21.(b). Write a note on the transmission impairments

TRANSMISSION IMPAIRMENT

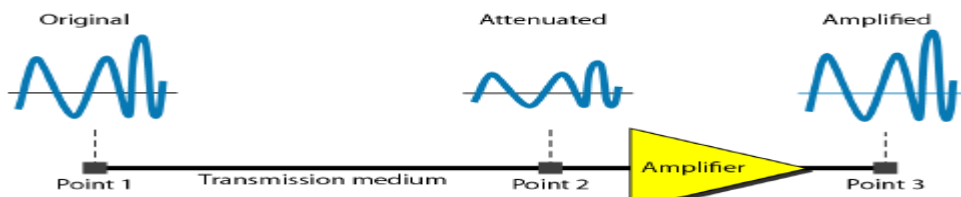
Signals travel through transmission media,, which are not perfect. The imperfection causes signal impairment. This means that the signal at the beginning of the medium is not the same as the signal at the end of the medium.. What is sent is not what is received. Three causes off impairment are attenuation, distortion, and noise.

Attenuation

Attenuation means loss of energy.

When a signal travels through a medium, it loses some of its energy so that it can overcome the resistance of the medium.

To compensate for this loss, amplifiers are used to amplify the signal.



Decibel

To show that a signal has lost or gained strength, we use the concept of the decibel (dB).

- The decibel measures the relative strengths of two signals or a signal at two different points.
- The decibel is negative if a signal is attenuated and positive if a signal is amplified.

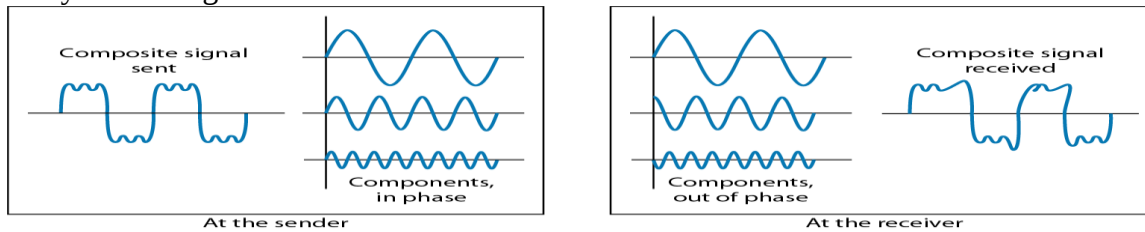
$$dB = 10 \log_{10}(P_2/P_1)$$

where P1 and P2 are the powers of a signal at points 1 and 2, respectively

Distortion

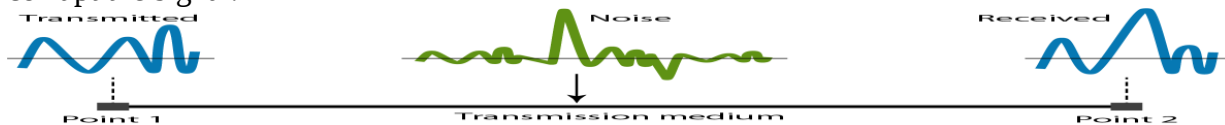
Distortion means that the signal changes its form or shape.

- Distortion occurs in a composite signal made of different frequencies.
- Each signal component has its own propagation speed through a medium and therefore its own delay in arriving at the final destination



Noise

Several types of noise such as thermal noise, induced noise, crosstalk and impulse noise may corrupt the signal.



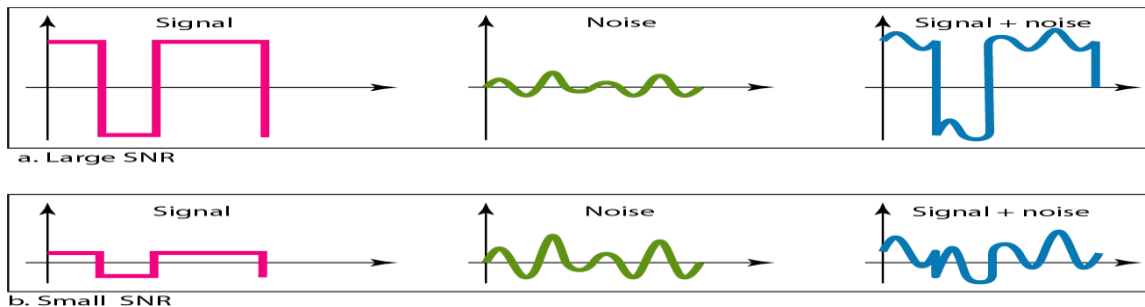
Signal-to-Noise ratio (SNR)

SNR is the statistical ratio of power of the signal to the power of the noise

- In decibels it can be expressed as follows:

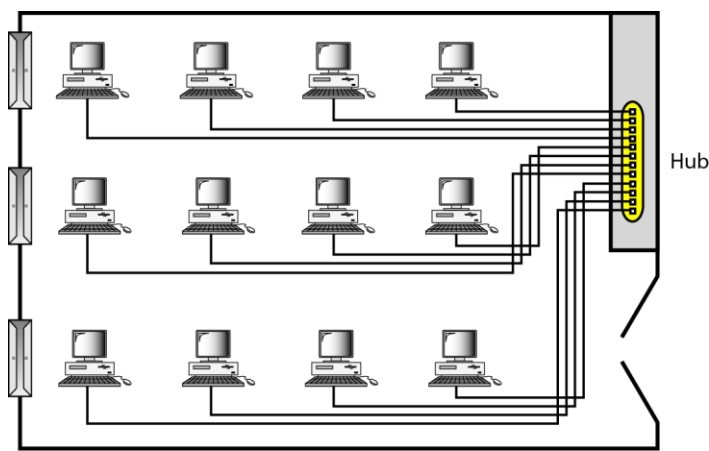
$$\text{SNRdB} = 10 \log_{10} \text{SNR}$$

Two cases of SNR: a high SNR & a low SNR



22. (a). Describe the architecture of LAN, MAN, WAN and their differences with neat sketches.

• Local Area Networks (LANs)

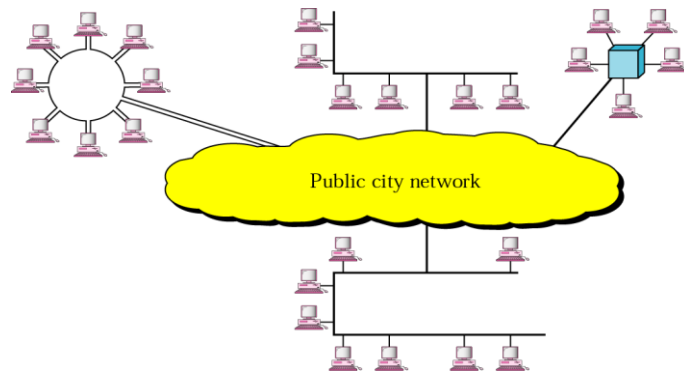


A local area network (LAN) is usually privately owned and links the devices in a single office, building, or campus. Depending on the needs of an organization and the type of technology used, a LAN can be as simple as two PCs and a printer in

someone's home office; or it can extend throughout a company and include audio and video peripherals. Currently, LAN size is limited to a few kilometers.

LANs are designed to allow resources to be shared between personal computers or workstations. The resources to be shared can include hardware (e.g., a printer), software (e.g., an application program), or data. A common example of a LAN, found in many business environments, links a workgroup of task-related computers, for example, engineering workstations or accounting PCs. One of the computers may be given a large capacity disk drive and may become a server to clients. Software can be stored on this central server and used as needed by the whole group. In this example, the size of the LAN may be determined by licensing restrictions on the number of users per copy of software, or by restrictions on the number of users licensed to access the operating system.

In addition to size, LANs are distinguished from other types of networks by their transmission media and topology. In general, a given LAN will use only one type of transmission medium. The most common LAN topologies are bus, ring, and star. Early LANs had data rates in the 4 to 16 megabits per second (Mbps) range. Today, however, speeds are normally 100 or 1000 Mbps

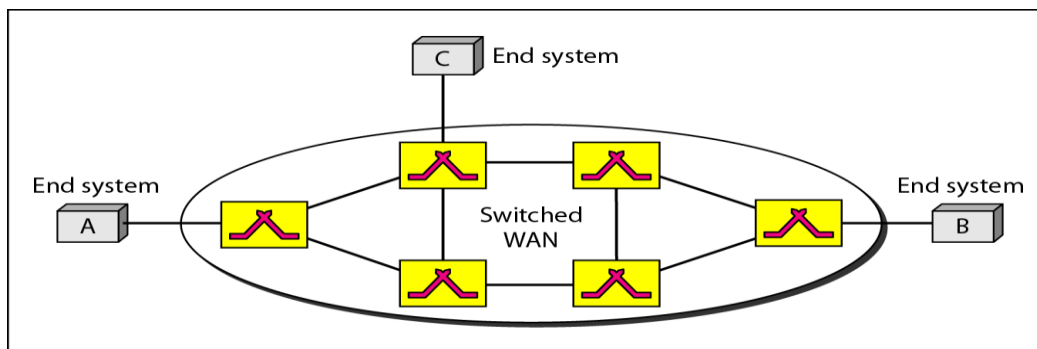


Metropolitan Area Networks (MANs)

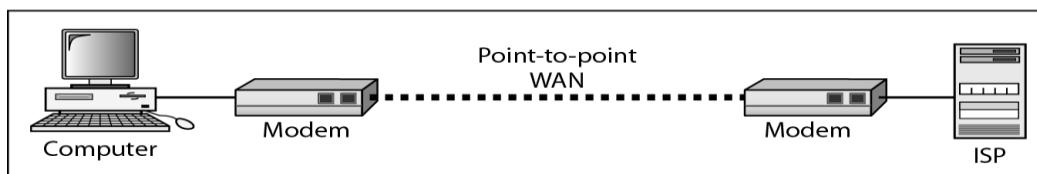
A metropolitan area network (MAN) is a network with a size between a LAN and a WAN. It normally covers the area inside a town or a city. It is designed for customers who need a high-speed connectivity, normally to the Internet, and have endpoints spread over a city or part of city. A good example of a MAN is the part of the telephone company network that can provide a high-speed DSL line to the

customer. Another example is the cable TV network that originally was designed for cable TV, but today can also be used for high-speed data connection to the Internet

• Wide Area Networks (WANs)



a. Switched WAN



b. Point-to-point WAN

A wide area network (WAN) provides long-distance transmission of data, image, audio, and video information over large geographic areas that may comprise a country, a continent, or even the whole world. A WAN can be as complex as the backbones that connect the Internet or as simple as a dial-up line that connects a home computer to the Internet. We normally refer to the first as a switched WAN and to the second as a point-to-point WAN.

- The switched WAN connects the end systems, which usually comprise a router (internetworking connecting device) that connects to another LAN or WAN.
- The point-to-point WAN is normally a line leased from a telephone or cable TV provider that connects a home computer or a small LAN to an Internet service provider (ISP). This type of WAN is often used to provide Internet access.

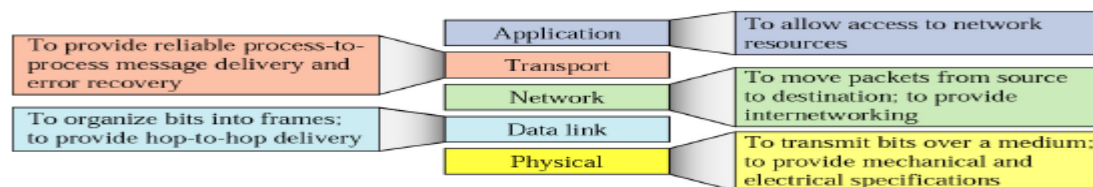
[OR]

22.(b). List out the types of addresses used in TCP/IP protocol suite and explain the association with the layers.

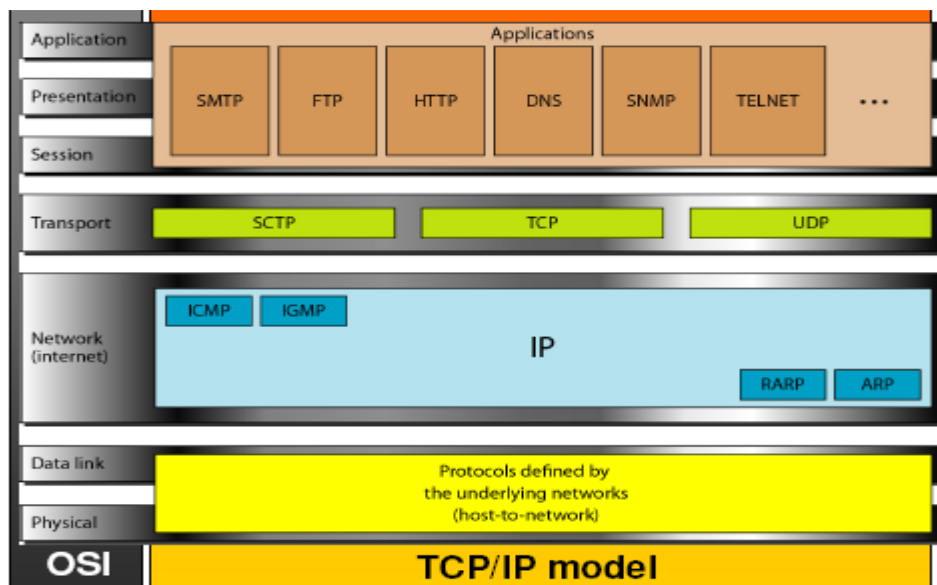
TCP/IP PROTOCOL SUITE

The layers in the TCP/IP protocol suite do not exactly match those in the OSI model. The original TCP/IP protocol suite was defined as having four layers: host-to-network, internet, transport, and application. However, when TCP/IP is compared to OSI, we can say that the TCP/IP protocol suite is made of five layers: physical, data link, network, transport, and application.

TCP/IP layers



TCP/IP and OSI model



ADDRESSING

Four levels of addresses are used in an internet employing the TCP/IP protocols: physical, logical, port, and specific.



Physical & Logical address

• Physical address

In computer networks a physical address means a MAC (Medium Access Control) address. Also known as Ethernet Hardware Address (EHA) or hardware address or **adapter address**. It is a number that acts like a name for a particular network adapter, eg. the network cards

• Logical address

—In computer networks, a logical address refers to a network layer address such as an IP address
 —An IP address (Internet Protocol address) is a unique address that certain electronic devices use in order to identify and communicate with each other on a computer network utilizing the Internet Protocol standard (IP)

Port & specific address

• Port address

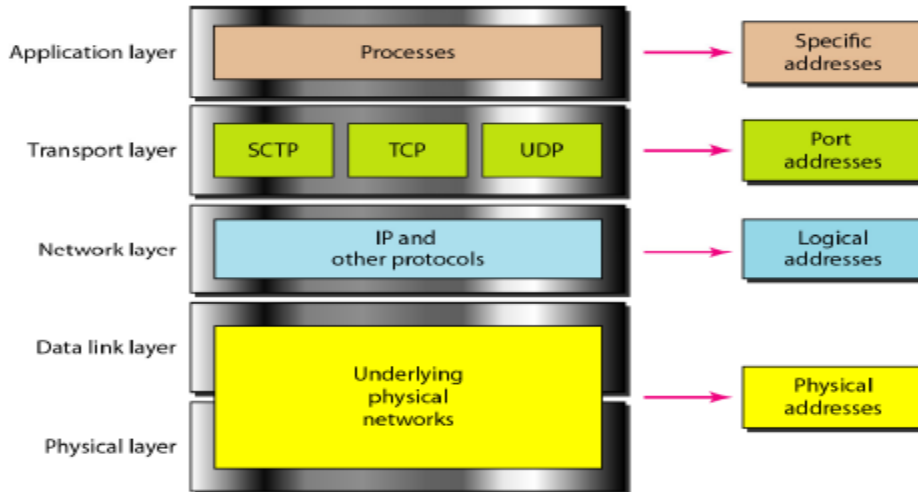
—TCP and UDP are transport protocols used for communication between computers via ports
 —The port numbers are divided into three ranges.

- The Well Known Ports are those in the range 0–1023.
- The Registered Ports are those in the range 1024–49151.
- The Dynamic and/or Private Ports are those in the range 49152–65535. These ports are not used by any defined application.

• Specific address

—This address is used by application processes

Relationship of layers-addresses in TCP/IP

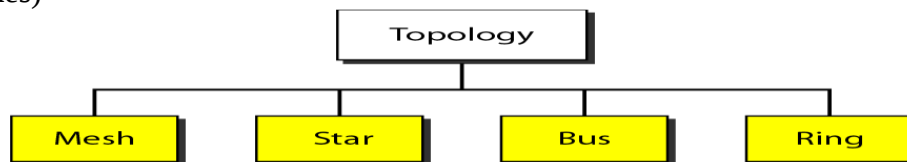


23. (a). Explain the various network topologies with a neat diagram

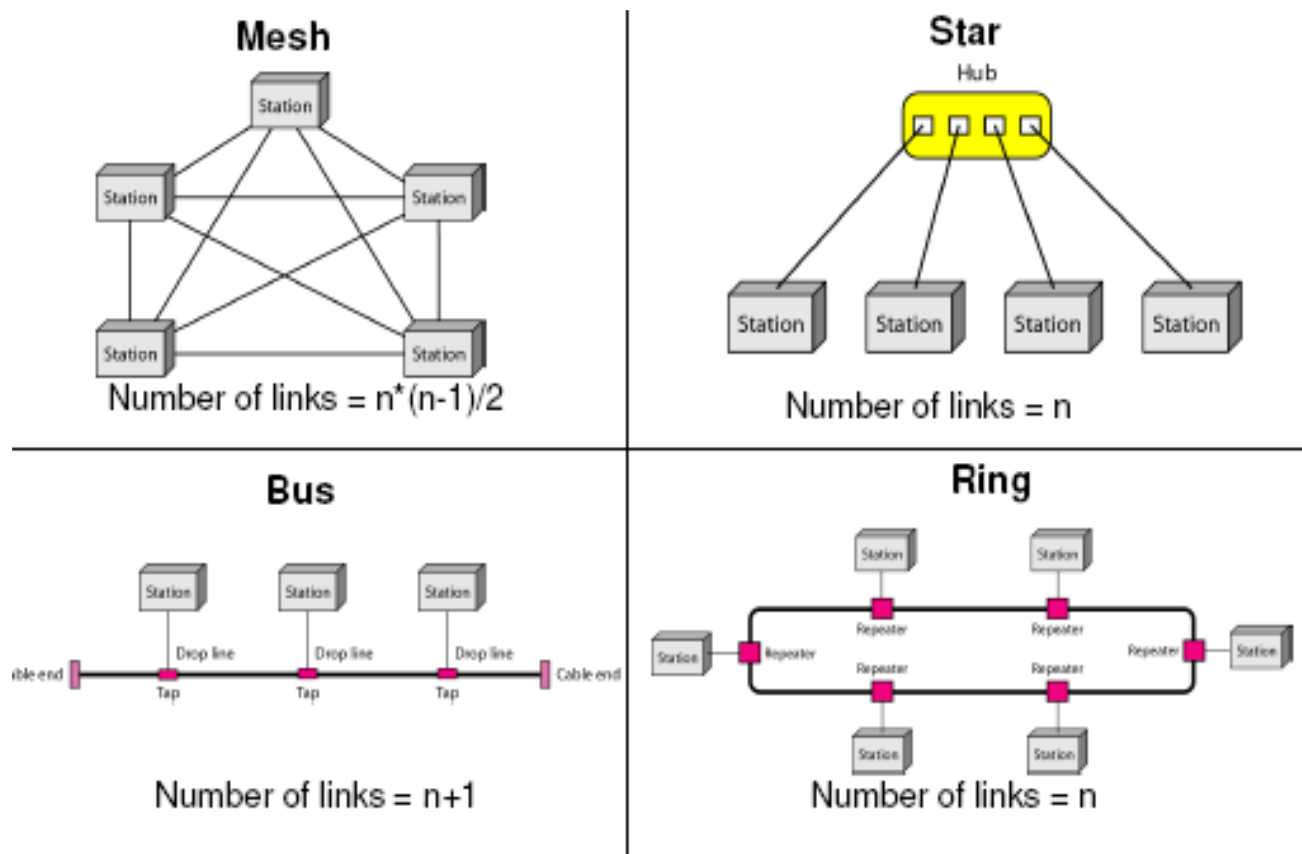
Physical Topology

Physical topology refers to the way in which a network is laid out physically.

Network topology is the geometric representation of the relationship of all the links and linking devices (nodes)



Topology categories



Mesh Topology : In a mesh topology, every device has a dedicated point-to-point link to every other device. The term dedicated means that the link carries traffic only between the two devices it connects. To find the number of physical links in a fully connected mesh network with n nodes, we first consider that each node must be connected to every other node. Node 1 must be connected to $n - 1$ nodes, node 2 must be connected to $n - 1$ nodes, and finally node n must be connected to $n - 1$ nodes. We need $n(n - 1)$ physical links. However, if each physical link allows communication in both directions (duplex mode), we can divide the number of links by 2. In other words, we can say that in a mesh topology, we need $n(n - 1) / 2$ duplex-mode links.

Advantages of Mesh Topology

- A dedicated links guarantees that each connection can carry its own data load, thus eliminating the traffic problems that can occur when links must be shared by multiple devices.
- A mesh topology is robust. If one link becomes unusable, it does not incapacitate the entire system.
- There is the advantage of privacy or security. When every message travels along a dedicated line, only the intended recipient sees it. Physical boundaries prevent other users from gaining access to messages.
- Finally, point-to-point links make fault identification and fault isolation easy. Traffic can be routed to avoid links with suspected problems. This facility enables the network manager to discover the precise location of the fault and aids in finding its cause and solution.

Disadvantages of Mesh Topology

The main disadvantages of a mesh are related to the amount of cabling and the number of I/O ports required.

- First, because every device must be connected to every other device, installation and reconnection are difficult.

- Second, the sheer bulk of the wiring can be greater than the available space (in walls, ceilings, or floors) can accommodate.
- Finally, the hardware required to connect each link (I/O ports and cable) can be prohibitively expensive.
- For these reasons a mesh topology is usually implemented in a limited fashion, for example, as a backbone connecting the main computers of a hybrid network that can include several other topologies.

Star Topology

In a star topology, each device has a dedicated point-to-point link only to a central controller, usually called a hub. The devices are not directly linked to one another. Unlike a mesh topology, a star topology does not allow direct traffic between devices. The controller acts as an exchange: If one device wants to send data to another, it sends the data to the controller, which then relays the data to the other connected device.

Advantages of Star Topology

- A star topology is less expensive than a mesh topology. In a star, each device needs only one link and one I/O port to connect it to any number of others which also makes it easy to install and reconfigure.
- Other advantages include robustness. If one link fails, only that link is affected. All other links remain active. This factor also lends itself to easy fault identification and fault isolation.
- As long as the hub is working, it can be used to monitor link problems and bypass defective links.

Disadvantages of Star Topology

- One big disadvantage of a star topology is the dependency of the whole topology on one single point, the hub. If the hub goes down, the whole system is dead.
- Although a star requires far less cable than a mesh, each node must be linked to a central hub. For this reason, often more cabling is required in a star than in some other topologies (such as ring or bus).

Bus Topology

The preceding examples all describe point-to-point connections. A **bus topology**, on the other hand, is multipoint. One long cable acts as a **backbone** to link all the devices in a network. Nodes are connected to the bus cable by drop lines and taps. A drop line is a connection running between the device and the main cable. A tap is a connector that either splices into the main cable or punctures the sheathing of a cable to create a contact with the metallic core. As a signal travels along the backbone, some of its energy is transformed into heat. Therefore, it becomes weaker and weaker as it travels farther and farther. For this reason there is a limit on the number of taps a bus can support and on the distance between those taps.

Advantages of Bus Topology

- Advantages of a bus topology include ease of installation. Backbone cable can be laid along the most efficient path, then connected to the nodes by drop lines of various lengths. In this way, a bus uses less cabling than mesh or star topologies.
- In a star, for example, four network devices in the same room require four lengths of cable reaching all the way to the hub. In a bus, this redundancy is eliminated. Only the backbone cable stretches through the entire facility. Each drop line has to reach only as far as the nearest point on the backbone.

Disadvantages

- Disadvantages include difficult reconnection and fault isolation.

- A bus is usually designed to be optimally efficient at installation. It can therefore be difficult to add new devices.
- Signal reflection at the taps can cause degradation in quality. This degradation can be controlled by limiting the number and spacing of devices connected to a given length of cable.
- Adding new devices may therefore require modification or replacement of the backbone.
- In addition, a fault or break in the bus cable stops all transmission, even between devices on the same side of the problem.

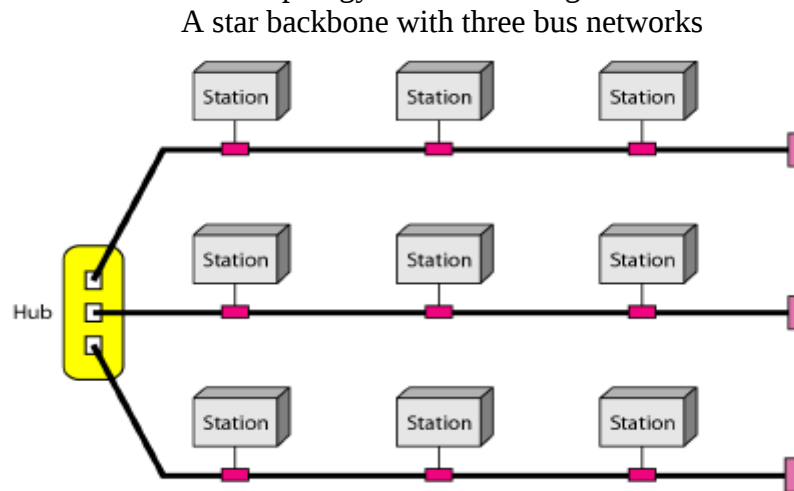
Ring Topology

Ring Topology In a ring topology, each device has a dedicated point-to-point connection with only the two devices on either side of it. A signal is passed along the ring in one direction, from device to device, until it reaches its destination. Each device in the ring incorporates a repeater. When a device receives a signal intended for another device, its repeater regenerates the bits and passes them along.

A ring is relatively easy to install and reconfigure. Each device is linked to only its immediate neighbors (either physically or logically). To add or delete a device requires changing only two connections. The only constraints are media and traffic considerations (maximum ring length and number of devices). In addition, fault isolation is simplified. Generally in a ring, a signal is circulating at all times. If one device does not receive a signal within a specified period, it can issue an alarm. The alarm alerts the network operator to the problem and its location. However, unidirectional traffic can be a disadvantage. In a simple ring, a break in the ring (such as a disabled station) can disable the entire network. This weakness can be solved by using a dual ring or a switch capable of closing off the break.

A hybrid topology

A network can be hybrid. For example, we can have a main star topology with each branch connecting several stations in a bus topology as shown in Figure below.



[OR]

23 (b). Illustrate the working of frequency division multiplexing with a neat diagram.

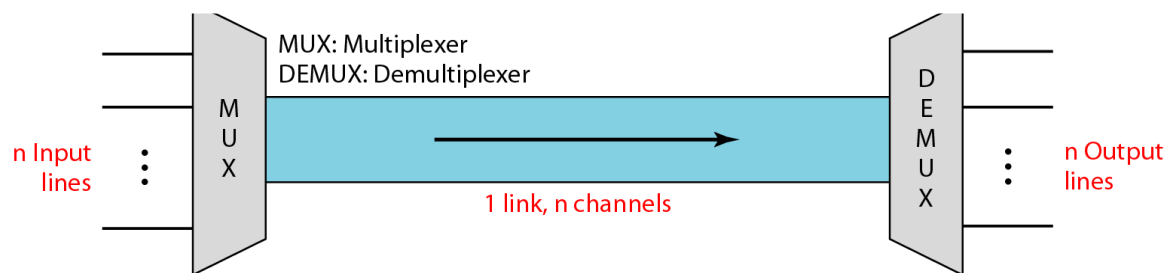
MULTIPLEXING

Whenever the bandwidth of a medium linking two devices is greater than the bandwidth needs of the devices, the link can be shared. Multiplexing is the set off techniques that allows the

simultaneous transmission of multiple signals across a single data link. As data and telecommunications use increases, so does traffic.

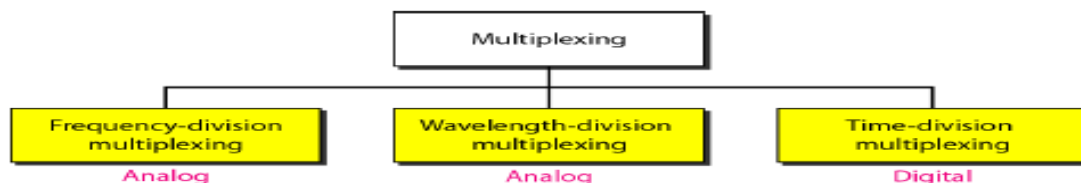
Dividing a link into channels

In a multiplexed system, n lines share the bandwidth of one link. Figure shows the basic format of a multiplexed system. The lines on the left direct their transmission streams to a multiplexer (MUX), which combines them into a single stream (many-to-one). At the receiving end, that stream is fed into a demultiplexer (DEMUX), which separates the stream back into its component transmissions (one-to-many) and directs them to their corresponding lines. In the figure, the word link refers to the physical path. The word channel refers to the portion of a link that carries a transmission between a given pair of lines. One link can have many (n) channels.



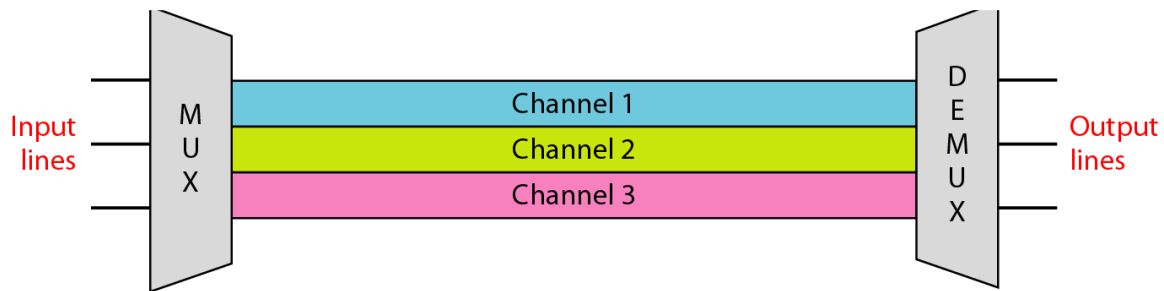
Categories of multiplexing

There are three basic multiplexing techniques: frequency-division multiplexing, wavelength-division multiplexing, and time-division multiplexing. The first two are techniques designed for analog signals, the third, for digital signals



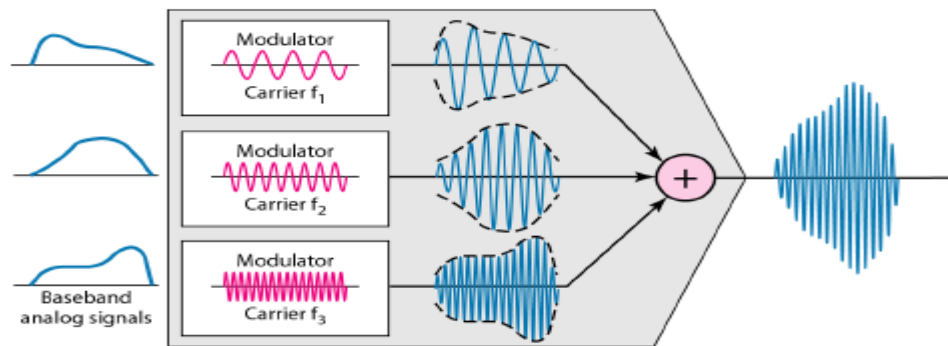
Frequency Division Multiplexing (FDM)

Frequency-division multiplexing (FDM) is an analog technique that can be applied when the bandwidth of a link (in hertz) is greater than the combined bandwidths of the signals to be transmitted. In FDM, signals generated by each sending device modulate different carrier frequencies. These modulated signals are then combined into a single composite signal that can be transported by the link. Carrier frequencies are separated by sufficient bandwidth to accommodate the modulated signal. These bandwidth ranges are the channels through which the various signals travel. Channels can be separated by strips of unused bandwidth-guard bands-to prevent signals from overlapping. In addition, carrier frequencies must not interfere with the original data frequencies. Figure gives a conceptual view of FDM. In this illustration, the transmission path is divided into three parts, each representing a channel that carries one transmission.



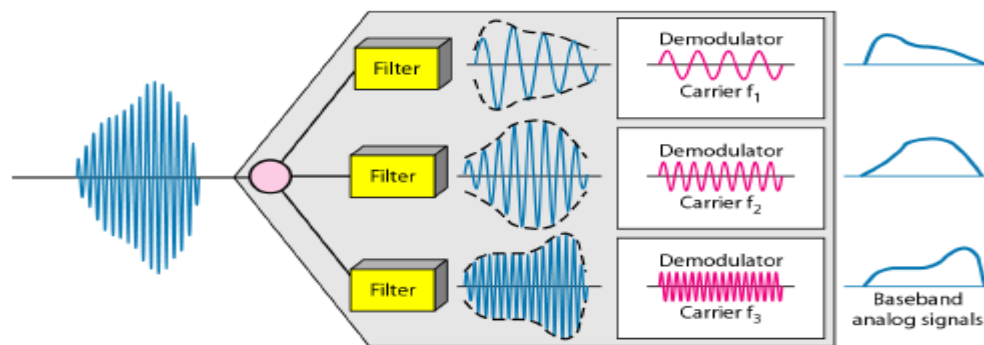
Multiplexing Process

Figure below is a conceptual illustration of the multiplexing process. Each source generates a signal of a similar frequency range. Inside the multiplexer, these similar signals modulates different carrier frequencies. The resulting modulated signals are then combined into a single composite signal that is sent out over a media link that has enough bandwidth to accommodate it.



Demultiplexing Process

The demultiplexer uses a series of filters to decompose the multiplexed signal into its constituent component signals. The individual signals are then passed to a demodulator that separates them from their carriers and passes them to the output lines.



A very common application of FDM is AM and FM radio broadcasting. Radio uses the air as the transmission medium. A special band from 530 to 1700 kHz is assigned to AM radio. All radio stations need to share this band. Each AM station needs 10kHz of bandwidth. Each station uses a different carrier frequency, which means it is shifting its signal and multiplexing. The signal that goes to the air is a combination of signals. A receiver receives all these signals, but filters (by tuning) only the one which is desired. Without multiplexing, only one AM station could broadcast to the common link, the air. The situation is similar in FM broadcasting. However, FM has a wider band of 88 to 108MHz because each station needs a bandwidth of 200 kHz. Another common use of FDM is in television broadcasting. Each TV channel has its own bandwidth of 6 MHz.