



S.Y.B.Sc. (C. S.)
SEMESTER - IV (CBCS)

COMPUTER NETWORK

SUBJECT CODE: USCS403

Prof. (Dr.) D. T. Shirke

Offg. Vice Chancellor
University of Mumbai, Mumbai

Prin. Dr. Ajay Bhamare

Offg. Pro Vice-Chancellor,
University of Mumbai

Prof. Prakash Mahanwar

Director,
IDOL, University of Mumbai

Programme Co-ordinator

: Shri Mandar Bhanushe

Head, Faculty of Science and Technology IDOL,
University of Mumbai – 400098

Course Co-ordinator

: Ms. Mitali Vijay Shewale

Doctoral Researcher,
Veermata Jijabai Technological Institute
HR Mahajani road, Matunga, Mumbai

Editor

: Mr. Anish Raut

Assistant Manager,
Dahua Technology India Pvt. Ltd., Mumbai.

Course Writers

: Dr. Vithal Nanabhau Patange

Assistant Professor ,
Shri Shivaji Education Society
Amravati Science college, Pauni, Bhandara.

: Mr. Mohamed Hasan phudinawala

Assistant Professor,
Royal College of Arts, Science and Commerce
Mumbai.

: Mr. Vijay Kothawade

Assistant Professor,
Madhavi College of Science and Commerce,
Dombivali East, Thane.

: Ms. Mitali Vijay Shewale

Doctoral Researcher,
Veermata Jijabai Technological Institute
HR Mahajani road, Matunga, Mumbai

: Ms. Trupti Kulkarni Kaujalgi

HOD IT CS,
ICLES' Motilal Jhunjhunwala College, Vashi.

August 2023, Print - 1

Published by : Director,

Institute of Distance and Open Learning,
University of Mumbai,
Vidyanagari, Mumbai - 400 098.

DTP composed and Printed by: Mumbai University Press

CONTENTS

Unit No.	Title	Page No.
1	Introduction of Data Communication	1
2	Introduction of Network Models	43
3	Digital and Analog Transmission	79
4	Transmission Modes	95
5	Analog to Analog Conversion and Multiplexing	109
6	Transmission Media and Switching	129
7	Introduction to Data Link Layer	156
8	Media Access Control (Mac)	166
9	Connecting Devices	187
10	Introduction to Network Layer	202
11	Unicast Routing	221
12	Next Generation IP	243
13	Introduction to the Transport Layer	261
14	Introduction to Application Layer	283
15	Standard Client-Server Protocols	298

S.Y.B.Sc. (C. S.)
SEMESTER - IV (CBCS)
COMPUTER NETWORK
SYLLABUS

Course: USCS403	TOPICS (Credits :02 Lectures/Week:03) Computer Networks	
Objectives: In this era of Information, its computation and its exchange techniques, Learner should be able to conceptualize and understand the framework and working of communication networks. And on completion, will be able to have a firm grip over this very important segment of Internet. Expected Learning Outcomes : 1. Learner will be able to understand the concepts of networking, which are important for them to be known as a '<i>networking professionals</i>'. 2. Useful to proceed with industrial requirements and International vendor certifications.		
Unit I	Introduction Network Models: Introduction to data communication, Components, Data Representation, Data Flow, Networks, Network Criteria, Physical Structures, Network types, Local Area Network, Wide Area Network, Switching, The Internet, Accessing the Internet, standards and administration Internet Standards. Network Models, Protocol layering, Scenarios, Principles of Protocol Layering, Logical Connections, TCP/IP Protocol Suite, Layered Architecture, Layers in	15L
	the TCP/IP Protocol Suite, Encapsulation and Decapsulation, Addressing, Multiplexing and Demultiplexing. Detailed introduction to Physical Layer, Detailed introduction to Data-Link Layer, Detailed introduction to Network Layer, Detailed introduction to Transport Layer, Detailed introduction to Application Layer. Data and Signals, Analog and Digital Data, Analog and Digital Signals, Sine Wave Phase, Wavelength, Time and Frequency Domains, Composite Signals, Bandwidth, Digital Signal, Bit Rate, Bit Length, Transmission of Digital Signals, Transmission Impairments, Attenuation, Distortion, Noise, Data Rate Limits, Performance, Bandwidth, Throughput, Latency (Delay)	

Unit II	Introduction to Physical Layer and Data-Link Layer: Digital Transmission digital-to-digital conversion, Line Coding, Line Coding Schemes, analog-to-digital conversion, Pulse Code Modulation (PCM), Transmission Modes, Parallel Transmission, Serial Transmission. Analog Transmission, digital-to-analog Conversion, Aspects of Digital-to-Analog Conversion, Amplitude Shift Keying, Frequency Shift Keying, Phase Shift Keying, analog-to-analog Conversion, Amplitude Modulation (AM), Frequency Modulation (FM), Phase Modulation (PM), Multiplexing, Frequency-Division Multiplexing, Wavelength-Division Multiplexing, Time-Division Multiplexing. Transmission Media, Guided Media, Twisted-Pair Cable, Coaxial Cable, Fiber-Optic Cable. Switching, Three Methods of Switching , Circuit Switched Networks, Packet Switching, Introduction to Data-Link Layer, Nodes and Links, Services, Two Sub-layers, Three Types of addresses, Address Resolution Protocol (ARP). Error Detection and Correction, introduction, Types of Errors, Redundancy, Detection versus Correction,	15L
Unit III	Network layer, Transport Layer Media Access Control (MAC), random access, CSMA, CSMA/CD, CSMA/CA, controlled access, Reservation, Polling, Token Passing, channelization, FDMA, TDMA, CDMA. Connecting Devices and Virtual LANs, connecting devices, Hubs, Link-Layer	15L

INTRODUCTION OF DATA COMMUNICATION

Unit Structure:

- 1.0 Objectives
- 1.1 Introduction to data communication
 - 1.1.1 Characteristics of Data Communication
 - 1.1.2 Components of Data Communication
- 1.2 Data Representation
- 1.3 Data flow
- 1.4 Networks
- 1.5 Physical Structure
 - 1.5.1 Type of Connection (Line configuration)
 - 1.5.2 Physical Topology
- 1.6 Network Type
 - 1.6.1 LAN (Local Area Network)
 - 1.6.2 PAN (Personal Area Network)
 - 1.6.3 MAN (Metropolitan Area Network)
 - 1.6.4 WAN (Wide Area Network)
- 1.7 Switching
 - 1.7.1 Circuit Switching
 - 1.7.2 Message Switching
 - 1.7.3 Packet Switching
- 1.8 The Internet
 - 1.8.1 Internet Accessing
- 1.9 Internet Standard
 - 1.9.1 Internet Administration
- 1.10 List of References
- 1.11 Exercises

1.0 OBJECTIVES

After going through this unit, you will be able to:

- Define data communication, Networks study, Network Criteria and Structure approach
- State the common Network types
- Describe the basic in The Internet and its standards
- Classify Network Model and TCP/IP systems
- Explain what is Encapsulation, Decapsulation and Multiplexing Concept
- Illustrate the DLL, Network layer, Transport layer and Application Layer

1.1 INTRODUCTION TO DATA COMMUNICATION

Humans are the only creatures on earth who are capable of speak with every different thru the medium of language. But human beings take this present to another volume. Distance, time, and physical existence of the individual don't count in conversation nowadays because they build a communicate gadget thru which they are able to communicate or proportion records like pictures, videos, text, files, and so forth. With their loved ones anytime anywhere. Communication is defined as a process wherein multiple pc transfers data, commands to each different and for sharing sources. Or in different words, conversation is a method or act wherein we can send or receive records. A community of computers is defined as an interconnected series of self-reliant computers. Independent manner no pc can start stop or manipulate another laptop. Statistics verbal exchange is a manner of changing information or facts. In case of pc networks this alternate is achieved among devices over a transmission medium. This system entails a verbal exchange system that's made of hardware and software. The hardware element entails the sender and receiver devices and the intermediate gadgets through which the statistics passes. The software program component involves certain guidelines which specify what's to be communicated, how it's far to be communicated and while. It's also called as a Protocol.

1.1.1 Characteristics of Data Communication

The effectiveness of any facts communications gadget depends upon the following 4 fundamental traits:

- 1 **Delivery:** The facts should be introduced to the best destination and accurate consumer.
- 2 **Accuracy:** The conversation machine ought to deliver the data correctly, without introducing any mistakes. The information might

also get corrupted all through transmission affecting the accuracy of the delivered information.

- 3 **Timeliness:** Audio and Video facts needs to be added in a timely way with none put off; any such facts shipping is known as actual time transmission of facts.
- 4 **Jitter:** It is the version within the packet arrival time. Un-even Jitter may additionally have an effect on the timeliness of information being transmitted.

1.1.2 Components of Data Communication

A verbal exchange machine is made of the following components:

1. **Message:** A message is a chunk of information that is to be transmitted from one person to any other. it could be a text document, an audio report, a video record, and so on.
2. **Sender:** its miles in reality a tool that sends facts messages. It could be a laptop, mobile, cellphone, laptop, video digital, or workstation, and many others.
3. **Receiver:** it is a device that gets messages. it is able to be a computer, smartphone mobile, notebook, and so forth.
4. **Transmission Medium / Communication Channels:** verbal exchange channels are the medium that connect or greater workstations. Workstations can be linked through both wired media and Wi-Fi media.
5. **Set of rules (Protocol):** Set of policies (Protocol): when someone sends the information (The sender), it has to be comprehensible to the receiver additionally in any other case it's miles meaningless. For example, Suresh sends a message to Vithal. If Suresh writes in Hindi and Vithal cannot understand Hindi, it is a meaningless communique.

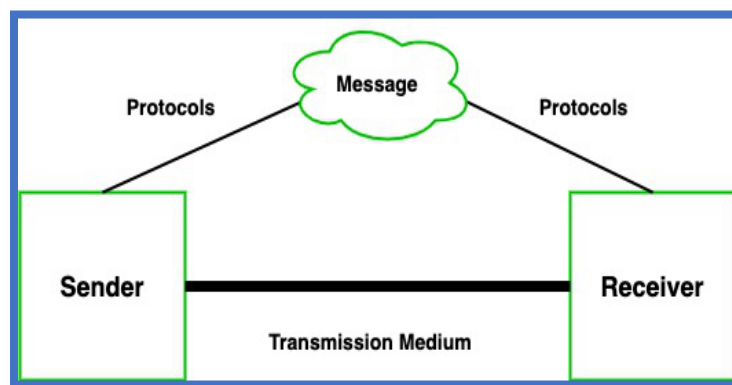


Figure: Data Communication

Therefore, there is a certain set of rules (protocols) that governs every computer connected to the Internet, namely:

- **TCP (Transmission Control Protocol):** It is responsible for splitting messages into packets on the source computer and reassembling the received packet on the destination or receiving computer. It also ensures that the packets have information about the source of the message data, the destination of the message data, the order in which the message data should be reassembled, and checks that the message was sent correctly to specific destinations.
- **IP (Internet Protocol):** IP is responsible for processing the address of the destination computer so that each packet is sent to the correct destination.

1.2 DATA REPRESENTATION

Data is a collection of raw facts that are processed to derive information. There can be different forms in which data can be represented. Some of the forms of data used in communication are as follows:

1. **Text :** The text contains a combination of lowercase and uppercase alphabets. It is stored as a bit pattern. Prevailing encoding system: ASCII, Unicode
2. **Numbers :** Numbers include a combination of digits from 0 to 9. It is stored as a bit pattern. Prevailing encoding system: ASCII, Unicode
3. **Pictures :** A picture is worth a thousand words is a very famous saying. In computers, images are stored digitally. A pixel is the smallest element of an image. Simply put, an image or picture is a matrix of pixel elements. Pixels are represented in the form of bits. Depending on the type of image (black n white or color), each pixel would require a different number of bits to represent the pixel value. The size of an image depends on the number of pixels (also called resolution) and the bit pattern used to represent the value of each pixel. Example: if an image is pure black and white (two-color), each pixel can be represented by a value of either 0 or 1, so an image composed of 10 x 10 pixel elements would only require 100 bits to store in memory. On the other hand, an image that contains gray may require 2 bits to represent each pixel value (00 - black, 01 - dark gray, 10 - light gray, 11 - white). So the same 10 x 10-pixel image would now require 200 bits of memory to store. Commonly used image formats: jpg, png, bmp, etc.
4. **Video :** Video refers to the transmission of data in the form of an image or movie

As we know data communication is communication in which we can send or receive data from one device to another. Data communication is divided into three types:

1. Simplex Communication:

This is a one-way communication, or one-way communication, where one device only receives and the other only sends data, and the device uses its entire capacity for transmission. For example, IoT, data input using keyboard, music output using speaker, etc.

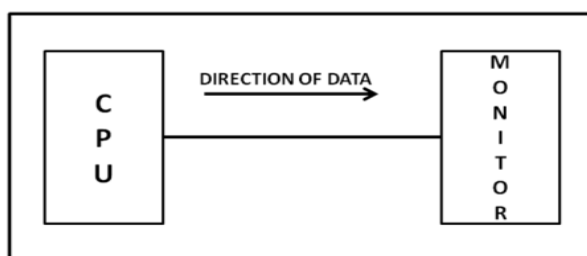


Figure: Simplex Communication

2. Half Duplex communication:

It is a two-way communication or we can say it is a two-way communication in which both devices can send and receive data but not at the same time. When one device sends data, the other device only receives and vice versa. For example, a walkie-talkie.

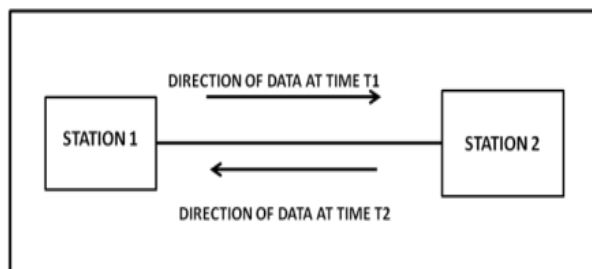


Figure: Half Duplex communication

3. Full-duplex communication:

It's miles a two-way communication or we are able to say that it is a bidirectional communicate in which both the devices can ship and obtain information at the equal time. as an instance, mobile telephones, landlines, and so forth.

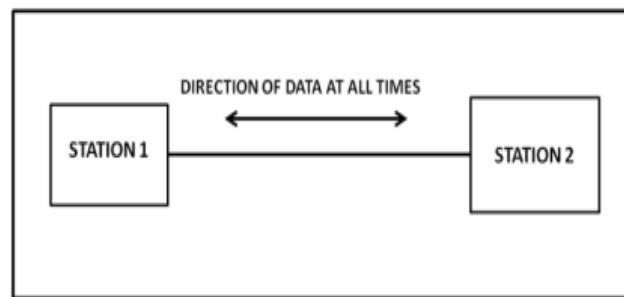


Figure: Full-duplex communication

1.4 NETWORKS

A network consists of or extra computer systems which might be linked with the intention to proportion resources (which includes printers and CDs), trade documents, or allow digital communications. The computers on a network may be connected thru cables, cellphone strains, radio waves, satellites, or infrared mild beams. Maximum networks use allotted processing, wherein a undertaking is split among a couple of computer systems. Instead of one unmarried big device being liable for all aspects of method, separate computers (normally a private pc or pc) handle a subset.

1.4.1 Network Criteria

1. Performance
2. Reliability
3. Security

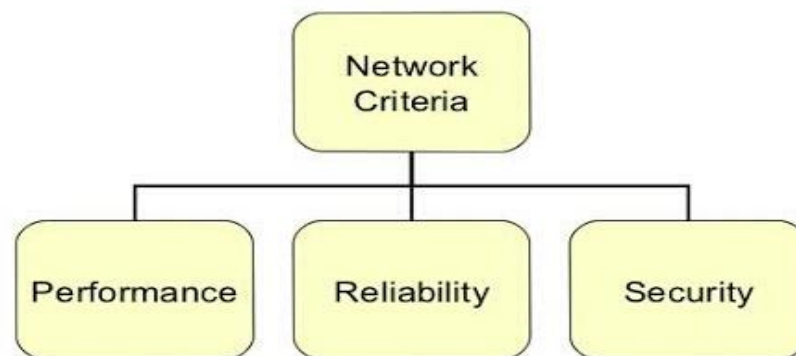


Figure Network Criteria

Communities need to be capable of meet a sure range of standards. The most vital of those are overall performance, reliability, and protection.

1. Performance: -

Overall performance is frequently evaluated by using networking metrics: transit time, reaction time, throughput and delay. We regularly need greater throughputs and much less postpone. However, these two criteria are regularly contradictory. If we attempt to ship greater statistics to the community, we may growth throughput but we boom the delay due to site visitor's congestion inside the network.

Transit time is the amount of time required for a message to tour from one tool to every other. Reaction time is the elapsed time between an inquiry and a reaction. The performance of a community relies upon on a range of factors, inclusive of:

- The number of users
- The type of transmission medium
- Connected hardware
- Software

2. **Reliability:** -

In addition to accuracy of delivery, community reliability is measured through the frequency of failure, the time it takes a hyperlink to get over a failure and the network's robustness in a disaster. "Catastrophe" - community must be protected from catastrophic events which include hearth, earthquake, or theft.

3. **Securities:** -

Community safety troubles encompass protecting facts from unauthorized access, protective facts from harm and improvement, and implementing policies and approaches for recovery from breaches and data losses

1.5 PHYSICAL STRUCTURES

In physical structures, we need to define some network attributes.

1 Type of Connection

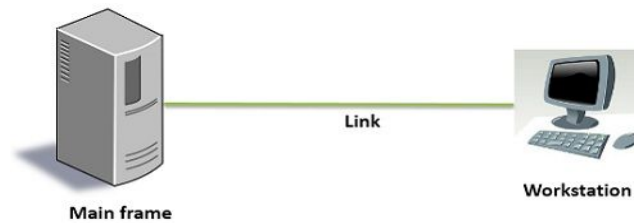
2 Physical [Topology](#)

1.5.1 Type of Connection (Line configuration)

There are two possible types of connections:

1. **Point-to-Point**

A factor-to-point connection affords a dedicated link between devices. The complete ability of the link is reserved for transmission among the ones two gadgets. Maximum factor-to-factor connections use a real duration of wire or cable to attach the 2 ends, but different alternatives, inclusive of microwave or satellite links, also are feasible. When you change TV channels by way of infrared far flung control, you are establishing a point-to-point connection between the far flung manipulates and the TV's manage machine.



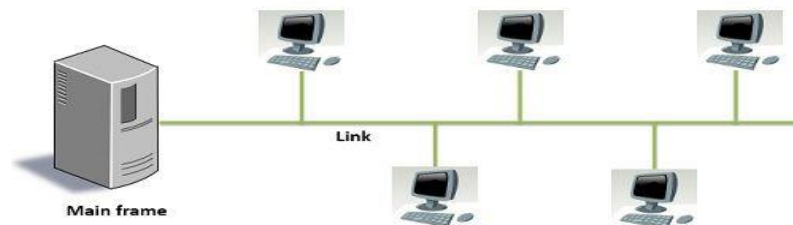
Point-to-point Connection

Example

Communication among television and far flung control for converting the channels through infrared ray. A pc connected by means of a cellphone line.

2. Multipoint

A multipoint (multi drop) connection is one wherein more than unique gadgets percentage a unmarried hyperlink. In a multipoint environment, the potential of the channel is shared, both spatially and temporally.



Multipoint Connection

Example

ATM is an example of Multipoint connection

Two types of Multipoint Connections

1. **Spatial Sharing:** If several computers can share the link simultaneously, it's called spatially shared line configuration.
2. **Temporal (Time) Sharing:** If users must take turns, it is a timeshared connection.

1.5.2 Physical Topology

The term physical topology refers back to the way in which a network is laid out bodily two or extra gadgets connect or extra links form a topology. The topology of a community is the geometric representation of the relationship of all the hyperlinks and linking devices (typically known as nodes) to each other.

There are four basic topologies possible:

1. Mesh Topology
2. Star Topology

3. Bus Topology and
4. Ring Topology
1. **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.

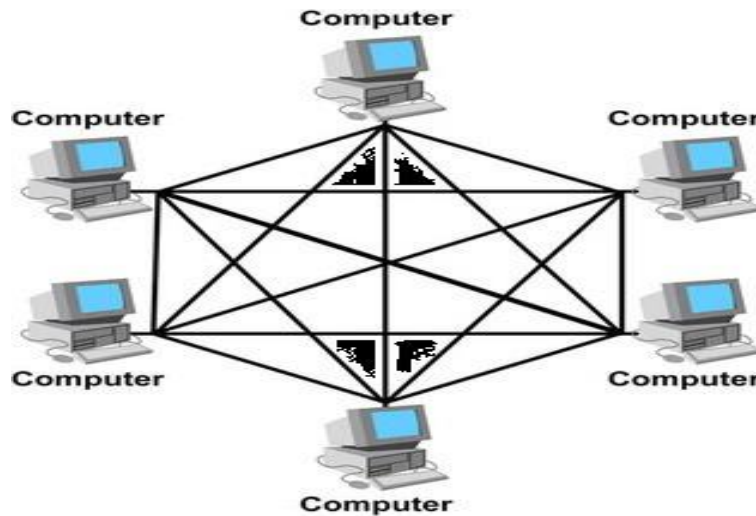


Figure Mesh Topology

To discover the quantity of bodily links in a fully related mesh community with n nodes, we first don't forget that every node needs to be related to each other node. Node 1 have to be related to $n - 1$ nodes, node 2 need to be linked to $n - 1$ nodes, and sooner or later node n ought to be linked to $n - 1$ nodes. We need $n(n - 1)$ bodily hyperlinks. However, if each bodily link allows communication in both directions (duplex mode), we are able to divide the quantity of links by 2.

In different phrases, we can say that during a mesh topology, we need $n(n - 1) / 2$ duplex-mode hyperlinks.

Characteristics of a mesh topology

- In a mesh topology, on the premise of the provision of connection between nodes, all devices determine the direction of the information float and work as a router.
- If a ruin occurs in a section of cable, the traffic load of the network is redistributed between all nodes, which preserve the supply of the network.
- It's far a type of community topology that gives redundant hyperlinks throughout the community, however it is hardly ever used due to work worried in having a network and large fee, as the community additives are immediately linked to each other aspect

- Furthermore, for installing partial mesh topologies, the mesh network setup is ideal as it balances the want for redundancy as well as fee.

Types of mesh topology

The total mesh and partly-related mesh are sorts of the mesh topology, that are mentioned under:

1. **Full Mesh Topology:** In a full mesh topology, all of the gadgets are related with all different devices. Full Mesh is a community where every node could have an $n-1$ wide variety of connections if there are n numbers of nodes available inside the community. A complete mesh topology is typically reserved for community backbones, which offers a splendid deal of redundancy. However, it could be excessively high-priced to put in force. Even though it can be steeply-priced to put in force, it affords an advantage that if one of the nodes goes down, the site visitor's load of the community is redistributed to different nodes.
2. **Partial Mesh Topology:** only some nodes, within the partial mesh topology, are connected with all of the other nodes. It method that during this network, it isn't always necessary to connect all the gadgets are connected with other. Compared to full mesh topology, it's far much less high-priced, and it offers fundamental redundancy to control the failure of any nodes. The partial mesh topology is used in peripheral networks thru which they work with a complete-mesh spine in tandem.

Protocols in Mesh Topology

Protocols are the set of regulations; these are in shape in layer three of the OSI version that defines requirements for statistics communicate among two nodes. The 3 kinds of protocols, Proactive, Hybrid, and Reactive, are used inside the mesh topology. Every protocol performs a crucial position in coping with the network thru its very own capabilities and influences performance and scalability.

- **Proactive protocol:** in the network direction, this protocol gives constant self-tracking of the nodes with the help of comments from the nodes. Additionally, if any node gets fails, it allows reroute the community course. It ensures maximum uptime of the community, and gives quick restoration from any failure, and presents robust performance. It consumes greater sources in a dynamic surroundings and has a chance of collusion, however within the static surroundings, it is much ideal in which the network path does now not frequently change. Consequently, its miles higher to use this protocol for the right environment.
- **Hybrid Protocol:** This protocol offers the exceptional aggregate based totally on the surroundings and the verbal exchange desires and uses reactive methods and the traits of the

proactive protocol. The value of community operations is optimized via this technique.

- **Reactive Protocol:** This protocol helps to decide the network direction on the time of request for records transmission. It determines the premier direction and scans with the help of the complete community. It has higher scalability that makes it able to righting in shape for dynamic surroundings.

Advantages of Mesh Topology

There are various benefits of mesh topology, which are discussed below:

- **Scalable:** In a mesh topology, every node acts like Router. but, there are no exclusive Routers. It easy to add an extra node in this topology and connect it to the network. moreover, to scale up the community, more effort isn't always wished.
- **Robust:** If any single node receives fails in the system, the network availability will now not be affected and could be maintained. And, sturdy functions are protected on this topology to triumph over any situation. Moreover, this topology has no general shutdown.
- **Decrease value:** It needs less funding in infrastructure as it's far a rather decentralized machine. Additionally, to manipulate the community, there aren't any valuable servers.
- **Redundancy:** This topology constructed a whole lot of redundancy to preserve maximum uptime and gives several paths to reach the destination.
- As any failure does not disrupt its approaches, consequently, information transmission is greater consistent.
- This topology offers the advantage of adding a new node without disrupting the facts transmissions.
- **Disadvantages of Mesh Topology**

There are numerous demerits of mesh topology that are given under:

- **Complex:** in this topology, every node works as a router that will increase complexity.
- **Planning:** This topology gives flexibility and scalability because it allows the addition of new gadgets within the current community, in order to must make certain uniform latency across all nodes. For this reason, community planning makes it a touch hard.

- Strength intake: all of the time in this network setup; each node ought to remain energetic that caused excessive energy intake and growth the load.
- As compared to other network topologies, along with point to factor, megastar, bus, the fee of mesh topology is excessive.
- With the mesh topology, maintenance is challenging, and all nodes need an in addition utility price to think about.
- Furthermore, inside the mesh topology, the installation is tons hard.

2. Star Topology

In a celeb topology, each device has a devoted factor-to-point hyperlink best to a central controller, generally called a hub. The gadgets are not directly linked to one another. Not like a mesh topology, a celebrity topology does not permit direct traffic between gadgets. The controller acts as an exchange: If one device desires to send statistics to any other, it sends the information to the controller, which then relays the information to the alternative connected tool.

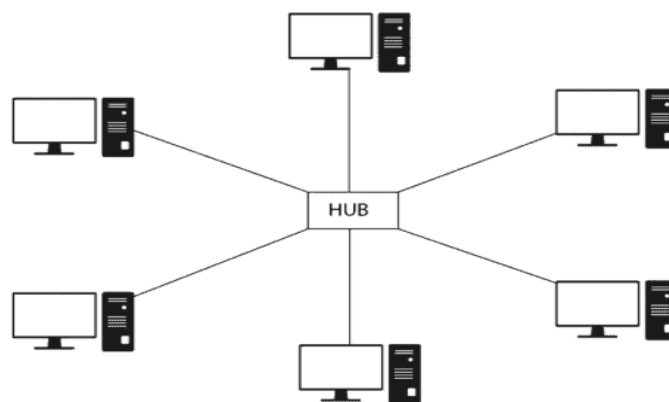


Figure Star Topology

Applications of Star Topology

In networking, superstar topology may be used a variety of places. You could find its uses all round yourself due to its easy availability and cheap famous person Topology device. But, some makes use of superstar topology are as follows:

- Most pc labs in educational institutions utilize this layout to attach nodes within the lab.
- Our domestic networks are definitely configured in this community Topology.

- Some other use of star topology is the banking zone, wherein all banking customers are related with every different with the help of this form of Topology.

Characteristics of Star Topology

The features or traits of famous person topology are as follows:

- The smooth-to-installation star topology may be hired in nearly any kind of computer community, whether or not small, medium, or large.
- Compared to bus topology, celebrity topology needs greater cable. Also, in this type of linked network, there is no dependency.
- In celebrity topology, to increase the entire network, you may use the daisy chain arrangement.
- As compared to different forms of topological structures, the shape of superstar topology is extra cozy in phrases of dropping records.
- It offers an advantage; the entire network does now not disturb in case you get rid of or join devices.
- Advantages of Star Topology
- A celeb topology is less luxurious than a mesh topology.
- In a celebrity, each device wishes simplest one link and one I/O port to attach it to any quantity of others. This thing also makes it easy to install and reconfigure.
- A ways less cabling desires to be housed, and additions, actions, and deletions involve handiest one connection: among that tool and the hub.
- It consists of robustness. If one link fails, most effective that hyperlink is affected. All other hyperlinks continue to be lively. This issue additionally lends itself to smooth fault identity and fault isolation. As long as the hub is working, it may be used to display link issues and pass defective hyperlinks.

Disadvantages of Star Topology

- A star topology is a dependency of the entire topology on one unmarried factor, the hub. If the hub goes down, the whole system is dead.
- Despite the fact that a celeb calls for a way less cable than a mesh, every node needs to be connected to a vital hub. Because of this,

often more cabling is required in a star than in a few other topologies (which include ring or bus).

3. 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.

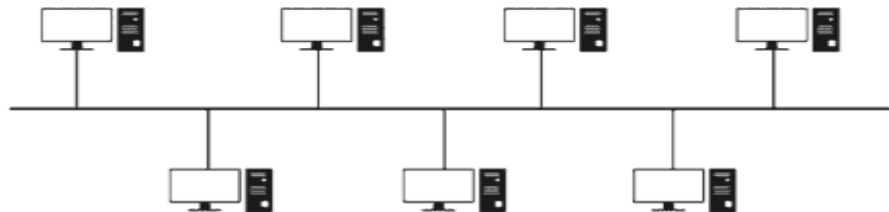


Figure Bus Topology

Nodes are connected to the bus cable by way of drop strains and taps. A drop line is a connection walking among the device and the principle cable. A tap is a connector that either splices into the primary cable or punctures the sheathing of a cable to create a contact with the metal core. As a sign travels alongside the backbone, a number of its energy is transformed into warmth. Consequently, it turns into weaker and weaker because it travels farther and farther. Because of this there is a restrict at the wide variety of faucets a bus can guide and on the gap among the ones taps. CSMA: it is a media get admission to manage used to govern the records glide so that information integrity is maintained, i.e., the packets do now not wander away. There are opportunity approaches of dealing with the troubles that arise while two nodes ship the messages concurrently.

- CSMA CD: CSMA CD (Collision detection) is an get right of entry to technique used to stumble on the collision. As soon as the collision is detected, the sender will stop transmitting the records. Consequently, it really works on "healing after the collision".
- CSMA CA: CSMA CA (Collision Avoidance) is an access technique used to keep away from the collision by using checking whether the transmission media is busy or not. If busy, then the sender waits till the media will become idle. This approach correctly reduces the possibility of the collision. It does now not work on "restoration after the collision".

Advantages of Bus Topology

- Benefits of bus topology consist of ease of installation. Spine cable can be laid along the most efficient direction, after which linked to the nodes through drop lines of numerous lengths.

- On this way, a bus uses much less cabling than mesh or megastar topologies. In a celeb, as an example, four community devices inside the same room require 4 lengths of cable achieving the entire manner to the hub.
 - In a bus, this redundancy is eliminated. Only the spine cable stretches via the whole facility. Each drop line has to reach only as some distance as the nearest factor on the spine.
- Disadvantages of Bus Topology

Disadvantages include difficult reconnection and fault isolation.

- Dangers encompass tough reconnection and fault isolation.
- A bus is usually designed to be optimally green at installation. It is able to consequently be difficult to feature new gadgets.
- Signal mirrored image at the faucets can purpose degradation in fine. This degradation may be controlled via proscribing the wide variety and spacing of gadgets linked to a given length of cable. Including new gadgets can also therefore require modification or substitute of the spine.
- In addition, a fault or damage inside the bus cable stops all transmission, even among devices at the equal facet of the hassle. The broken region displays alerts back in the course of foundation, creating noise in both guidelines.

4. Ring Topology

In a hoop topology, each device has a committed factor-to-point reference to only the 2 devices on either side of it. A sign is handed alongside the ring in one route, from device to device, until it reaches its destination. Each tool within the ring carries a repeater. While a tool receives a signal supposed for any other device, its repeater regenerates the bits and passes them along.

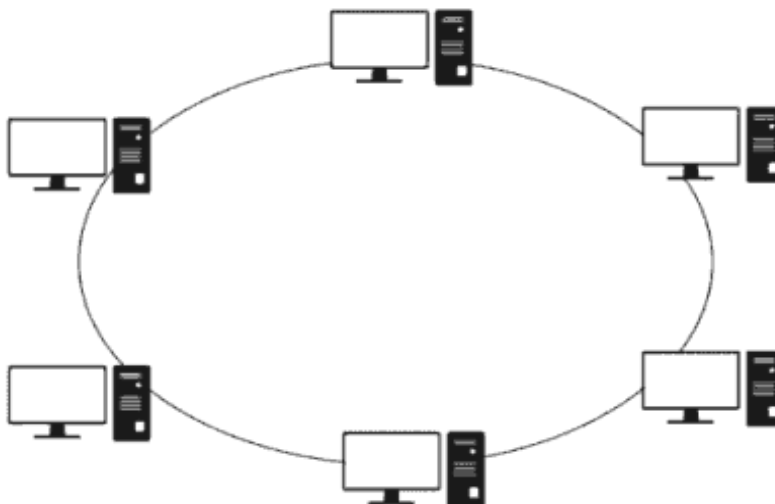


Figure Ring Topology

Advantages of Ring Topology

- A hoop is notably clean to put in and reconfigure. Every device is linked to best it's on the spot friends (both physically and logically). To feature or delete a tool requires converting simplest connections. The only constraints are media and visitors concerns (maximum ring period and number of devices).
- Similarly, fault isolation is simplified. Typically, in a ring, a signal is circulating always. If one device does no longer obtain a signal within a specified length, it cans difficulty an alarm. The alarm alerts the network operator to the hassle and its location.
- Disadvantages of Ring Topology
- In Uni-directional Ring, a know-how packet (token) ought to go through all of the nodes.
- In an easy ring, a break within the ring (which includes a disabled station) can disable the complete community. This weak point can be solved by means of the use of a twin ring or a switch able to closing off the ruin.
- Ring topology changed into typical whilst IBM introduced its local-place community Token Ring. These days, the need for higher-velocity LANs has made this topology less popular.

5. Hybrid Topology

It may be a combination of quite two topologies. Computer networking, a network structure that contains quite two topologies is understood as hybrid topology. It inherits the benefits and drawbacks of included topologies.

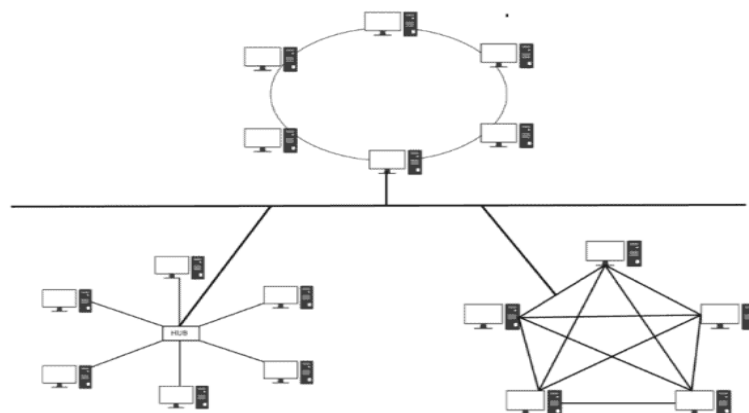


Figure Hybrid Topology

The hybrid topology includes a mixture of bus, mesh, ring topology, star, and tree topology. The mix of topologies depends on the necessity of a corporation.

Types of Hybrid Network Topologies

There are two types of hybrid network topologies counting on the essential requirement of a corporation. Still, the foremost commonly used one is Star-Ring and Star-Bus topologies that structure the hybrid.

Advantages of Hybrid Topology

The hybrid network combines the advantages of various sorts of topologies can be modified as per requirement

- It is very reliable
- It is extremely flexible.
- It is expensive
- It is easily scalable
- Design is complex.

Disadvantages of Hybrid Topology

- Hardware changes are required to attach topology to a different topology.

Differences between Star Topology and Bus Topology

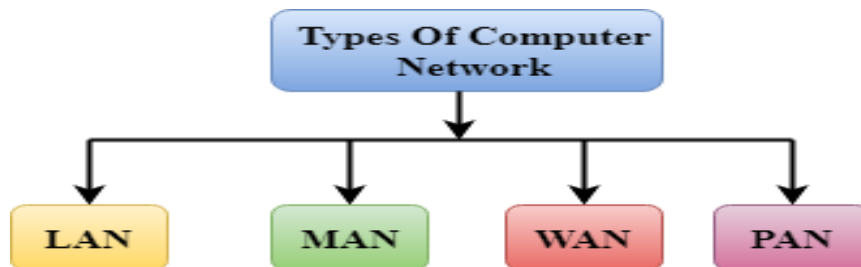
There are several distinctions between star and bus topologies. A table with all of the distinctions between star topology and bus topology is shown below:

STAR TOPOLOGY	BUS TOPOLOGY
A star topology is a network topology in which all devices are connected to a single central hub or switch.	A bus topology is a network topology in which all devices are connected by a single central connection.
In star topology, if the central core fails, the whole system will get affected, and even you cannot use the Computer Network.	In a Bus topology, if the network cable fails, the whole network will also fail.
The performance and the management of high traffic of the network are dependent on the central hub in a star topology.	When there is a lot of traffic on the network with a Bus topology, the network performance suffers. As a

	result, it is unable to adequately manage a large volume of traffic.
Any terminator is not included in the star topology.	At both ends of the network, the terminators are included in the bus topology.
Because of the need for extra wires and a central hub for connection, the cost of implementation star topology is high.	As compared to a star topology, bus topology is less costly.
In star topology, the rate of data transmission is fast.	The rate of data transmission in a bus topology is slower than in a star topology.
The nodes in a star architecture communicate through the central hub. The message is forwarded to the receiver node after it arrives at the central hub from the sender.	The process of data transmission in star topology is something different. In a bus topology, the sender's message is sent directly to the recipient.

1.6 NETWORK TYPES

A computer network is a group of computers linked to each other that enables the computer to communicate with another computer and share their resources, data, and applications. A computer network can be categorized by their size. A **computer network** is mainly of **four types**:



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

1.6.1 Local Area Network

A local location network (LAN) is a network that is restrained to fairly small vicinity. It's miles generally constrained to a geographic region together with a writing lab, college, or constructing. Computer systems connected to a community are widely labeled as servers or workstations. Servers are

typically no longer utilized by human beings directly, but alternatively run continuously to offer "offerings" to the opposite computer systems (and their human users) on the network. offerings supplied can include printing and faxing, software program web hosting, record storage and sharing, messaging, information storage and retrieval, entire get right of entry to control (protection) for the network's assets, and many others.



Figure: Local Area Network

Workstations are referred to as such due to the fact they usually do have a human consumer which interacts with the network via them. Workstations had been historically taken into consideration a computing device, such as a laptop, keyboard, show, and mouse, or a computer, with incorporated keyboard, show, and touchpad. With the advent of the pill pc, and the contact display devices along with iPad and iPhone, our definition of computing device is speedy evolving to include the ones devices, because of their capability to interact with the network and utilize community services.

Servers have a tendency to be more effective than workstations, even though configurations are guided by way of needs. For instance, a collection of servers is probably placed in a secure place, far from people, and only accessed thru the network. In such cases, it'd be commonplace for the servers to function without a dedicated show or keyboard. However, the dimensions and velocity of the server's processor(s), hard force, and main memory might upload dramatically to the cost of the machine. On the other hand, a laptop might not need as plenty storage or running reminiscence, however may require an high priced display to deal with the desires of its person. Every pc on a community should be correctly configured for its use.

On a single LAN, computers and servers may be related through cables or wirelessly. Wi-Fi get entry to a stressed out network is made viable with the aid of wireless access factors (WAPs). Those WAP devices provide a bridge between computer systems and networks. An average WAP would possibly have the theoretical capability to connect loads or maybe heaps of Wi-Fi users to a community, despite the fact that sensible potential is probably a way less.

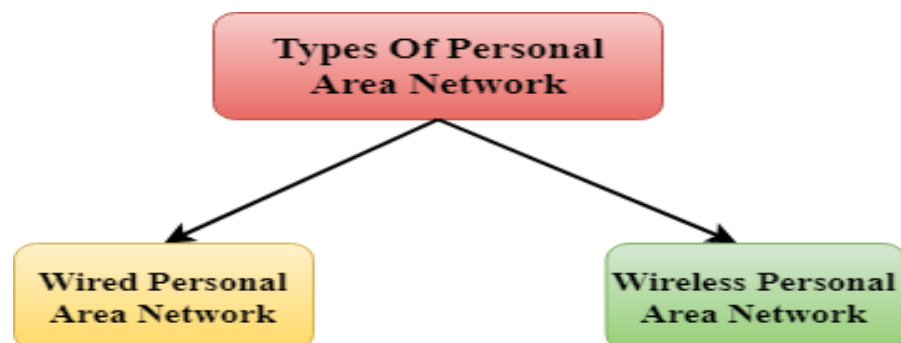
Almost constantly servers can be linked by using cables to the community, because the cable connections stay the quickest. Workstations which are stationary (computers) are also generally connected through a cable to the community, despite the fact that the fee of Wi-Fi adapters has dropped to the factor that, when installing workstations in an current facility with inadequate wiring, it is able to be simpler and much less highly-priced to use wireless for a computer.

1.6.2 PAN (Personal Area Network)

- Non-public region community is a community organized inside a person character, commonly within quite a number 10 meters.
- Private place community is used for connecting the computer gadgets of personal use is referred to as private location network.
- Thomas Zimmerman changed into the primary research scientist to deliver the concept of the private area community.
- Non-public region network covers an area of 30 ft.
- Personal computer devices which are used to increase the personal area community are the laptop, mobile phones, media player and play stations.



There are two types of Personal Area Network:



- Wired Personal Area Network
- Wireless Personal Area Network

Wireless Personal Area Network: Wi-Fi wireless personal region network is evolved by way of simply using Wi-Fi wireless technologies along with Bluetooth. It's far a low variety network

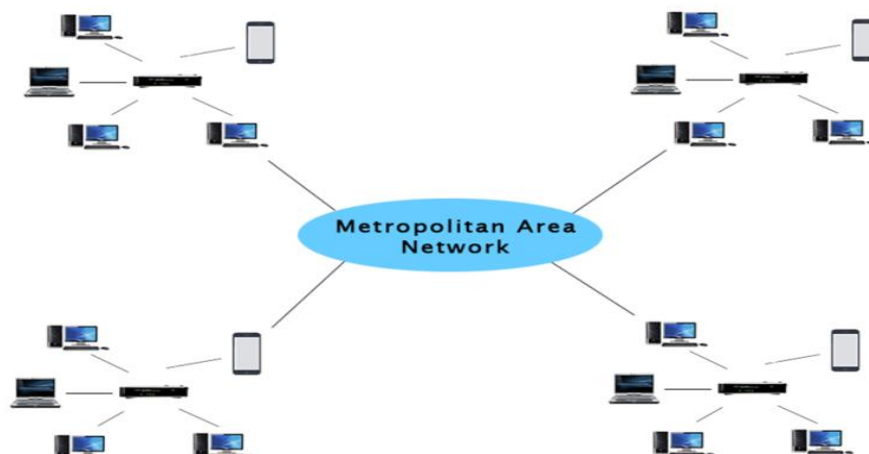
Wired Personal Area Network: Wired Personal Area Network is created by using the USB.

Examples of Personal Area Network

- Body vicinity community: body location network is a network that moves with a person. As an instance, a cellular network actions with a person. Think someone establishes a community connection after which creates a reference to every other device to proportion the statistics.
- Offline community: An offline community may be created inside the home, so it's also known as a home community. A home community is designed to combine the gadgets along with printers, pc, and television however they may be no longer connected to the net.
- Small domestic workplace: it's far used to attach a variety of devices to the net and to a company network the use of a VPN

1.6.3 MAN (Metropolitan Area Network)

- A metropolitan area community is a network that covers a bigger geographic place through interconnecting a distinctive LAN to form a larger network.
- Government corporations use man to hook up with the residents and private industries.
- In man, numerous LANs are connected to every different thru a smartphone change line.
- The most broadly used protocols in man are RS-232, frame Relay, ATM, ISDN, OC-three, ADSL, and so on.
- It has a higher variety than neighborhood region network (LAN).

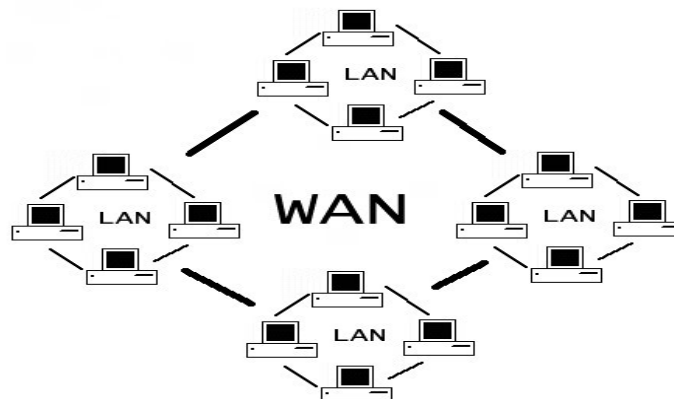


Uses of Metropolitan Area Network

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

1.6.4 WAN (Wide Area Network)

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



Examples of Wide Area Network:

- **Mobile Broadband:** A 4G network is widely used across a region or country.
- **Last mile:** A telecom company is used to provide the internet services to the customers in hundreds of cities by connecting their home with fiber.
- **Private network:** A bank provides a private network that connects the 44 offices. This network is made by using the telephone leased line provided by the telecom company.

Advantages of Wide Area Network:

Following are the advantages of the Wide Area Network:

- **Geographical location:** A extensive place community presents a big geographical vicinity. Suppose if the branch of our office is in a one-of-a-kind town then we can hook up with them via WAN. The net gives a leased line through which we will connect to every other branch.
- **Centralized statistics:** In case of WAN community, facts are centralized. Consequently, we do no longer want to shop for the emails, documents or again up servers.
- **Get up to date documents:** software program agencies work on the live server. Consequently, the programmers get the updated files inside seconds.
- **Trade messages:** In a WAN network, messages are transmitted fast. The web application like fb, Whatsapp, and Skype allows you to communicate with buddies.
- **Sharing of software program and resources:** In WAN community, we can share the software and different resources like a hard pressure, RAM.
- **International business:** we are able to do the commercial enterprise over the internet globally.
- **High bandwidth:** If we use the leased strains for our employer then this offers the excessive bandwidth. The excessive bandwidth will increase the statistics transfer rate which in flip increases the productivity of our company.

Disadvantages of Wide Area Network:

The following are the disadvantages of the Wide Area Network:

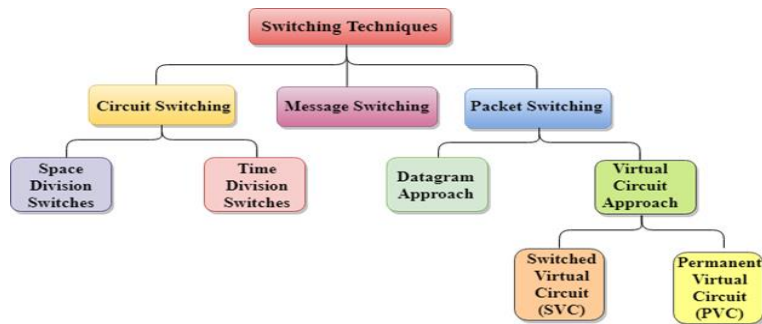
- **Security issue:** A WAN network has more security issues as compared to LAN and MAN network as all the technologies are combined together that creates the security problem.
- **Needs Firewall & antivirus software:** The data is transferred on the internet which can be changed or hacked by the hackers,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

1.7 SWITCHING

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

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

Classification of Switching Techniques



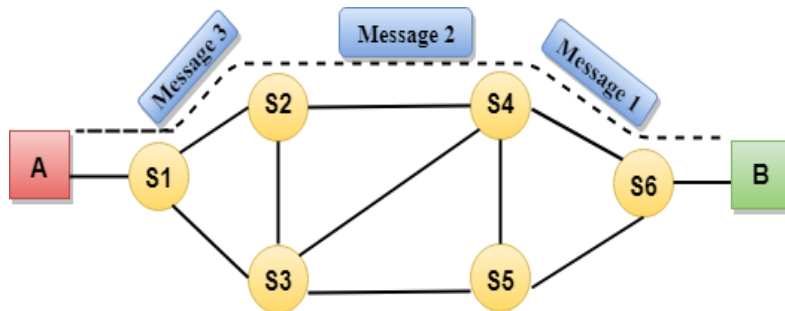
1.7.1 Circuit Switching

- Circuit switching is a switching technique that establishes a dedicated course among sender and receiver.
- Within the Circuit Switching technique, once the relationship is hooked up then the devoted course will remain to exist until the relationship is terminated.
- Circuit switching in a network operates in a comparable way as the telephone works.
- A whole quit-to-stop path ought to exist earlier than the conversation takes area.
- In case of circuit switching method, when any consumer desires to send the statistics, voice, video, a request sign is dispatched to the receiver then the receiver sends lower back the acknowledgment to make certain the supply of the dedicated course. After receiving the acknowledgment, devoted route transfers the facts.
- Circuit switching is used in public smartphone community. it's far used for voice transmission.
- Fixed information can be transferred at a time in circuit switching generation.

Communication through circuit switching has 3 phases:

- Circuit establishment
- Data transfer

- Circuit Disconnect



Circuit Switching can use either of the two technologies:

Space Division Switches:

- Area department Switching is a circuit switching technology wherein a single transmission course is completed in a switch by way of using a physically separate set of go factors.
- Area department Switching may be executed with the aid of the use of crossbar transfer. A crossbar switch is a metallic go point or semiconductor gate that may be enabled or disabled by way of a control unit.
- The Crossbar switch is made via the use of the semiconductor. For example, Xilinx crossbar switches using FPGAs.
- Space department Switching has excessive velocity, high ability, and no blocking off switches.

Space Division Switches can be categorized in two ways:

- Crossbar Switch
- Multistage Switch

Crossbar Switch

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

Disadvantage of Crossbar switch:

The quantity of pass factors will increase because the range of stations is accelerated. Therefore, it will become very high-priced for a huge switch. The answer to that is to use a multistage switch.

Multistage transfer

- Multistage switch is made via splitting the crossbar switch into the smaller devices and then interconnecting them.
- It reduces the number of go factors.

- If one course fails, then there might be an availability of every other direction.

Advantages of Circuit Switching:

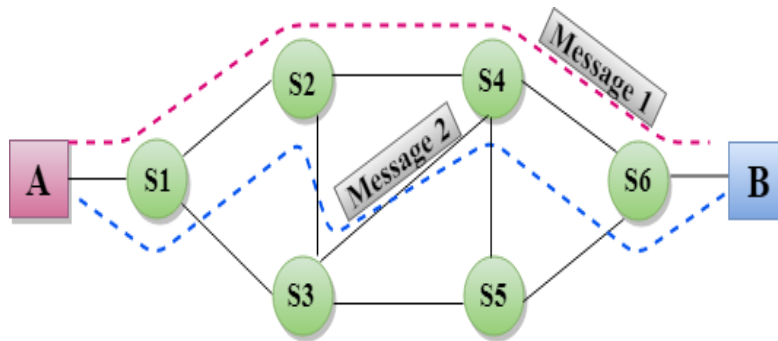
- In the case of Circuit Switching technique, the communication channel is dedicated.
- It has fixed bandwidth.

Disadvantages of Circuit Switching:

- Once the dedicated course is established, the only postpone occurs within the speed of data transmission.
- It takes a long term to establish a connection approx. 10 seconds at some stage in which no statistics can be transmitted.
- Its miles greater high priced than different switching strategies as a devoted course is required for each connection.
- It is inefficient to use due to the fact as soon as the direction is established and no information is transferred, then the capability of the direction is wasted.
- In this situation, the relationship is devoted therefore no other data may be transferred despite the fact that the channel is loose.

1.7.2 Message Switching

- Message Switching is a switching technique wherein a message is transferred as a entire unit and routed through intermediate nodes at which it's far saved and forwarded.
- In Message Switching method, there may be no status quo of a committed route between the sender and receiver.
- The vacation spot address is appended to the message. Message Switching presents a dynamic routing because the message is routed through the intermediate nodes based at the facts available inside the message.
- Message switches are programmed in this sort of manner so one can provide the maximum efficient routes.
- Every and each node shops the complete message after which forwards it to the next node. This sort of community is called shop and forward community.
- Message switching treats every message as an impartial entity.



Advantages of Message Switching

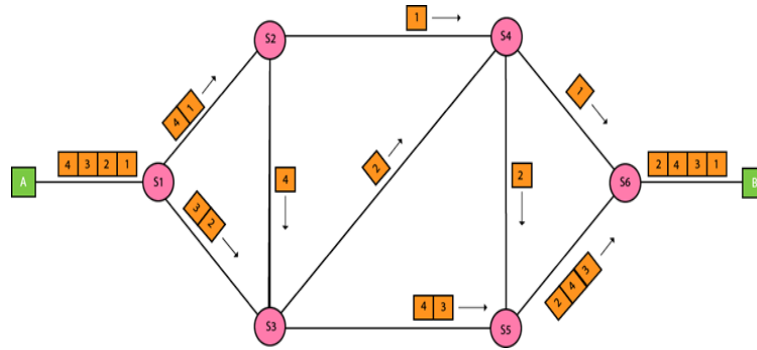
- Data channels are shared a number of the communicating gadgets that improve the performance of using available bandwidth.
- Site visitor's congestion can be reduced due to the fact the message is temporarily stored within the nodes.
- Message priority may be used to manage the community.
- The scale of the message that's sent over the community may be numerous. consequently, it supports the records of unlimited length.

Disadvantages of Message Switching

- The message switches should be geared up with sufficient storage to permit them to store the messages till the message is forwarded.
- The lengthy delay can occur due to the storing and forwarding facility supplied by the message switching technique.

1.7.3 Packet Switching

- The packet switching is a switching method in which the message is sent in a single move, however its miles divided into smaller portions, and they are sent personally.
- The message splits into smaller portions referred to as packets and packets are given a completely unique quantity to perceive their order on the receiving stop.
- Each packet includes some statistics in its headers together with source deal with, destination address and collection range.
- Packets will travel throughout the community, taking the shortest course as viable.
- All of the packets are reassembled at the receiving end in accurate order.
- If any packet is missing or corrupted, then the message might be sent to resend the message.
- If the precise order of the packets is reached, then the acknowledgment message can be dispatched.



Approaches of Packet Switching

There are two approaches to Packet Switching:

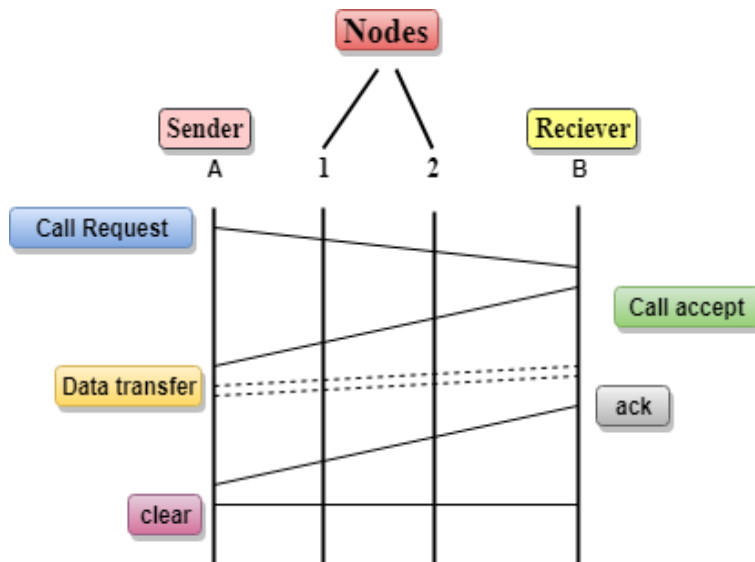
1.7.3.1 Datagram Packet switching:

- It's far a packet switching generation in which packet is referred to as a datagram, is taken into consideration as an impartial entity. Every packet includes the data approximately the destination and transfer makes use of these facts to ahead the packet to the proper vacation spot.
- The packets are reassembled on the receiving end in correct order.
- In Datagram Packet Switching approach, the direction isn't fixed.
- Intermediate nodes take the routing choices to forward the packets.
- Datagram Packet Switching is also known as connectionless switching.

1.7.3.2 Virtual Circuit Switching

- Virtual Circuit Switching is also known as connection-oriented switching.
- In the case of Virtual circuit switching, a preplanned route is established before the messages are sent.
- Call request and call accept packets are used to establish the connection between sender and receiver.
- In this case, the path is fixed for the duration of a logical connection.

Let's understand the concept of virtual circuit switching through a diagram:



- In the above diagram, A and B are the sender and receiver respectively. 1 and a couple of are the nodes.
- Call request and make contact with receive packets are used to set up a connection among the sender and receiver.
- Whilst a course is set up, records may be transferred.
- After transmission of facts, an acknowledgment signal is sent with the aid of the receiver that the message has been received.
- If the consumer desires to terminate the relationship, a clean sign is dispatched for the termination.

Differences b/w Datagram approach and Virtual Circuit approach

Datagram approach	Virtual Circuit approach
Node takes routing decisions to forward the packets.	Node does not take any routing decision.
Congestion cannot occur as all the packets travel in different directions.	Congestion can occur when the node is busy, and it does not allow other packets to pass through.
It is more flexible as all the packets are treated as an independent entity.	It is not very flexible.

Advantages of Packet Switching:

- cost-effective: In packet switching method, switching gadgets do now not require large secondary garage to shop the packets, so value is minimized to a point. Consequently, we can say that the packet switching approach is a cost-effective approach.

- **Reliable:** If any node is busy, then the packets can be rerouted. This ensures that the Packet Switching method presents dependable communication.
- **Efficient:** Packet Switching is an effective method. It does no longer require any hooked up route prior to the transmission, and many users can use the same communicate channel concurrently, consequently makes use of to be had bandwidth very efficaciously.

Disadvantages of Packet Switching:

- Packet Switching technique cannot be implemented in those applications that require low delay and high-quality services.
- The protocols used in a packet switching technique are very complex and requires high implementation cost.
- If the network is overloaded or corrupted, then it requires retransmission of lost packets. It cannot also lead to the loss of critical information if errors are nor recovered.

1.8 THE INTERNET

Net is a worldwide community that connects billions of computer systems the world over with every other and to the arena wide net. It makes use of standard net protocol suite (TCP/IP) to connect billions of laptop users worldwide. it is set up by using cables together with optical fibers and other Wi-Fi and networking technologies. At gift, net is the quickest mean of sending or changing facts and facts among computer systems the world over. it's far believed that the internet changed into developed by using "defense superior tasks business enterprise" (DARPA) branch of the USA. And, it was first related in 1969.

Why is the Internet Called a Network?

Internet is known as a community as it creates a network through connecting computer systems and servers the world over the use of routers, switches and phone traces, and different conversation gadgets and channels. So, it can be considered a global network of physical cables together with copper cellphone wires, wi-fiber optic cables, TV cables, etc. moreover, even Wi-Fi connections like 3G, 4G, or c084d04ddacadd4b971ae3d98fecfb2a make use of these cables to get admission to the internet.

Internet isn't the same as the world wide net as the sector extensive net is a network of computer systems and servers created by using connecting them through the internet. So, the net is the backbone of the internet as it gives the technical infrastructure to establish the WWW and acts as a medium to transmit statistics from one laptop to another laptop. It uses net browsers to show the records on the patron, which it fetches from net servers.

The internet is not owned with the aid of a single character or organization entirely. it is a idea primarily based on physical infrastructure that connects networks with other networks to create a global network of billions of

computers. As of 12 August 2016, there had been greater than 300 crores of net users across the world.

How does internet work?

Before understanding this let us understand some basics related to internet:

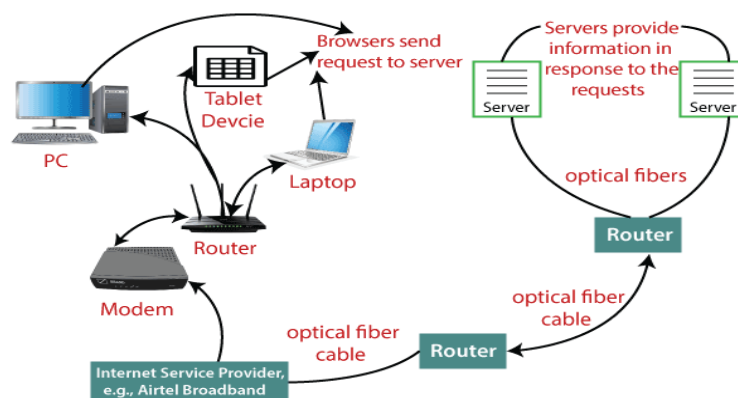
The internet works with the assist of clients and servers. A tool together with a laptop, which is related to the internet, is called a patron, not a server because it isn't without delay linked to the internet. However, their mile in a roundabout way linked to the net through an internet carrier provider (ISP) and is identified by using an IP deal with, that's a string of numbers. Similar to you have a cope with for your home that uniquely identifies your house, an IP cope with acts because the transport address of your tool. The IP address is furnished via your ISP, and you could see what IP cope with your ISP has given for your machine.

A server is a huge laptop that stores web sites. It additionally has an IP cope with. a place where a huge range of servers are saved is referred to as a facts middle. The server accepts requests send by way of the consumer through a browser over a community (net) and responds therefore.

To get entry to the net we need a site name, which represents an IP address quantity, i.e., every IP cope with has been assigned a domain name. For example, youtube.com, fb.com, paypal.com are used to represent the IP addresses. Domains are created as it's far hard for someone to take into account an extended string of numbers. however, net does now not apprehend the area call; it is aware the IP address, so while you input the area name inside the browser seek bar, the net has to get the IP addresses of this area name from a big phone e book, that's called DNS (area call Server).

As an instance, if you have a person's name, you could find his smartphone wide variety in a telephone book through searching his name. The net makes use of the DNS server inside the same manner to locate the IP deal with of the domain name. DNS servers are managed by ISPs or similar companies.

Now after understanding the basics, let us see how internet works?



While you switch to your pc and sort a domain name in the browser seek bar, your browser sends a request to the DNS server to get the corresponding IP cope with. After you have the IP address, the browser forwards the request to the respective server.

Once the server gets the request to provide records about a particular wireless internet site, the statistics starts off evolved flowing. The records are transferred via the optical wirelesses cables in virtual format or inside the shape of light pulses. Maybe as the servers are positioned at distant places, the statistics additionally ought to tour heaps of miles via optical wi-fiber cable to reach your computer.

The optical wi-fiber is attached to a router, which converts the mild alerts into electric indicators. These electrical signals are transmitted on your pc the usage of an Ethernet cable. For that reason, you receive the desired information through the net, that is simply a cable that connects you with the server.

Moreover, in case you are using Wi-Fi net the usage of cellular statistics, the indicators from the optical cable are wirelesses sent to a mobile tower and from wherein it reaches in your cellular phone in the shape of electromagnetic waves.

The internet is managed by means of ICANN (net enterprise for Assigned Names and Numbers) located inside the United States of America. It manages IP addresses task, area name registration, and many others.

The statistics switch could be very fast at the net. The moment you press enter you get the statistics from a server positioned heaps of miles away from you. The purpose for this speed is that the records is sent within the binary shape (0, 1), and those zeros and ones are divided into small portions known as packets, which can be dispatched at excessive place.

Advantages of the Internet:

- **Immediately Messaging:** you could send messages or talk to all people the usage of net, together with email, voice chat, video conferencing, and so forth.
- **Get instructions:** the usage of GPS technology, you can get directions to almost each region in a metropolis, United States, and so forth. You may discover eating places, shops, or any other carrier close to your place.
- **Online buying:** It allows you to save online which include you may be clothes, shoes, e-book movie tickets, railway tickets, flight tickets, and more.
- **Pay payments:** you may pay your payments on line, such as strength bills, fuel payments, college charges, and so on.

- On-line Banking: It lets in you to use net banking in which you may take a look at your balance, acquire or switch money, get an announcement, request cheque-e-book, and many others.
- Online promoting: you may promote your services or products online. It facilitates you attain more clients and consequently will increase your sales and profit.
- Earn a living from home: in case you want to work at home, you can do it the usage of a device with net access. Today, many companies permit their employees to earn a living from home.
- Enjoyment: you can pay attention to online song, watch motion pictures or films, play online video games.
- Cloud computing: It allows you to connect your computer systems and internet-enabled devices to cloud offerings together with cloud garage, cloud computing, and many others.
- Career building: you can look for jobs online on different task portals and send you CV thru e mail if required.

1.8.1 Internet Accessing

Net access is frequently furnished at domestic, schools, places of work, public locations, internet cafes, libraries and different locations. The net started to advantage recognition with dial-up net get right of entry to. In a pretty short time, internet access technology modified, presenting faster and more dependable alternatives. Presently, broadband technologies such as cable internet and ADSL are the most extensively used strategies for net get right of entry to. The speed, value, reliability and availability of net get entry to depend at the region, internet Provider Company and form of connection. There are many different ways to obtain internet access, including:

- Wireless connection
- Mobile connection
- Hotspots
- Dial-up
- Broadband
- DSL
- Satellite

Access to computer systems or clever gadgets is one of the essential factors for information the level of net get admission to for a region. However, net get right of entry to is not uniformly distributed within or among countries. A virtual divide exists among many nations and areas. Accurate net get admission to be related to regions with excessive-profits populations, a high improvement index and high technological development.

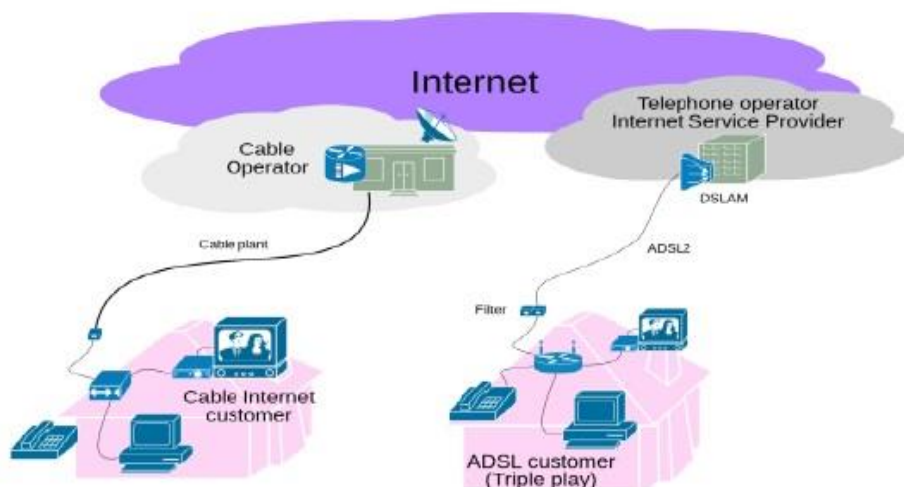
1.8.1.1 Dial-up Connections

In dial-up connection, pc makes use of its modem to dial a telephone number given to the consumer by using a web provider. This launches a connection among personal pc and ISP server. The process starts off evolved when the ISP server answers, and ceases whilst your laptop or the server "hangs up". This is similar to a conventional phone name. Most ISP servers disconnect mechanically after certain duration of inaction. As soon as a connection is configured at the user's pc, he/she can use the relationship. It far comfy and de-allocates unused memory robotically.



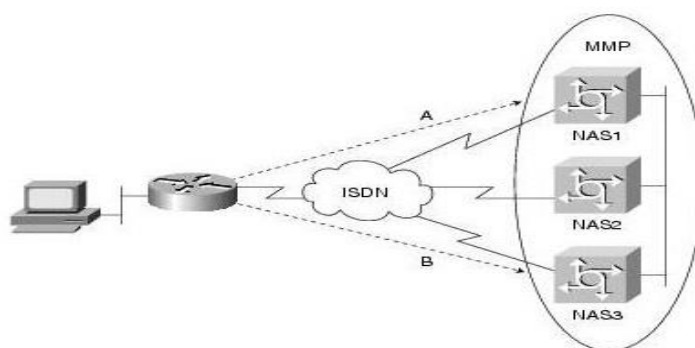
1.8.1.2 Broadband Connection

Broadband connections are considered as high pace connections, as they use modes that may cope with numerous signals right now, which include fiber optics, twisted pair cables, coaxial cable and different technology. Even with masses of users at the network, those connections permit huge files and complicated web pages to down load fast. To be taken into consideration as a broadband, the relationship should be able to transmit facts at a charge faster than is possible with the fastest dial-up connection. Downloading and importing content will be rapid.



1.8.1.3 Integrated Services Digital Network (ISDN) Service

Incorporated offerings virtual network (ISDN) is a digital service that simultaneously transmits voice & facts, and controls alerts over a single smartphone line. ISDN provider operates on a general cellphone line, however requires a special modem and speaks to provider, which adds to the value. An ISDN statistics connection can switch data up to 128,000 bits in line with 2d (128 Kbps). It facilitates to connect a computer, smartphone and fax to a single ISDN line and use them simultaneously.



1.8.1.4 Digital Subscriber Line (DSL)

Virtual Subscriber Line is just like that of ISDN in using cellphone community, but it makes use of more advanced virtual signal processing and algorithms to squeeze maximum range of indicators through cellphone traces. DSL also requires changes in components of cellphone community before it is able to be presented in any area. Like ISDN, DSL presents simultaneous information, voice and fax transmission on the equal line. Numerous variations of DSL offerings are available for domestic and enterprise use; every model presents 24/7 complete-time connection at different tiers of service, velocity, bandwidth and distance.

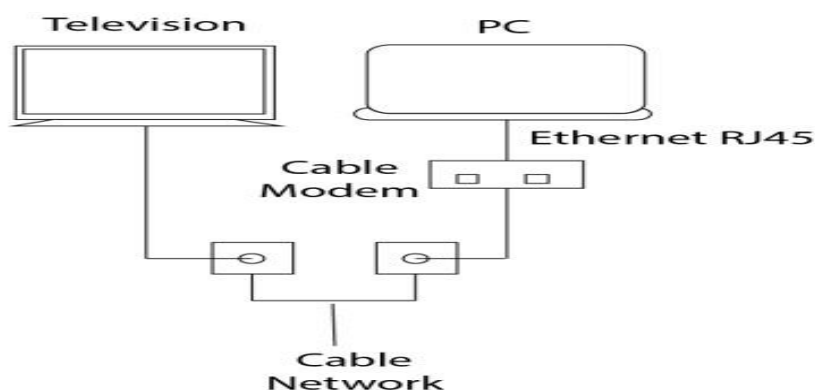


DSL Type	Maximum Sending speed	Maximum Receiving speed	Maximum Distance	Lines Required	Phone Support
ADSL	800 Kbps	8 Mbps	5,500 m	1	Yes

HDSL	1.54 Kbps	1.54 Mbps	3,650 m	2	No
IDSL	144 Kbps	144 Mbps	10,700 m	1	No
MSDSL	2 Mbps	2 Mbps	8,800 m	1	No
RADSL	1 Mbps	7 Mbps	5,500 m	1	Yes
SDSL	2.3 Mbps	2.3 Mbps	6,700 m	1	No
VDSL	16 Mbps	52 Mbps	1,200 m	1	Yes

1.8.1.5 Cable Modem Service

Now-a-days many cable TV agencies use a few percentages in their network's bandwidth to offer internet get entry to through prevailing cable television connections. Seeing that this connection uses a special cable modem, it's far called "Cable Modem provider". Cable television systems transmit facts through coaxial cable that can transmit statistics as a lot as a hundred instances faster than not unusual telephone traces. Coaxial cable lets in transmission through numerous channels simultaneously, i.e., the internet statistics can be transmitted on one channel, whilst audio, video and control indicators are transmitted one after the other. The user can get right of entry to internet and watch television concurrently, with non-interfering statistics streams.



1.8.1.6 Wireless LAN (WLAN) Connections

Wireless LAN connections are very common these days, which are based on the technology that is often cited as Wi-Fi (Wireless Fidelity). The distance covered by WLAN is usually measured in meters rather than miles. Therefore, this is not a technology that connects directly to an ISP but can be used to connect to another LAN or device through which internet access is achieved.



Process

- To connect to internet, the wireless access point is connected to a wired LAN like any other devices, and then computers with wireless NICs can access the wired LAN.
- "Wireless access point" is a device that acts as a hub or switch.
- "NIC" refers to a Network Interface Card which helps to identify a computer on a network.

1.8.1.7 Wireless WAN (WWAN) Connections

A WWAN is a digital network that spans over a large geographical place. A WWAN accepts and transmits data the usage of radio signals through cell web sites and satellites. On the switching middle, the WWAN divides off into segments after which connect to either remoted or public community thru cellphone or different excessive pace verbal exchange links. The facts are then linked to an employer's current LAN/WAN infrastructure. The coverage vicinity for WWAN is generally measured in miles (kilometers) with a records transmission rate of one hundred Mbps.



1.8.1.8 Satellite Services

Satellite services provide a mutual (two-way) communication between user and the internet. This provides a full-time connection which is used in armed forces, business, etc. It includes two parts

Transceiver- A satellite dish that is placed outdoors in direct line of sight to one of the several satellites in geostationary orbit.

Modem-like device - It is connected to a dish, placed indoors and connected to a LAN or computer.

1.9 INTERNET STANDARD

Net preferred is a working report or a proposed specification that profits the reputation of the standard after proper validation and is applicable to the global net. This net widespread is fashioned beneath the supervision of the net Society (ISOC) which is answerable for the introduction and implementation of net standards. In this phase, we are able to talk the ISOC business enterprise, and we are able to additionally see the step-via-step formation of internet requirements.

Formation of net Society

Earlier there have been no such corporations or committees that could outline the net requirements. But, human beings round the arena had been doing their own factor to set up global communication. This raises the want for requirements that may be time-honored across the world.

even as running on ARPANET the branch of defense (DoD) set up an informal committee i.e., the net activities Board (IAB) which fits to installation a network of the related computer so they can proportion their findings. This will additionally lessen or even gets rid of the cost of duplication effort.

The acronym of IAB becomes modified to internet structure Board (IAB). The IAB consisted of approximately ten individuals in which each member become assigned a project, which they used to document to DOD. DOD and NSF used to raise the fund for IAB.

On every occasion a brand new popular changed into required the contributors of IAB use it to discuss the requirements and announce this to the software graduate college students to put into effect it with the help of coding. The communique between IAB and students changed into completed with the help of series of on-line technical reviews i.e., RFCs (Request For comments). Those feedbacks are prepared in serial order online and can be fetched by means of everybody who's inquisitive about them.

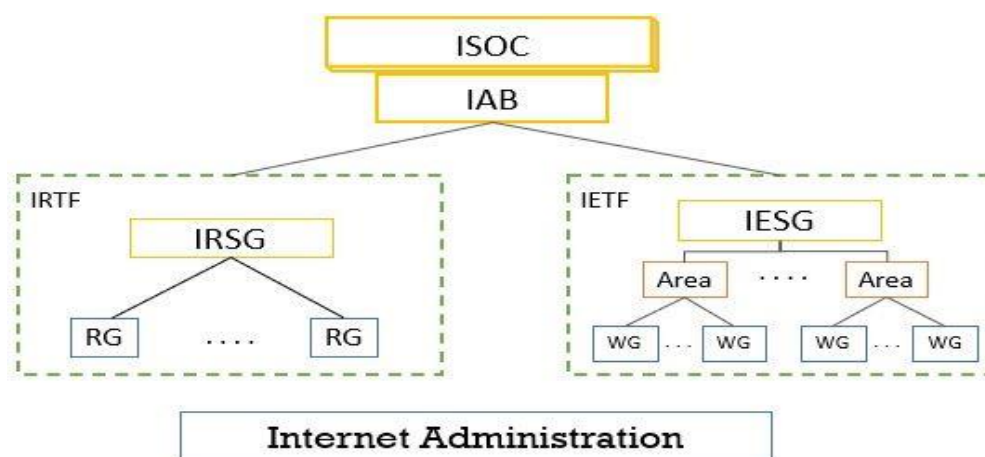
With the non-stop boom of the internet, IAB was reorganized in the 12 months 1989. It changed into repopulated to form an corporation which was then split into two organizations, IRTF (net research assignment force) and IETF (net Engineering task pressure). IRTF is held for undertaking lengthy-term studies and IETF is held to address quick-time period problems.

A year later an internet society was for and populated with who were inquisitive about the net. The net Society emerges as an agency that became capable of creating standards for the net. The trustees of the net Society then have been responsible to employ the individuals of IAB.

That is how the net Society (ISOC) became fashioned which thoroughly check and authorized the proposed net standards.

1.9.1 Internet Administration

Just now we have seen the formation of ISOC (Internet Society). Well under ISOC there are various groups that coordinate to create and maintain the internet standards. In the figure below you can see the organization of all the groups under the ISOC.



Internet Society (ISOC)

The ISOC is an international organization that thoroughly tests the proposed specifications before it is realized into the standard. Several administrative bodies such as IEFT, IA, etc. work cooperatively under the ISOC which is responsible to address all the issues that decide the future of the Internet. ISOC encourages scholarly activities and research in the field of the Internet.

IAB (Internet Architecture Board)

IAB issues technical advice to the Internet Society in order to ensure continuous growth of the Internet. IAB oversees the evolution of the Internet so that it becomes a global platform for communication.

IAB supervises two task forces namely IRTF (Internet Research Task Force) and IETF (Internet Engineering Task Force). IAB edits and manages the RFCs. IAB makes a connection between the Internet and other organizations that operate in the same direction.

Internet Research Task Force (IRTF)

IRTF conducts long-term researches related to the Internet. The research topic is related to Internet protocols, Internet applications, about their architecture. All the research groups involved in the research must have long-term membership. The research group (RG) or the working group under the IRTF is managed by the IRSG (Internet Research Steering Group).

Internet Engineering Task Force (IEFT)

IETF tackles short-term engineering issues subjected to the internet. The running organizations beneath IEFT are controlled by using the IESG (internet Engineering steerage group). The IEFT identifies the troubles at the net and proposes a option to resolve the trouble.

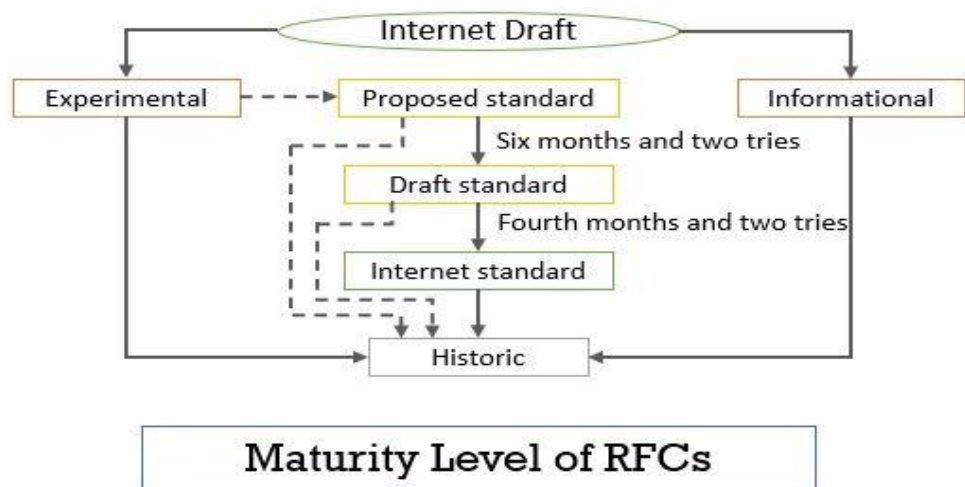
IETF additionally evaluations the proposed specs that are able to turning into the internet standards. each working institution is assigned a selected topic. The context of the assigned topic can be divided into nine classes

together with application, net protocols, user offerings, net operation, routing, community control, shipping, internet protocol subsequent technology, and safety.

1.9.1.1 Internet Standardization Process

Internet standards are those proposed specifications that are thoroughly tested for their successful implementation before they gain the status of Internet standard. Let us see the step-by-step procedure of the formation of Internet standards.

In the figure below we can observe level by level maturity of Internet standard.



A specification or announcement or a working report is taken into consideration as a web draft that has no professional fame and has an entire life of best six months. Upon approval through a few net governments, the net draft is posted as Request for remark (RFC). The RFCs are assigned a particular wide variety and are to be had on-line for humans inquisitive about the net. Every RFC before gaining the status of well-known falls at least into one of the maturity ranges underneath.

Maturity Level

1. Experimental

The published RFC is classified as experimental if it is applicable only in an experimental situation.

2. Informational

The published RFC is classified as informational if it is consisting of any kind of information regarding the Internet. Usually, they put forth by non-Internet organizations such as vendors.

3. Proposed Standard

The RFC specification is classified as a proposed standard if it is well defined and has the work which seems to be of the Internet community's interest. The specification here is examined and implemented by the different groups.

4. Draft Standard

Two successful implementations independent of each other and still interoperable promotes the proposed standard promote it to draft standard.

5. Internet Standard

Validation of successful implementation of a draft standard elevates it to Internet standard.

6. Historic

The Historic RFC are the specifications that are replaced by the newer specifications or those specifications that never reached the maturity levels that are required to gain the status of Internet standard.

The RFC itself can be of five different kinds on the basis of requirement.

Requirement Level

1. Required

The RFC is considered as required if it is mandatory to implement on all Internet systems to check whether it meets the specified standard. For example, ICMP an IP

2. Recommended

The RFC is considered as recommended if it is useful but is not mandatory to implement it on all Internet systems to check whether it meets the specified standard. For example, FTP and TELNET

3. Elective

The RFC is considered an elective if it can be used by an Internet system only for its own benefit.

4. Limited Use

The RFC is considered limited if can be used only in a limited situation. The RFC that are classified as experimental in maturity level is limited RFCs.

5. Not Recommended

The RFC is not meant for general use is not recommended RFC. Historical RFCs fall in this category.

So this is all about the Internet standards how they are formed, which committee is responsible for the creation, implementation, and maintenance of the Internet standards.

1.10 LIST OF REFERENCES

<https://www.globalspec.com/reference/23852/203279/internet-administration-governance-and-standards>

<https://www.geeksforgeeks.org/internet-administration/>

<https://www.ietf.org/standards/>

<https://binaryterms.com/internet-standards.html>

<https://www.javatpoint.com/internet>

<https://www.javatpoint.com/computer-network-switching-techniques>

<https://www.javatpoint.com/types-of-computer-network>

<https://www.geeksforgeeks.org/types-of-computer-networks/>

https://www.tutorialspoint.com/data_communication_computer_network/index.htm

<https://www.kdkce.edu.in/pdf/Data-comm-UNIT-I.pdf>

https://www.tutorialspoint.com/computer_concepts/computer_concepts_representation_data_information.htm

1.11 EXERCISES

- 1 What is meant by Data Communication and explain its characteristics?
- 2 What are the components of Data communication?
- 3 Explain different Data flow directions.
- 4 What is Network and explain characteristics of Networks?
- 5 Write about different types of connections.
- 6 Write about Protocol and Standards.
- 7 Explain different types of topologies.
- 8 Explain different types of Networks.
- 9 Write note on following
 - 1 Packet switching
 - 2 Message Switching
 - 3 Circuit Switching
- 10 Define Switching, Internet, Network and Topology

INTRODUCTION OF NETWORK MODELS

Unit Structure:

- 2.0 Objectives
- 2.1 Introduction of Network Model
- 2.2 Protocol Layering
 - 2.2.1 Scenarios
 - 2.2.2 Principles of Protocol Layering
- 2.3 Logical Connection
- 2.4 TCP/IP Protocol Suite
- 2.5 Layered Architecture
- 2.6 Layers in the TCP/IP Protocol Suite
- 2.7 Encapsulation and Decapsulation
- 2.8 Addressing
 - 2.8.1 Classful Addressing
- 2.9 Multiplexing and Demultiplexing
- 2.10 Introduction to Physical Layer
- 2.11 Introduction to Data-Link Layer
- 2.12 Introduction to Network Layer
- 2.13 Introduction to Transport Layer
- 2.14 Introduction to Application Layer
- 2.15 List of References
- 2.16 Exercises

2.0 OBJECTIVES

After going through this unit, you will be able to:

- Define Network Model, and Protocols
- State the TCP/IP Model
- Explain what Encapsulation, Decapsulation Concept
- Illustrate the DLL, Network layer, Transport layer and Application Layer

2.1 INTRODUCTION OF NETWORK MODEL

Networking engineering is a complex challenge, which entails software, firmware, chip degree engineering, hardware, and electric pulses. To ease network engineering, the entire networking idea is split into multiple layers. Every layer is worried in some precise undertaking and is impartial of all different layers. However, as a whole, nearly all networking obligations rely on all of those layers. Layers proportion facts among them and they rely on each other only to take input and ship output. A communication subsystem is a complex piece of hardware and software. Early tries for imposing the software for such subsystems have been based totally on a single, complex, unstructured program with many interacting additives. The consequent software program became very difficult to check and alter. To overcome such hassle, the ISO has advanced a layered technique. In a layered method, networking idea is divided into numerous layers, and every layer is assigned a specific project. Consequently, we can say that networking responsibilities rely upon the layers.

Why Use a model?

Before we go to a long way, let's perform a little truth take a look at. A model describes the entire structure. At the start of the bankruptcy, I stated that many networking fashions "have long gone the manner of the dodo." There may also have been accurate thoughts in each, but all of us ended up using one model mainly—TCP/IP. For instance, both Apple and IBM initially advanced their very own protocol suites, however transformed to TCP/IP because of its popularity. This segment explains the historical use of models and affords a more present day viewpoint. Even an easy conversation machine is a complex environment wherein heaps or even millions of transactions arise day by day. Interconnected systems are appreciably greater complicated. A single electric disturbance or software program configuration error can prevent crowning glory of these transactions. Fashions offer a place to begin for determining what should be executed to allow verbal exchange or to discern out how systems the usage of specific protocols may hook up with each other. In addition, they help in troubleshooting troubles. as an example, how might a Novell NetWare client strolling IPX/SPX communicate with an IBM AS/four hundred over a TCP/IP-primarily based community? Depicts a scenario in which numerous one-of-a-kind platforms might be required to have interaction with each different. Windows nodes are primarily based at the TCP/IP protocol suite but, if required, can run Novell NetWare patron software program for network authentication. Novell evolved internetworking and shipping protocols called IPX and SPX. At the alternative end of the network, the IBM mainframe communicates via the protocols used in the SNA version. Imagine the programming and further attempt required to keep all the transactions between those separate architectures.

2.2 PROTOCOL LAYERING

A protocol is a fixed of regulations and standards that often outline a language that devices will use to talk. There are a fantastic range of protocols in use significantly in networking, and that they're commonly carried out in several layers.

It affords a communication provider wherein the process is used to change the messages. When the communication is easy, we are able to use handiest one simple protocol.

Whilst the conversation is complicated, we have to divide the assignment among specific layers, so, we want to observe a protocol at each layer, this approach we used to call protocol layering. This layering lets in us to separate the services from the implementation.

Each layer wishes to receive a set of offerings from the lower layer and to present the offerings to the top layer. The change completed in any person layer will not have an effect on the alternative layers.

Basic Elements of Layered Architecture

The basic elements of the layered architecture are as follows –

- Service – Set of actions or services provided from one layer to the higher layer.
- Protocol – It defines a set of rules where a layer uses to exchange the information with its peer entity. It is concerned about both the contents and order of the messages used.
- Interface – It is a way through that the message is transferred from one layer to another layer.

Reasons

The reasons for using layered protocols are explained below –

- Layering of protocols provides well-defined interfaces between the layers, so that a change in one layer does not affect an adjacent layer.
- The protocols of a network are extremely complicated and designing them in layers makes their implementation more feasible.
- Advantages

The blessings of layered protocols are as follows –

- Assists in protocol fashion, due to protocols that characteristic at a particular layer have referred to records those they paintings and a described interface to the layers on top of and underneath.
- Foster's competition due to the fact products from sincerely fantastic companies will work along.

- Prevents technology or functionality adjustments in a single layer from touching first-rate layers above and beneath.
- Gives an average language to provide an explanation for networking functions and abilities.

Disadvantages

The disadvantages of layered protocols are as follows –

- The main hazards of layered systems consist by and large of overhead each in computation and in message headers resulting from the abstraction boundaries between layers. Because a message normally should bypass through several (10 or greater) protocol layers the overhead of those limitations is normally extra than the computation being finished.
- The upper-stage layers cannot see what's inside the lower layers, implying that a utility cannot correct wherein in an incredibly connection a problem is or exactly what the matter is.
- The better-level layers can't manage all elements of the decrease layers, in order that they can't adjust the switch gadget if beneficial (like controlling windowing, header compression, CRC/parity checking, et cetera), nor specify routing, and should depend upon the lower protocols running, and can not specify alternatives when there are issues.

Scenario

Dividing the task between different layers is called Protocol layering. Scenarios Two simple scenarios are available to understand the need for protocol layering.

First Scenario

In the first scenario, communication is simple that it can occur in only one layer. Assume Maria and Ann is neighbors with a lot of common ideas. Communication between Maria and Ann takes place in one layer, face to face, in the same language, as shown in Figure



Figure: Single layer protocol

Set of rules followed in this scenario:

First, Maria and Ann realize that they have to greet each different after they meet. 2d, they recognize that they ought to confine their vocabulary to the level in their friendship. 1/3, each birthday party is aware of that she should refrain (not talking) from speak me while the alternative party is talking. Fourth, each celebration knows that the conversation needs to be a conversation. Fifth, they ought to trade a few best words after they leave.

2nd state is that

Inside the second situation, we assume that Ann is obtainable a higher-degree role in her agency, however wishes to move to some other department positioned in a town very some distance from Maria. The two pals still want to retain their communicate and exchange ideas because they have given you an innovative challenge to start a brand new business when they both retire. They decide to retain their communicate the usage of regular mail through the post workplace. They agree on an encryption/decryption approach. The sender of the letter encrypts it to make it unreadable by an interloper; the receiver of the letter decrypts it to get the original letter.

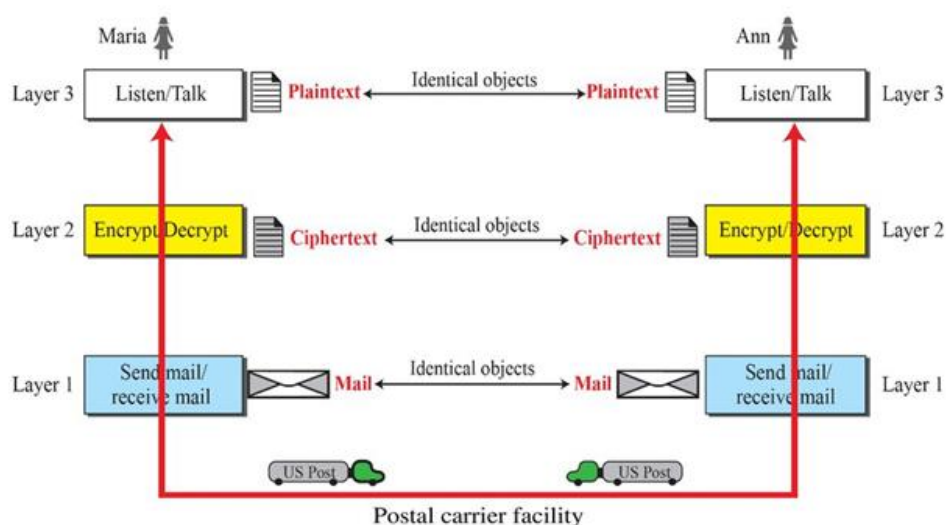


Figure: Three layer protocol

Don't forget that Maria sends the primary letter to Ann. Maria talks to the device at the third layer as although the device is Ann and is taking note of her. The third layer machine listens to what Maria says and creates the plaintext (a letter in English), that's passed to the second one layer system. The second layer device takes the plaintext, encrypts it, and creates the cipher textual content, which is passed to the primary layer device. The first layer gadget, takes the cipher textual content, puts it in an envelope, provides the sender and receiver addresses, and mails it. Protocol layering enables us to divide a complicated venture into numerous smaller and less complicated tasks. As an instance, inside the determine 1.3.2, we should have used simplest one device to do the task of all three machines. However, if Maria and Ann determine that the encryption/decryption achieved with

the aid of the device isn't sufficient to guard their secrecy; they could must alternate the complete gadget. Within the present state of affairs, they need to alternate only the second one layer device; the other two can continue to be the equal. This is known as modularity.

2.2.2 Principles of Protocol Layering

First Principle

If we want bidirectional communication, we need to make each layer so that it is able to perform two opposite tasks, one in each direction. For example, the third layer task is to listen (in one direction) and talk (in the other direction). The second layer needs to be able to encrypt and decrypt. The first layer needs to send and receive mail.

Second Principle

The second precept that we want to follow in protocol layering is that the two objects under every layer at each website ought to be same. As an instance, the item beneath layer 3 at both sites needs to be a plaintext letter. The item under layer 2 at each sites need to be a cipher textual content letter. The object below layer 1 at each site should be a chunk of mail. Logical connection among every layer is shown in discern. We've layer-to-layer verbal exchange. Maria and Ann can think that there is a logical (imaginary) connection at every layer via which they could send the object constructed from that layer.

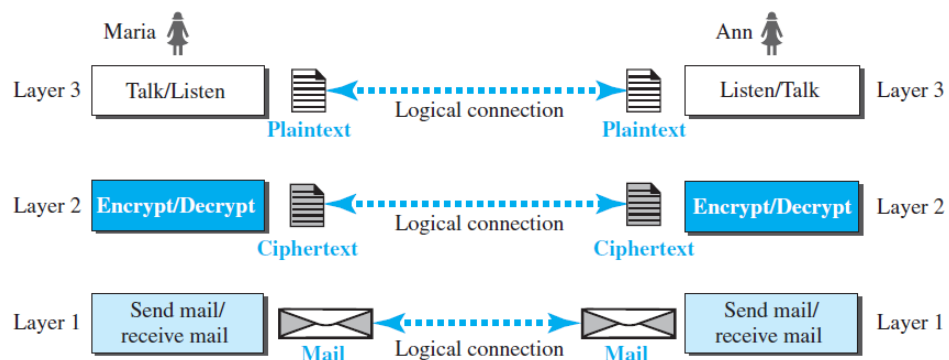
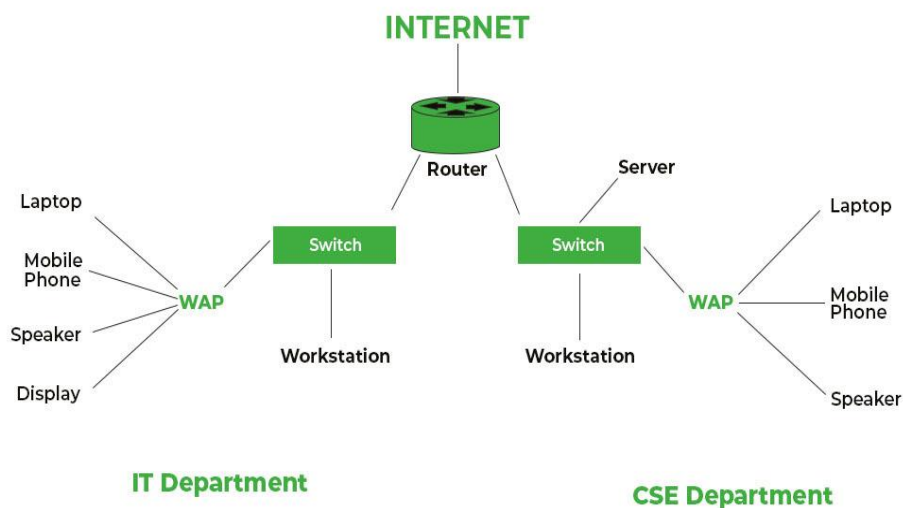


Figure: The concept of logical connection between layers

2.3 LOGICAL CONNECTION

A logical network is a version of the relationship among entities in which each entity is described by means of a node, and the hyperlinks between nodes constitute the connections. The aim of the usage of this model is to understand how specific components of a business enterprise are associated with one another. In other phrases, it may be used to apprehend how responsibilities or sources are allocated within an organization primarily based on its ability (availability) and needs (priority). Logical networks additionally help us see styles in our records that could no longer be glaring while looking at man or woman pieces of facts.

The supply and destination are linked by using a community that is made from various factors and elements. Physical and logical additives are the two classes used to symbolize the elements and elements of a community. Each detail is represented via its own bodily or logical issue, except hosts which don't have any corresponding physical or logical element due to the fact they can't be mapped to every other web hosting kind (additives like switches, routers, and gateways are speculated to be non-host components).



In the diagram, the logical components of a network refer to the information that travels from source to destination. The User information includes the data transported in a framework via the network. Frames have three parts:

- The header
- The data
- The trailer

Frames have destination addresses using which data is traveled to its intended destination.

And if we talk about the physical components of a network then network hardware devices like a switch and the cabling are utilized. These devices make it possible to carry the data from source to destination that makes up the complete physical network.

- The consumer thinks that it's far a single self-contained and impartial entity community although it is probably viable that it's far a part of a big network or a LAN (nearby region network).
- The mapping is one-to-one between a bodily network and a logical community interface/tool.
- Packets may be exchanged among logical interfaces at the same logical community.

- The cause of the use of subnets is to offer communication between logical interfaces that are sharing the same physical interface. For that reason, it can be decided whether or not sending packets from one logical community interface to another logical community interface on the equal logical community is possible the use of the subnet masks.
- as an instance, you may think about a logical community comprising devices from specific networks around the world as in a worldwide agency where the computers of web site managers from specific nations is probably connected as a unmarried logical network to assist quick and smooth communication even though they're separated by way of continents.
- It also reveals its packages in allotted applications due to the involved binding of dispensed additives as a single institution or unmarried entity. Hence, the arrangement of Logical community components is also used in the representation of business environments, departments like engineering, finance, and many others. For example, let us assume the IP address 192.168.0.1/24 could be the logical network.

Therefore, a logical network consists of the following:

Network ID - 192.168.0.1

Subnet mask- /24 = 255.255.255.0

Usable IP addresses= $2^8 - 2 = 254$ IP addresses

Broadcast address - 192.168.0.255

The above information is available in every logical network and is used to determine the usability of the design according to our requirements.

So, select the best possible IP numbering scheme and the subnet mask will find the count of usable IP addresses. In the case of multiple networks within the organization, then let not have overlapping of similar networks. The available networks must be unique to avoid routing issues.

2.4 TCP/IP PROTOCOL SUIT

TCP/IP stands for Transmission manipulate Protocol/net Protocol and is a suite of verbal exchange protocols used to interconnect community devices on the internet. TCP/IP is likewise used as a communications protocol in a private laptop network (an intranet or extranet).

The complete IP suite -- a set of regulations and tactics -- is typically called TCP/IP. TCP and IP are the two principal protocols, though others are blanketed in the suite. The TCP/IP protocol suite features as an abstraction layer between internet programs and the routing and switching fabric.

TCP/IP specifies how data is exchanged over the net with the aid of providing end-to-stop communications that identify the way it must be

damaged into packets, addressed, transmitted, routed and acquired at the destination. TCP/IP calls for little valuable control and is designed to make networks dependable with the capability to get better routinely from the failure of any tool at the community.

The two essential protocols within the IP suite serve precise capabilities. TCP defines how applications can create channels of communication throughout a network. It also manages how a message is assembled into smaller packets earlier than they're then transmitted over the net and reassembled within the right order on the destination address.

IP defines the way to address and path each packet to make sure it reaches the proper vacation spot. Every gateway pc on the community tests this IP deal with to decide wherein to forward the message.

A subnet masks tells a laptop, or different community device, what part of the IP deal with is used to symbolize the community and what element is used to represent hosts, or other computer systems, at the community.

Network cope with translation (NAT) is the virtualization of IP addresses. NAT enables improve safety and decrease the wide variety of IP addresses an organization wishes.

Common TCP/IP protocols include the following:

- **Hypertext Transfer Protocol (HTTP)** handles the communication between a web server and a web browser.
- **HTTP Secure** handles secure communication between a web server and a web browser.
- **File Transfer Protocol** handles transmission of files between computers.

2.5 LAYERED ARCHITECTURE

- The main aim of the layered architecture is to divide the design into small pieces.
- Each lower layer adds its services to the better layer to provide a complete set of offerings to control communications and run the packages.
- It gives modularity and clean interfaces, i.e., presents interplay among subsystems.
- It ensures the independence between layers with the aid of presenting the offerings from decrease to higher layer without defining how the offerings are carried out. Consequently, any modification in a layer will no longer have an effect on the other layers.
- The variety of layers, features, contents of every layer will range from network to community. However, the reason of every layer is to

provide the carrier from lower to a better layer and hiding the information from the layers of ways the services are implemented.

- The simple elements of layered structure are offerings, protocols, and interfaces.
- **Service:** It is a set of actions that a layer provides to the higher layer.
- **Protocol:** It defines a set of rules that a layer uses to exchange the information with peer entity. These rules mainly concern about both the contents and order of the messages used.
- **Interface:** It is a way through which the message is transferred from one layer to another layer.
- In a layer n architecture, layer n on one machine will have a communication with the layer n on another machine and the rules used in a conversation are known as a layer-n protocol.

Let's take an example of the five-layered architecture.

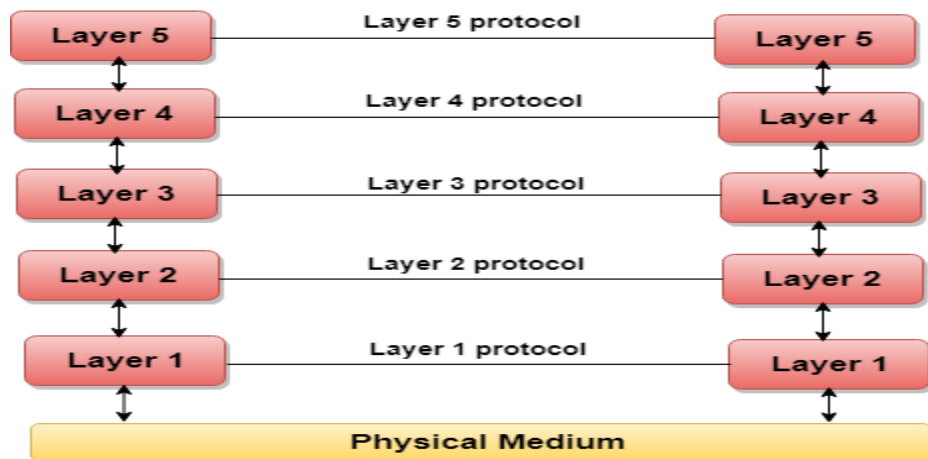


Figure: five-layered architecture

- In case of layered architecture, no information is transferred from layer n of 1 machine to layer n of any other system. Instead, each layer passes the facts to the layer right away just below it, until the lowest layer is reached.
- Under layer 1 is the physical medium through which the actual conversation takes area.
- In a layered structure, unmanageable responsibilities are divided into numerous small and viable tasks.
- The information is exceeded from the higher layer to decrease layer through an interface. Layered architecture offers a smooth-cut interface in order that minimum data is shared amongst distinct layers. It also ensures that the implementation of one layer can be effortlessly replaced with the aid of some other implementation.
- a set of layers and protocols is called network structure.

Why do we require Layered architecture?

- **Divide-and-conquer approach:** Divide-and-conquer approach makes a design process in such a way that the unmanageable tasks are divided into small and manageable tasks. In short, we can say that this approach reduces the complexity of the design.
- **Modularity:** Layered architecture is more modular. Modularity provides the independence of layers, which is easier to understand and implement.
- **Easy to modify:** It ensures the independence of layers so that implementation in one layer can be changed without affecting other layers.
- **Easy to test:** Each layer of the layered architecture can be analyzed and tested individually.

2.6 LAYERS IN THE TCP/IP PROTOCOL SUITE

Protocols are units of policies for message formats and strategies that allow machines and application packages to exchange statistics. Those regulations ought to be followed by every gadget concerned in the conversation in order for the receiving host as a way to apprehend the message. The TCP/IP suite of protocols may be understood in terms of layers (or degrees).

This discerns depicts the layers of the TCP/IP protocol. From the top they are software Layer, shipping Layer, network Layer, community Interface Layer, and hardware.

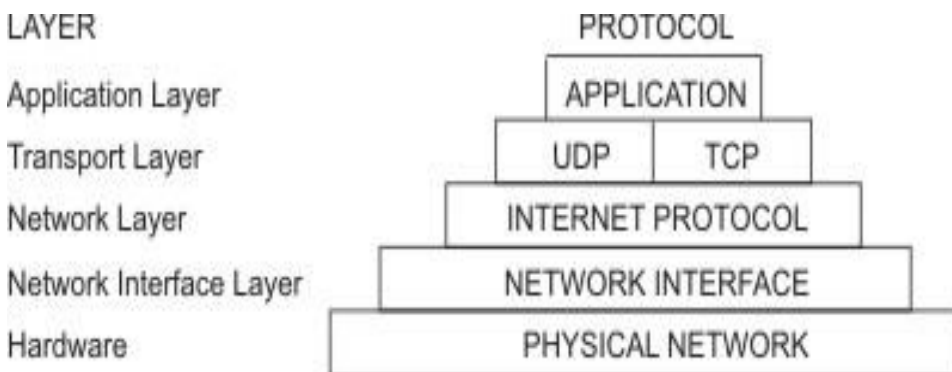


Figure TCP/IP suite of protocols

TCP/IP cautiously defines how facts move from sender to receiver. First, application applications send messages or streams of information to one of the internet transport Layer Protocols, both the consumer Datagram Protocol (UDP) and the Transmission control Protocol (TCP). These protocols acquire the statistics from the utility, divide it into smaller portions known as packets, and add a destination deal with, after which bypass the packets alongside to the next protocol layer, the internet community layer.

The net community layer encloses the packet in an internet Protocol (IP) datagram, places in the datagram header and trailer, comes to a decision where to send the datagram (either directly to a vacation spot in any other case to a gateway), and passes the datagram on to the community Interface layer.

The community Interface layer accepts IP datagrams and transmits them as frames over specific community hardware, consisting of Ethernet or Token-Ring networks.

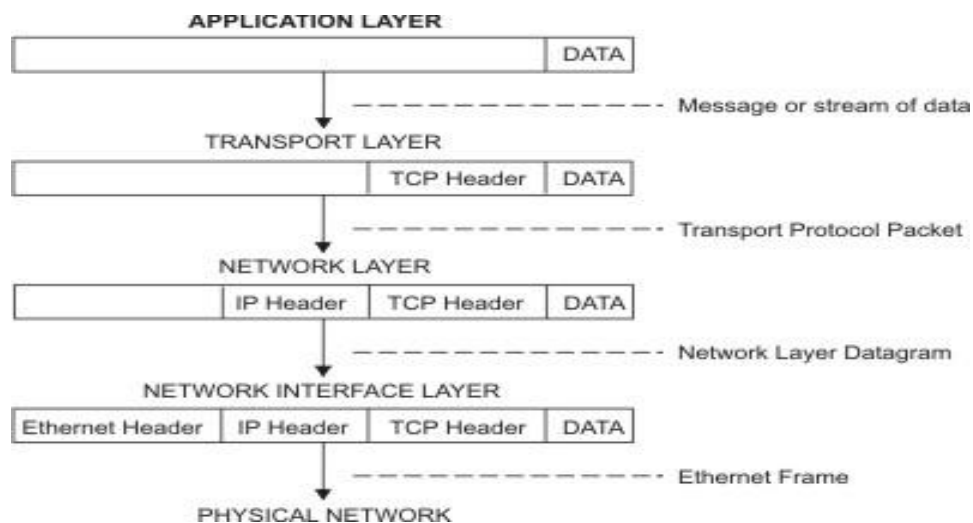


Figure Movement of information from sender application to receiver host

This figure shows the flow of information down the TCP/IP protocol layers from the Sender to the Host.

Frames received by a host go through the protocol layers in reverse. Each layer strips off the corresponding header information, until the data is back at the application layer.

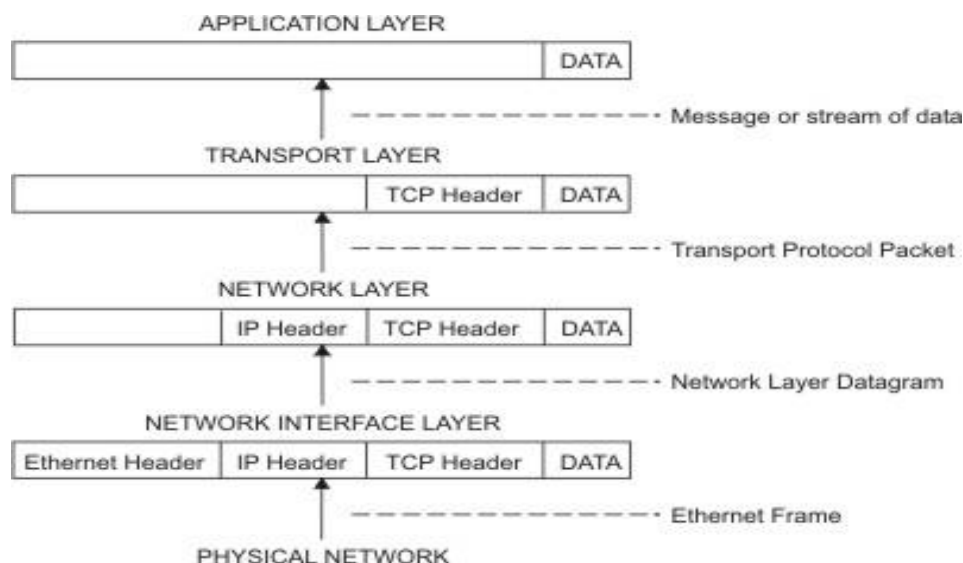


Figure Movement of information from host to application

This figure shows the flow of information up the **TCP/IP** protocol layers from the Host to the Sender. Frames are received by the Network Interface layer (in this case, an Ethernet adapter). The Network Interface layer strips off the Ethernet header, and sends the datagram up to the Network layer. In the Network layer, the Internet Protocol strips off the IP header and sends the packet up to the Transport layer. In the Transport layer, the **TCP** (in this case) strips off the **TCP** header and sends the data up to the Application layer. Hosts on a network send and receive information simultaneously. In Figure more accurately represents a host as it communicates.

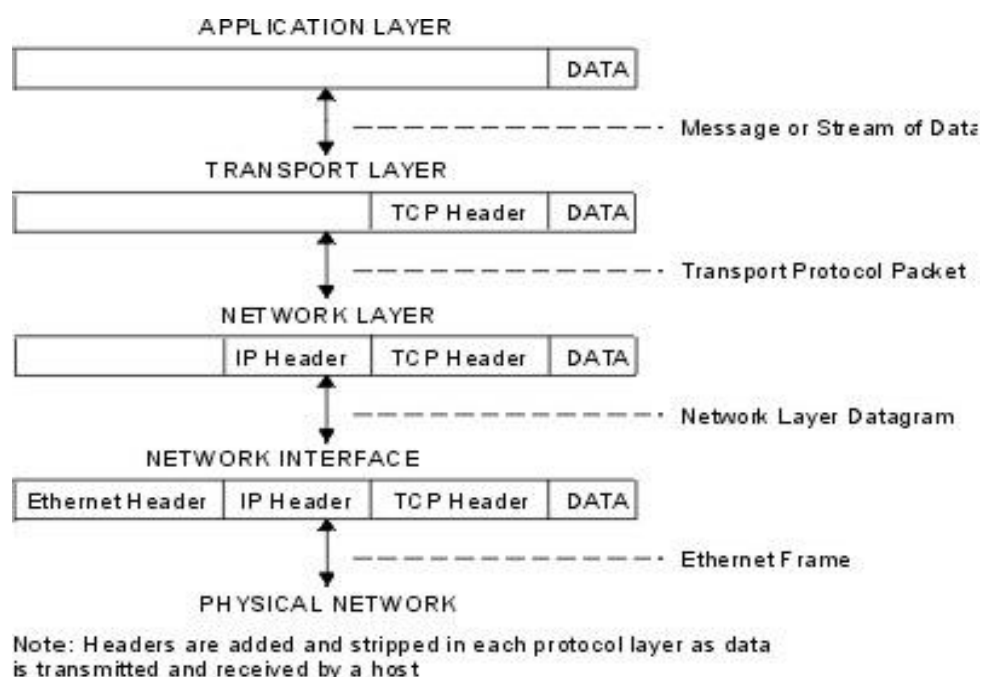
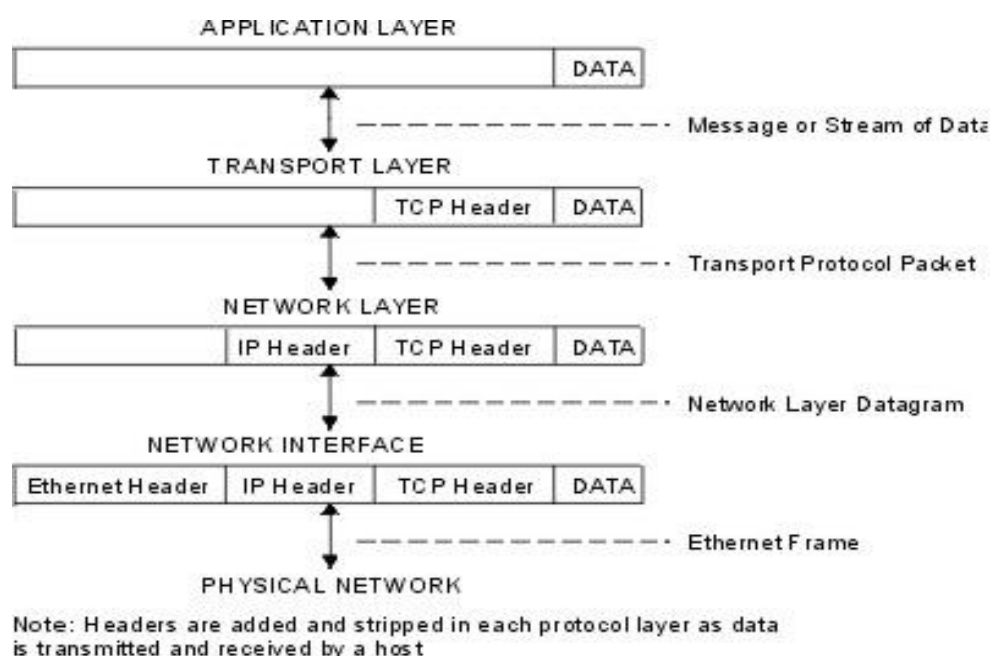


Figure Host data transmissions and receptions



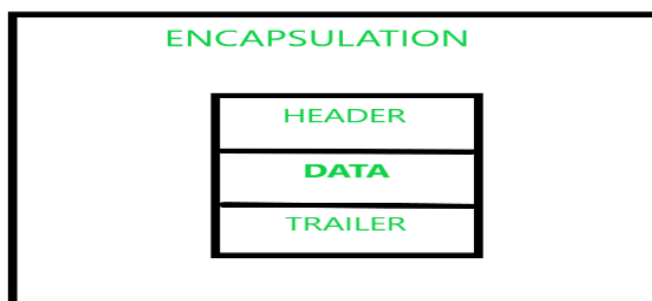
This figure shows data flowing both ways through the **TCP/IP** layers.

- Internet Protocol (IP) version 6
- **Internet Protocol (IP)** version 6 (**IPv6** or *IPng*) is the next generation of **IP** and has been designed to be an evolutionary step from **IP** version 4 (**IPv4**).
- Packet tracing
- Packet tracing is the process by which you can verify the path of a packet through the layers to its destination.
- Network Interface packet headers
- At the Network Interface layer, packet headers are attached to outgoing data.
- Internet network-level protocols
- The Internet network-level protocols handle machine-to-machine communication.
- Internet Transport-Level Protocols
- the **TCP/IP** transport-level protocols allow application programs to communicate with other application programs.
- Internet Application-Level Protocols
- **TCP/IP** implements higher-level Internet protocols at the application program level.
- Assigned Numbers
- for compatibility with the general network environment, well-known numbers are assigned for the Internet versions, networks, ports, protocols, and protocol options. Additionally, well-known names are also assigned to machines, networks, operating systems, protocols, services, and terminals.

2.7 ENCAPSULATION AND DECAPSULATION

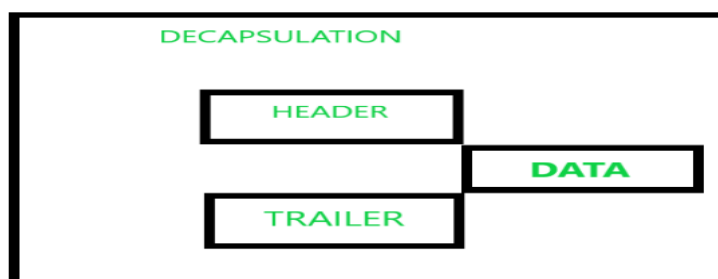
Encapsulation:

Encapsulation refers to attaching new statistics within the utility Layer statistics as it's far passed onto next layers in the TCP/IP version. These extra statistics essentially divided into two parts, Header and Trailer. These are factors attached in an effort to make the transmission extra smoother, on each layer a PDU (Protocol information Unit) is generated. The concept of Encapsulations may be summarized in the screenshot attached in advance.



Decapsulation

Decapsulation refers to the removal of all these additional information and extraction of originally existing data, and this process continues till the last layer i.e. the Application Layer. This process removes fragments of distinct information in each layer as it approaches that layer. Here is the pictorial representation of the whole process.



Difference between encapsulation and decapsulation:

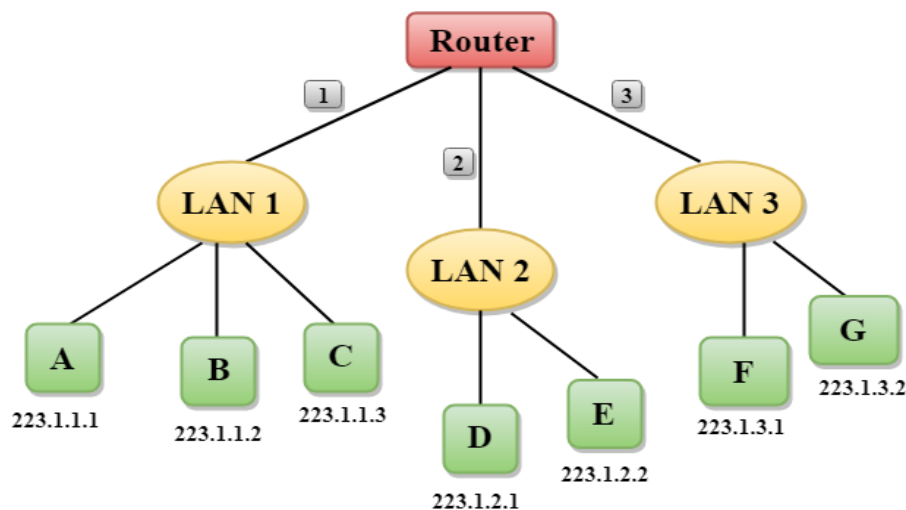
Now, Let us look at the differences between Encapsulation and Decapsulation.

ENCAPSULATION	DECAPSULATION
The data moment starts from the upper Layer and terminates finally on the lowest layer.	Whereas, here the data moves from The lower layer till the upper layer.
The process involves addition of header and Trailer section.	This process involves removal of header and trailer sections
This process executes first and is followed by decapsulation.	This process executes once encapsulation Is finally completed.
It occurs inside the source device.	It occurs inside the destination device.

2.8 ADDRESSING

- Network Addressing is one of the major obligations of the network layer.
- Network addresses are usually logical, i.e., software program-primarily based addresses.
- A host is likewise referred to as stop system that has one hyperlink to the network. The boundary among the host and link is referred to as an interface. Consequently, the host could have simplest one interface.
- A router isn't like the host in that it has two or extra links that connects with it. While a router forwards the datagram, then it forwards the packet to one of the links. The boundary between the router and link is known as an interface, and the router could have a couple of interfaces, one for every of its hyperlinks. Every interface is able to sending and receiving the IP packets, so IP calls for each interface to have an address.
- Every IP address is 32 bits lengthy, and they're represented in the shape of "dot-decimal notation" in which every byte is written within the decimal shape, and they're separated by means of the length. An IP deal with might appear to be 193.32.216.9 in which 193 represents the decimal notation of first eight bits of an deal with, 32 represents the decimal notation of 2d eight bits of an cope with

Let's understand through a simple example.



- In the above figure, a router has three interfaces labeled as 1, 2 & 3 and each router interface contains its own IP address.
- Each host contains its own interface and IP address.

- All the interfaces attached to the LAN 1 is having an IP address in the form of 223.1.1.xxx, and the interfaces attached to the LAN 2 and LAN 3 have an IP address in the form of 223.1.2.xxx and 223.1.3.xxx respectively.
- Each IP address consists of two parts. The first part (first three bytes in IP address) specifies the network and second part (last byte of an IP address) specifies the host in the network.

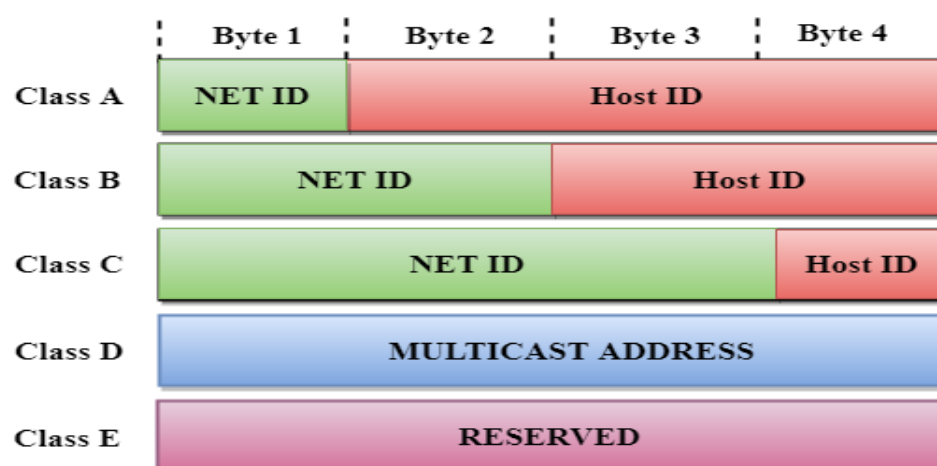
2.8.1 Classful Addressing

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

- Class A
- Class B
- Class C
- Class D
- Class E

An IP address is divided into two parts:

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



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

Class A

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

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

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

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

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



Class B

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

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

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

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

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



Class C

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

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

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

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

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



Class D

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



Class E

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



Rules for assigning Host ID:

The Host ID is used to determine the host within any network. The Host ID is assigned based on the following rules:

- The Host ID must be unique within any network.
- The Host ID in which all the bits are set to 0 cannot be assigned as it is used to represent the network ID of the IP address.
- The Host ID in which all the bits are set to 1 cannot be assigned as it is reserved for the multicast address.

Rules for assigning Network ID:

If the hosts are located within the same local network, then they are assigned with the same network ID. The following are the rules for assigning Network ID:

- The network ID cannot start with 127 as 127 is used by Class A.
- The Network ID in which all the bits are set to 0 cannot be assigned as it is used to specify a particular host on the local network.
- The Network ID in which all the bits are set to 1 cannot be assigned as it is reserved for the multicast address.

2.9 MULTIPLEXING AND DEMULTIPLEXING

Multiplexing

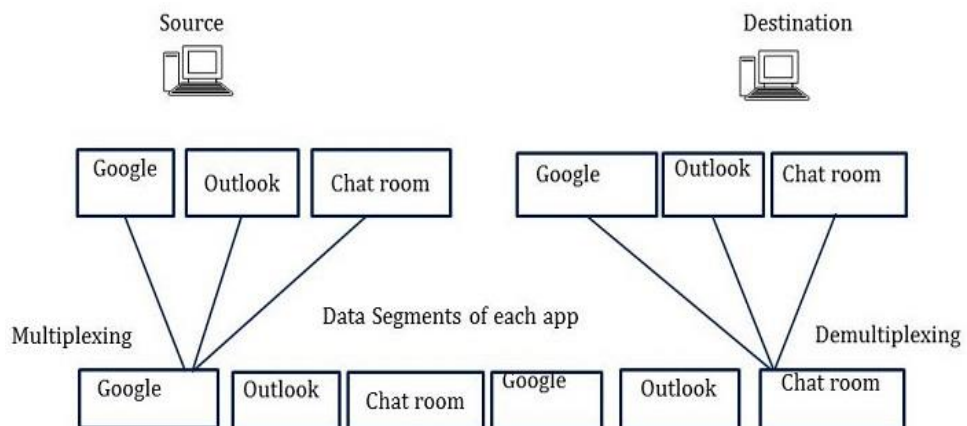
Multiplexing is the process of collecting the data from multiple application processes of the sender, enveloping that data with headers and sending them as a whole to the intended receiver.

- In Multiplexing at the Transport Layer, the data is collected from various application processes. These segments contain the source port number, destination port number, header files, and data.
- These segments are passed to the Network Layer which adds the source and destination IP address to get the datagram.

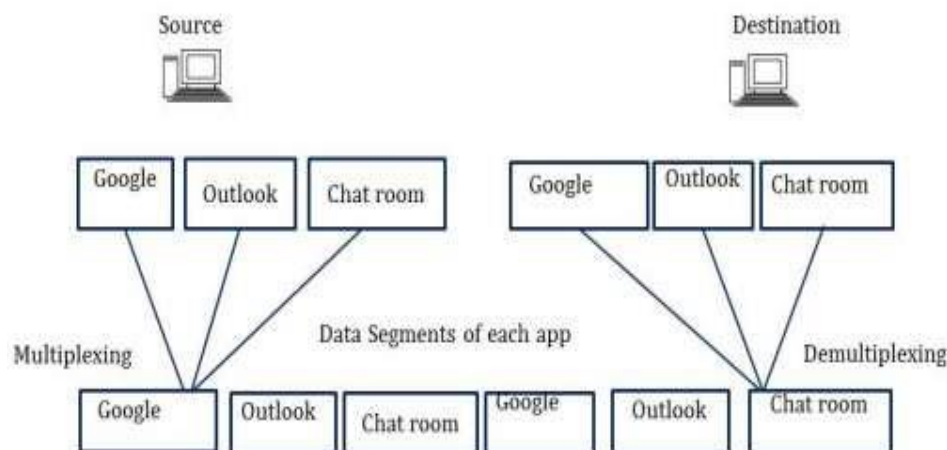
Demultiplexing

Delivering the received segments at the receiver side to the correct app layer processes is called demultiplexing.

- The vacation spot host gets the IP datagrams; each datagram has a source IP address and a destination IP address.
- Each datagram includes 1 shipping layer section.
- Each section has the supply and destination port number.
- The vacation spot host makes use of the IP addresses and port numbers to direct the segment to the proper socket.

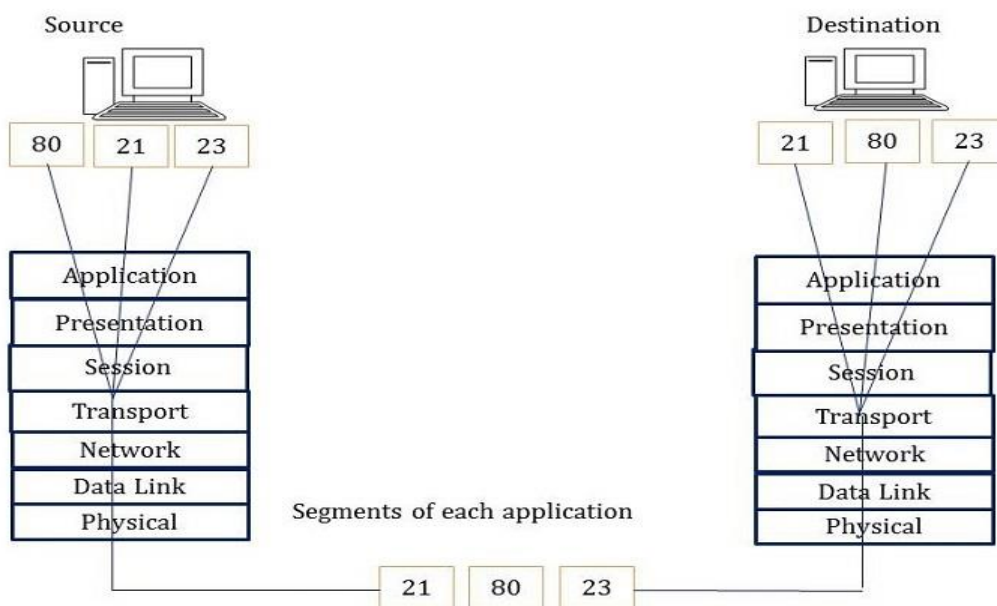


Multiplexing and demultiplexing are just principles that describe the procedure of the transmission of data generated through one of a kind programs concurrently. While the records arrive on the transport layer, every information section is independently processed and sent to its appropriate software in the vacation spot gadget.



The main objective of multiplexing and demultiplexing is to allow us to use a multitude of applications simultaneously.

- The above parent suggests that the source computer is the use of Google, Outlook, and Chat applications on the identical time.
- All the information is forwarded to a vacation spot laptop.
- Every software has a segment put on a twine to be transmitted. It indicates that every one programs are jogging simultaneously.
- Without multiplexing/demultiplexing exists, a consumer can use only one software at a time due to the fact most effective the segments of that utility are put on the cord and transmitted. For clarification, see the determine below –



In the above figure, the Application layer has generated data, and then passed it down to the Transport layer to be segmented.

- After segmenting the data, port numbers are given to each segment to be ready for transmission.
- Then the segments are put on a wire to travel across the network to the destination. This process is called "multiplexing".
- When the transmitted segments reach the Transport layer of the destination, they are automatically sent up to their appropriate applications. This process is called "demultiplexing".

2.10 INTRODUCTION TO PHYSICAL LAYER

The physical layer's principal functionality is transmitting statistics from one pc to any other. That is the bottom layer of the OSI version and is derived beneath the class of hardware layer.

This residue provides an electrical, mechanical, and procedural interface to the transmission medium. It consists of numerous network additives like connectors, receivers, cables and so forth. It defines the transmission of raw bits over the information hyperlink layer. This layer interfaces with the statistics hyperlink layer and carry out image encoding, transmission, reception and interpreting.

in line with community specifications, this accretion controls maximum of the network's physical connections, including wireless transmission, cabling, cabling standards and brands, connectors and types, community interface playing cards, and greater. on the other hand, the bodily layer does not address the actual physical medium (like copper).

Functions of Physical Layer

Following are the various functions performed by the Physical layer of the OSI model.

1. **Representation of Bits:** Data in this layer consists of stream of bits. The bits must be encoded into signals for transmission. It defines the type of encoding i.e. how 0's and 1's are changed to signal.
2. **Data Rate:** This layer defines the rate of transmission which is the number of bits per second.
3. **Synchronization:** It deals with the synchronization of the transmitter and receiver. The sender and receiver are synchronized at bit level.
4. **Interface:** The physical layer defines the transmission interface between devices and transmission medium.
5. **Line Configuration:** This layer connects devices with the medium: Point to Point configuration and Multipoint configuration.
6. **Topologies:** Devices must be connected using the following topologies: Mesh, Star, Ring and Bus.

7. **Transmission Modes:** Physical Layer defines the direction of transmission between two devices: Simplex, Half Duplex, Full Duplex.
8. Deals with baseband and broadband transmission.

Physical Layer Devices:

- Repeaters
- Hubs
- Network Interface cards (NICs)
- Cables and Connectors

Transmission Media

Transmission media is a communication channel that transfers the data from the sender to the recipient using electromagnetic signals. Its primary role is to transport data in the form of bits over a local area network (LAN). There are two different types of media:

- **Guided Media:** It includes all wired/cable communication mediums like optic fiber, twisted pair cable, or coaxial cable.
- **Unguided Media:** It includes wireless or open-air space communication because there is no physical connection between sender and receiver. It includes Wi-Fi communication.

Transmission Impairment

When signals travel through a medium, it tends to distort. Some of the reasons are given below:

- **Attenuation:** It is the loss of energy of the signal. The strength of the signal decreases while passing through the medium with the increase in distance traveled.

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

- **Dispersion:** It is the spread or the overlapping of the signal while passing through any media.
 - **Distortion:** It can be defined as the changes in the form or shape of the signal. This generally happens in composite signals, which are made of different frequencies.
 - **Noise:** Any random disturbance in an analog or digital signal which distorts the information carried by the signal is known as noise. There are several types of noise like crosstalk noise, thermal noise, and impulse noise.
- Signal to noise ratio is calculated as $\text{SNR} = \text{Avg signal power} / \text{Avg noise power}$

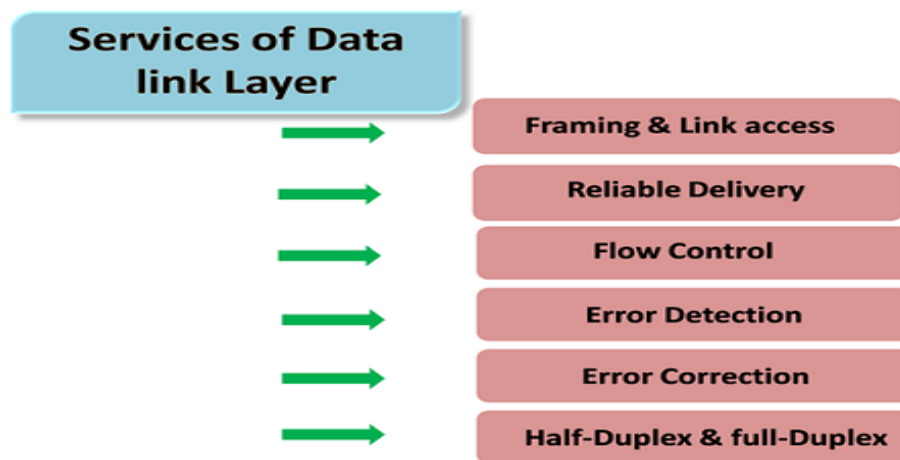
2.11 INTRODUCTION TO DATA-LINK LAYER

The Data Link Layer is the second layer of the OSI Layered model. This layer is one of the most complex layers and has complicated functionalities and liabilities. The Data Link layer hides the information of underlying hardware and represents itself to the upper layer as the medium to talk.

The Data Link layer works amongst hosts which might be directly linked in a few stories. This direct connection might be point-to-point or broadcast. Structures on broadcast network are stated to be on equal link. The work of the Data Link layer has a bent to get greater complicated even as it is coping with multiple hosts on an unshared collision area.

The Data Link layer is accountable for changing data into signals step by step and to send that over the underlying hardware. At the receiving stop, the Data Link layer picks up signals from hardware which can be in the shape of electrical indicators, assembles them in a recognizable frame layout, and passes over to the upper layer. The Data Link layer has sub-layers:

- **Logical Link Control:** It deals with protocols, flow-control, and error control
- **Media Access Control:** It deals with actual control of media



Functionality of Data-link Layer

Data link layer does many tasks on behalf of upper layer. These are:

- Framing
- Data-link layer takes packets from Network Layer and encapsulates them into Frames. Then, it sends each frame bit-by-bit on the hardware. At receiver's end, data link layer picks up signals from hardware and assembles them into frames.
- Addressing

- Data-link layer provides layer-2 hardware addressing mechanism. Hardware address is assumed to be unique on the link. It is encoded into hardware at the time of manufacturing.
- Synchronization
- When data frames are sent on the link, both machines must be synchronized in order to transfer to take place.
- Error Control
- Sometimes signals may have encountered problem in transition and the bits are flipped. These errors are detected and attempted to recover actual data bits. It also provides error reporting mechanism to the sender.
- Flow Control
- Stations on same link may have different speed or capacity. Data-link layer ensures flow control that enables both machine to exchange data on same speed.
- Multi-Access
- When host on the shared link tries to transfer the data, it has a high probability of collision. Data-link layer provides mechanism such as CSMA/CD to equip capability of accessing a shared media among multiple Systems.

Sub-layers of Data Link Layer:

The data link layer is further divided into two sub-layers, which are as follows:

Logical Link Control (LLC):

This sublayer of the data link layer deals with multiplexing, the flow of data among applications and other services, and LLC is responsible for providing error messages and acknowledgments as well.

Media Access Control (MAC):

MAC sublayer manages the device's interaction, responsible for addressing frames, and also controls physical media access.

The data link layer receives the information in the form of packets from the Network layer, it divides packets into frames and sends those frames bit-by-bit to the underlying physical layer.

2.12 INTRODUCTION TO NETWORK LAYER

- The Network Layer is the third layer of the OSI model.
- It handles the service requests from the transport layer and further forwards the service request to the data link layer.
- The network layer translates the logical addresses into physical addresses
- It determines the route from the source to the destination and also manages the traffic problems such as switching, routing and controls the congestion of data packets.
- The main role of the network layer is to move the packets from sending host to the receiving host.

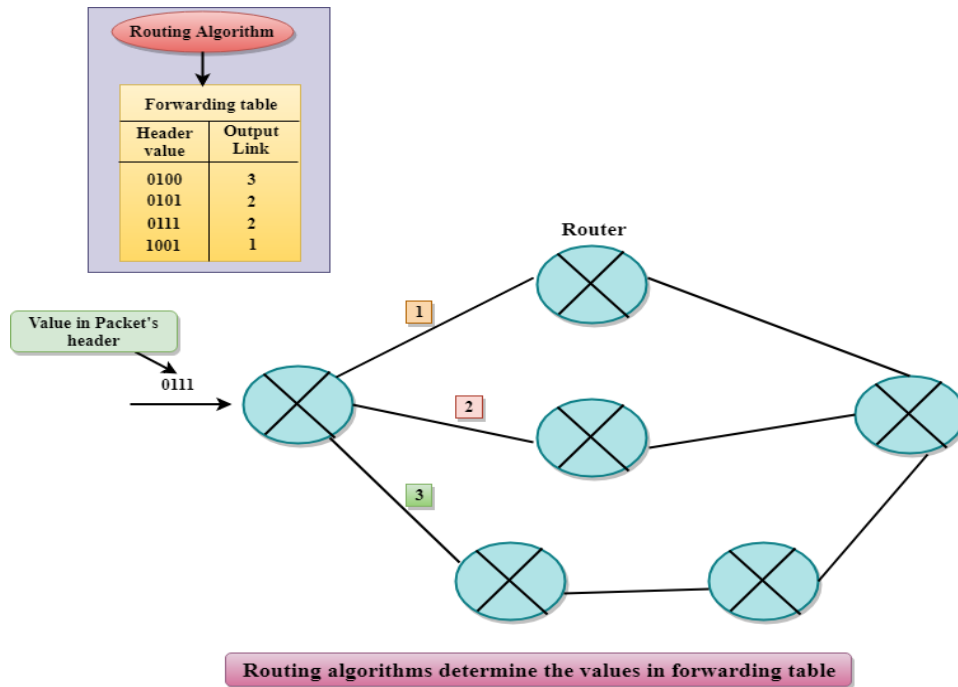
The main functions performed by the network layer are:

- **Routing:** When a packet reaches the router's input link, the router will move the packets to the router's output link. For example, a packet from S1 to R1 must be forwarded to the next router on the path to S2.
- **Logical Addressing:** The data link layer implements the physical addressing and network layer implements the logical addressing. Logical addressing is also used to distinguish between source and destination system. The network layer adds a header to the packet which includes the logical addresses of both the sender and the receiver.
- **Internetworking:** This is the main role of the network layer that it provides the logical connection between different types of networks.
- **Fragmentation:** The fragmentation is a process of breaking the packets into the smallest individual data units that travel through different networks.

Forwarding & Routing

In network layer, a router is used to ahead the packets. Every router has a forwarding table. A router forwards a packet by means of reading a packet's header vicinity and then the use of the header discipline fee to index into the forwarding desk. The price stored inside the forwarding table just like the header subject rate indicates the router's outgoing interface hyperlink to which the packet is to be forwarded.

As an instance, the router with a header subject price of 0111 arrives at a router, after which router indexes this header value into the forwarding desk that determines the output link interface is 2. The router forwards the packet to the interface 2. The routing set of rules determines the values which may be inserted inside the forwarding table. The routing set of rules can be centralized or decentralized.



Services Provided by the Network Layer

- **Guaranteed delivery:** This layer provides the service which guarantees that the packet will arrive at its destination.
- **Guaranteed delivery with bounded delay:** This service guarantees that the packet will be delivered within a specified host-to-host delay bound.
- **In-Order packets:** This service ensures that the packet arrives at the destination in the order in which they are sent.
- **Guaranteed max jitter:** This service ensures that the amount of time taken between two successive transmissions at the sender is equal to the time between their receipt at the destination.
- **Security services:** The network layer provides security by using a session key between the source and destination host. The network layer in the source host encrypts the payloads of datagrams being sent to the destination host. The network layer in the destination host would then decrypt the payload. In such a way, the network layer maintains the data integrity and source authentication services.

Advantages of Network Layer Services

Given below are some benefits of services provided by the network layer:

- By forwarding service of the network layer, the data packets are transferred from one place to another in the network.
- In order to reduce the traffic, the routers in the network layer create collisions and broadcast the domains.
- Failure in the data communication system gets eliminated by packetization.

Disadvantages of Network layer Services

- In the design of the network layer, there is a lack of flow control
- In the network layer, there is a lack of proper error control mechanisms; due to the presence of fragmented data packets the implementation of error control mechanism becomes difficult.

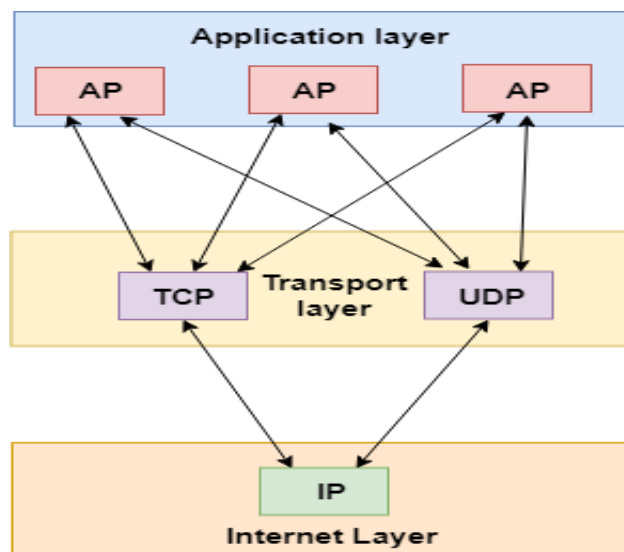
Design Issues with Network Layer

1. A key design issue is **determining how packets are routed from source to destination**. Routes can be based on static tables that are wired into the network and rarely changed. They can also be highly dynamic, being determined anew for each packet, to reflect the current network load.
2. If **too many packets** are present in the subnet at the same time, they will get into one another's way, forming **bottlenecks**. The **control of such congestion** also belongs to the network layer.
3. Moreover, the **quality of service** provided (delay, transmit time, jitter, etc.) is also a network layer issue.
4. When a packet has to **travel from one network to another to get to its destination**, many problems can arise such as:
 - The addressing used by the second network may be different from the first one.
 - The second one may not accept the packet at all because it is too large.
 - The protocols may differ, and so on.
 - 5 It is up to the network layer to overcome all these problems to allow heterogeneous Networks to be interconnected.

2.12 INTRODUCTIONS TO TRANSPORT LAYER

- The transport layer is a 4th layer from the top.
- The main role of the transport layer is to provide the communication services directly to the application processes running on different hosts.
- The transport layer provides a logical communication between application processes running on different hosts. Although the application processes on different hosts are not physically connected, application processes use the logical communication provided by the transport layer to send the messages to each other.
- The transport layer protocols are implemented in the end systems but not in the network routers.

- A computer network provides more than one protocol to the network applications. For example, TCP and UDP are two transport layer protocols that provide a different set of services to the network layer.
- All transport layer protocols provide multiplexing/demultiplexing service. It also provides other services such as reliable data transfer, bandwidth guarantees, and delay guarantees.
- Each of the applications in the application layer has the ability to send a message by using TCP or UDP. The application communicates by using either of these two protocols. Both TCP and UDP will then communicate with the internet protocol in the internet layer. The applications can read and write to the transport layer. Therefore, we can say that communication is a two-way process.

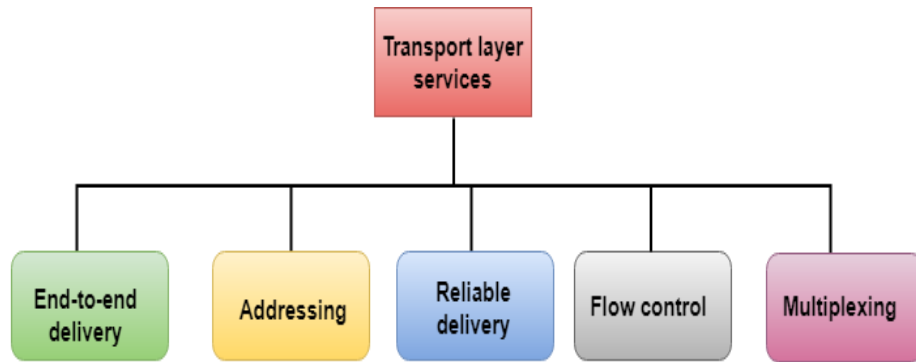


Services provided by the Transport Layer

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

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

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



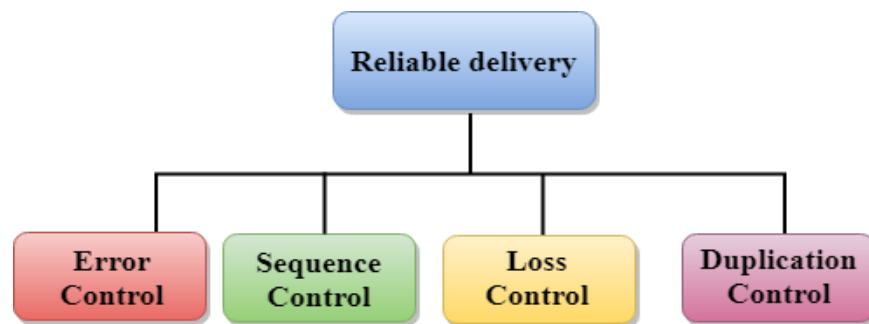
End-to-end delivery:

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

Reliable delivery:

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

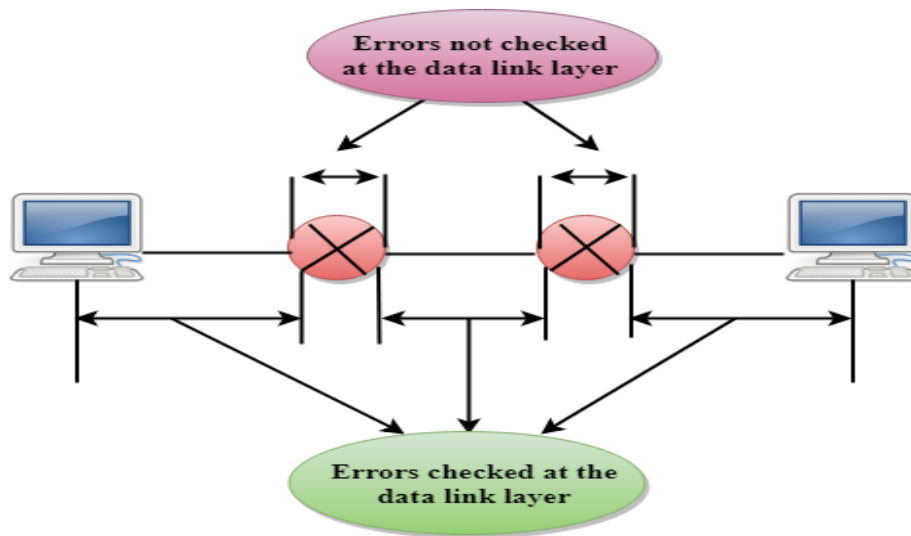
- Error control
- Sequence control
- Loss control
- Duplication control



Error Control

- The primary role of reliability is **Error Control**. In reality, no transmission will be 100 percent error-free delivery. Therefore, transport layer protocols are designed to provide error-free transmission.
- The data link layer also provides the error handling mechanism, but it ensures only node-to-node error-free delivery. However, node-to-node reliability does not ensure the end-to-end reliability.
- The data link layer checks for the error between each network. If an error is introduced inside one of the routers, then this error will not be caught by the data link layer. It only detects those errors that have

been introduced between the beginning and end of the link. Therefore, the transport layer performs the checking for the errors end-to-end to ensure that the packet has arrived correctly.



Sequence Control

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

Loss manipulate

Loss control is a third aspect of reliability. The delivery layer guarantees that each one the fragments of a transmission arrive at the vacation spot, not some of them. At the sending quit, all of the fragments of transmission are given series numbers through a transport layer. Those sequence numbers permit the receivers shipping layer to discover the lacking phase.

Duplication manipulate

Duplication manipulate is the fourth component of reliability. The shipping layer ensures that no reproduction records arrive at the vacation spot. Collection numbers are used to identify the misplaced packets; further, it permits the receiver to become aware of and discard duplicate segments.

Glide manage

Flow control is used to prevent the sender from overwhelming the receiver. If the receiver is overloaded with an excessive amount of records, then the receiver discards the packets and asking for the retransmission of packets. This increases network congestion and hence, reducing the system performance. The delivery layer is answerable for glide control. It makes use of the sliding window protocol that makes the information transmission

extra efficient as nicely as it controls the drift of records so that the receiver does now not grow to be crushed. Sliding window protocol is byte orientated rather than body orientated.

2.13 INTRODUCTION TO APPLICATION LAYER

The software Layer is topmost layer in the Open system Interconnection (OSI) model. This residue affords numerous approaches for manipulating the records (information) which without a doubt allows any form of person to get admission to network effectively. This residue additionally makes a request to its bottom layer, which is presentation layer for receiving diverse types of information from it. The software Layer interface at once interacts with utility and gives commonplace internet software offerings. This residue is largely highest degree of open gadget, which provides services without delay for application manner.

Present Layer=> Application Layer

Presentation Layer

Session Layer

Transport Layer

Network Layer

Data Layer

Physical Layer

Functions of Application Layer :

The Application Layer, as discussed above, being topmost layer in OSI model, performs several kinds of functions which are requirement in any kind of application or communication process. Following are list of functions which are performed by Application Layer of OSI Model –

Data from User <=> Application layer <=> Data from Presentation Layer

- Application Layer presents a facility by means of which customers can forward several e-mails and it additionally gives a garage facility.
- This layer lets in customers to get entry to, retrieve and manipulate documents in a remote laptop.
- It lets in customers to go browsing as a far flung host.
- This layer offers get admission to global information about various offerings.
- This residue affords offerings which consist of: shifting files, distributing results to the user, listing offerings, network assets and so forth.

- It presents protocols that allow software to ship and receive records and present significant facts to customers.
- It handles troubles such as community transparency, aid allocation and so on.
- This deposit serves as a window for customers and application procedures to get entry to community services.
- Application Layer is largely no longer a characteristic; however it plays software layer functions.
- The utility layer is truly an abstraction layer that specifies the shared protocols and interface methods utilized by hosts in a verbal exchange network.
- Utility Layer allows us to perceive verbal exchange partners, and synchronizing verbal exchange.
- This layer lets in customers to have interaction with other software program packages.
- on this layer, data is in visual form, which makes customers in reality recognize facts in place of remembering or visualize the statistics within the binary layout (zero's or 1's).
- This software layer basically interacts with running device (OS) and for this reason similarly preserves the facts in a suitable way.
- This accretion additionally receives and preserves statistics from its preceding layer, which is Presentation Layer (which incorporates in itself the syntax and semantics of the data transmitted).
- The protocols that are used on this application layer depend upon what information customers wish to ship or acquire.
- This software layer, in well known, performs host initialization accompanied by way of far off login to hosts.

Working of Application Layer in the OSI model

In the OSI model, this utility layer is narrower in scope. The software layer inside the OSI version typically acts most effective like the interface that's chargeable for communicating with host-primarily based and consumer-going through packages. this is in evaluation with TCP/IP protocol, in which the layers under the software layer, that's consultation Layer and Presentation layer, are clubbed collectively and shape a simple unmarried layer which is accountable for appearing the functions, which includes controlling the dialogues among computer systems, establishing as well as retaining as well as finishing a selected consultation, supplying facts compression and records encryption and so forth.

at first, customer sends a command to server and while server gets that command, it allocates port number to customer. Thereafter, the patron sends an initiation connection request to server and when server receives request, it offers acknowledgement (ACK) to patron through client has efficaciously hooked up a connection with the server and, consequently, now client has

access to server through which it could either ask server to ship any kinds of documents or different documents or it is able to upload a few documents or documents on server itself.

Features provided by Application Layer Protocols:

To make sure smooth conversation, application layer protocols are applied the same on supply host and vacation spot host.

The following are a number of the features which might be provided by using software layer protocols-

- The application Layer protocol defines system for both events that are involved in verbal exchange.
- These protocols define the kind of message being sent or received from any side (either source host or vacation spot host).
- Those protocols also define basic syntax of the message being forwarded or retrieved.
- Those protocols define the way to ship a message and the predicted response.
- Those protocols additionally define interaction with the following stage.

Application Layer Protocols: The application layer provides several protocols which allow any software program to easily send and acquire records and gift significant data to its customers. The subsequent are a number of the protocols that are furnished by way of the application layer.

- **TELNET:** Telnet stands for Telecommunications network. This protocol is used for managing documents over the internet. It permits the Telnet customers to get right of entry to the resources of Telnet server. Telnet makes use of port variety 23.
- **DNS:** DNS stands for domain call device. The DNS provider translates the area name (decided on via consumer) into the corresponding IP cope with. For example- if you pick the area name as www.abcd.com, then DNS ought to translate it as 192.36.20.eight (random IP cope with written only for know-how purposes). DNS protocol uses the port range fifty three.
- **DHCP:** DHCP stands for Dynamic Host Configuration Protocol. It provides IP addresses to hosts. Each time a bunch tries to register for an IP address with the DHCP server, DHCP server gives plenty of data to the corresponding host. DHCP uses port numbers sixty seven and sixty eight.
- **FTP:** FTP stands for document transfer Protocol. This protocol enables to transfer specific documents from one tool to another. FTP promotes sharing of files via far off computer gadgets with dependable, efficient statistics transfer. FTP uses port quantity 20 for records get admission to and port range 21 for records manipulate.

- SMTP: SMTP stands for simple Mail transfer Protocol. Its miles used to switch electronic mail from one consumer to every other user. SMTP is used by stop users to send emails effectively. SMTP uses port numbers 25 and 587.
- HTTP: HTTP stands for Hypertext transfer Protocol. It's far the inspiration of the arena huge web (WWW). HTTP works at the purchaser server version. This protocol is used for transmitting hypermedia files like HTML. This protocol changed into designed especially for the communications between the web browsers and web servers; however this protocol also can be used for numerous other purposes. HTTP is a stateless protocol (community protocol in which a customer sends requests to server and server responses again as in step with the given state), because of this the server is not chargeable for preserving the previous patron's requests. HTTP makes use of port number 80.
- NFS: NFS stands for community report machine. This protocol permits remote hosts to mount files over a network and interact with those report structures as although they may be installed domestically. NFS uses the port range 2049.
- • SNMP: SNMP stands for simple community management Protocol. This protocol gathers statistics via polling the devices from the community to the management station at fixed or random intervals, requiring them to disclose sure statistics. SNMP makes use of port numbers 161 (TCP) and 162 (UDP).

2.14 LIST OF REFERENCES

<https://www.globalspec.com/reference/23852/203279/internet-administration-governance-and-standards>

<https://www.geeksforgeeks.org/internet-administration/>

<https://www.ietf.org/standards/>

<https://binaryterms.com/internet-standards.html>

<https://www.javatpoint.com/internet>

<https://www.javatpoint.com/computer-network-switching-techniques>

<https://www.javatpoint.com/types-of-computer-network>

<https://www.geeksforgeeks.org/types-of-computer-networks/>

https://www.tutorialspoint.com/data_communication_computer_network/index.htm

<https://www.kdkce.edu.in/pdf/Data-comm-UNIT-I.pdf>

https://www.tutorialspoint.com/computer_concepts/computer_concepts_representation_data_information.htm

<https://www.geeksforgeeks.org/tcp-ip-model/>

<https://www.javatpoint.com/tcp-ip-full-form>

<https://www.javatpoint.com/multiplexing-in-computer-network>

<https://www.geeksforgeeks.org/physical-layer-in-osi-model/>

[https://www.tutorialspoint.com/data communication computer network/physical layer introduction.htm](https://www.tutorialspoint.com/data_communication_computer_network/physical_layer_introduction.htm)

<https://www.geeksforgeeks.org/data-link-layer/>

2.15 EXERCISES

1. Explain the concept of TCP/IP. ...
2. How many layers does TCP/IP architecture have? ...
3. What is an IP datagram? ...
4. How does IP protect data on the network? ...
5. How can we measure the performance of the IP link? ...
6. What is the difference between flow control and error control?
7. What is IP address?
8. What is the function of routing table?
9. Explain the purpose of various layers?
10. Explain various types of connecting devices?
11. What are the three criteria necessary for an effective and efficient network?
12. Why are protocols needed?

DIGITAL AND ANALOG TRANSMISSION

Unit Structure:

- 3.0 Introduction
- 3.1 Digital-To-Digital Conversion
 - 3.1.1 Line Coding
 - 3.1.2 Line Coding Schemes
 - 3.1.3 Unipolar Scheme
 - 3.1.4 Polar Schemes
 - 3.1.5 Return to Zero (Rz)
 - 3.1.6 Biphase
 - 3.1.7 Bipolar Schemes
 - 3.1.8 Multilevel transitions
 - 3.1.9 Multi transitions
 - 3.1.10 Block Coding
 - 3.1.11 4B/5B
 - 3.1.12 Scrambling
 - 3.1.13 HDB3
- 3.2 Analog-To-Digital Conversion
 - 3.2.1 Pulse Code Modulation (PCM)
 - 3.2.2 Sampling
 - 3.2.3 Quantization
 - 3.2.4 Encoding
- 3.3 Delta Modulation (DM)
 - 3.3.1 Modulator
 - 3.3.2 Demodulator
 - 3.3.3 Adaptive DM
- 3.4 Summary
- 3.5 Review Questions

3.0 INTRODUCTION

A computer network is designed to send information from one point to another. This Information needs to be converted to either a digital signal or an analog signal for transmission. In this chapter, we discuss the first choice, conversion to digital signals.

3.1 DIGITAL TO DIGITAL CONVERSION

We said that data can be either digital or analog. We also said that signals that represent data can also be digital or analog. In this section, we see how we can represent digital data by using digital signals. The conversion involves three techniques: line coding, block coding, and scrambling. Line coding is always needed; block coding and scrambling may or may not be needed.

3.1.1 Line Coding

Line coding is the process of converting digital data to digital signals. We assume that data, in the form of text, numbers, graphical images, audio, or video, are stored in computer memory as sequences of bits. Line coding converts a sequence of bits to a digital signal. At the sender, digital data are encoded into a digital signal; at the receiver, the digital data are recreated by decoding the digital signal. Figure 3.1.1 shows the process.

Characteristics

Before discussing different line coding schemes, we address their common characteristics

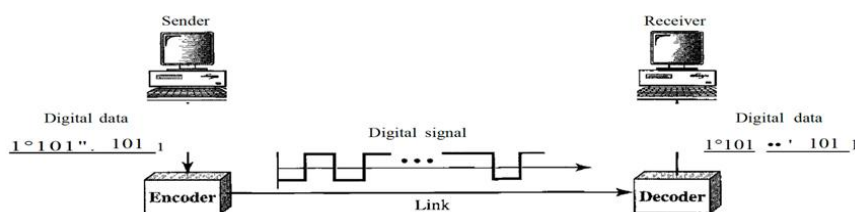


Figure 3.1.1 Line Coding

3.1.2 Line Coding Schemes

We can roughly divide line coding schemes into five broad categories, as shown in figure.

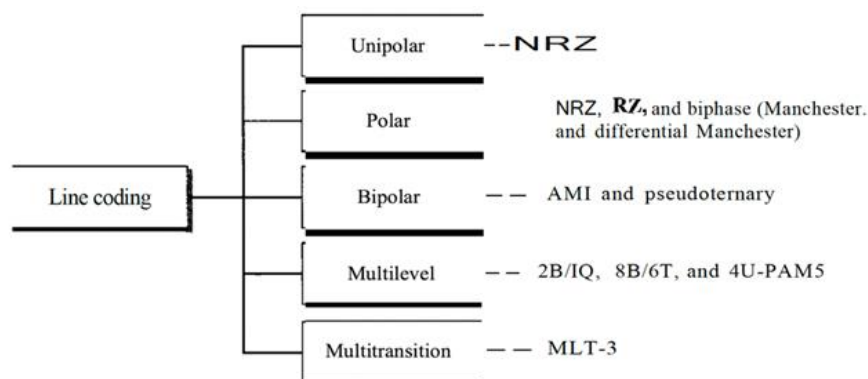


Figure 3.1.2 Line Coding Schemes

There are several schemes in each category. We need to be familiar with all schemes discussed in this section to understand the rest of the book. This section can be used as a reference for schemes encountered later.

3.1.3 Unipolar Scheme

In a unipolar scheme, all the signal levels are on one side of the time axis, either above or below.

NRZ (Non-Return-to-Zero) Traditionally, a unipolar scheme was designed as a non-return-to-zero (NRZ) scheme in which the positive voltage defines bit 1 and the zero voltage defines bit 0. It is called NRZ because the signal does not return to zero at the middle of the bit. Figure below shows a unipolar NRZ scheme.

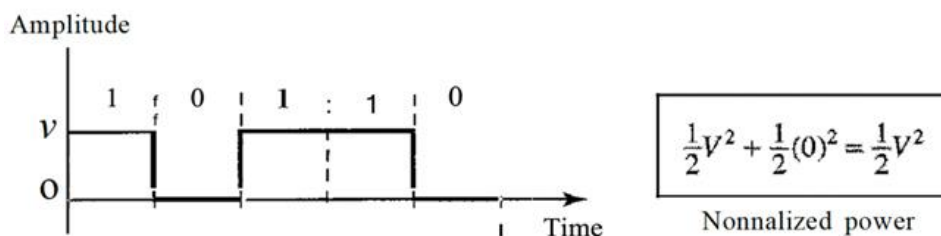


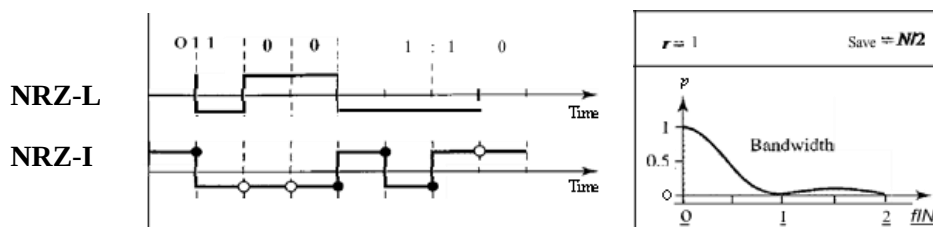
Figure 3.1.3 Unipolar Scheme

Compared with its polar counterpart (see the next section), this scheme is very costly. As we will see shortly, the normalized power (power needed to send 1 bit per unit line resistance) is double that for polar NRZ. For this reason, this scheme is normally not used in data communications today.

3.1.4 Polar Schemes

In polar schemes, the voltages are on both sides of the time axis. For example, the voltage level for 0 can be positive and the voltage level for 1 can be negative.

Non-Return-to-Zero (NRZ) In polar NRZ encoding, we use two levels of voltage amplitude. We can have two versions of polar NRZ: NRZ-L and NRZ-I, as shown in Figure 3.6. The figure also shows the value of r , the average baud rate, and the bandwidth. In the first variation, NRZ-L (NRZ-Level), the level of the voltage determines the value of the bit. In the second variation, NRZ-I (NRZ-Invert), the change or lack of change in the level of the voltage determines the value of the bit. If there is no change, the bit is 0; if there is a change, the bit is 1.



o No inversion: Next bit is 0 • Inversion: Next bit is 1

Figure 3.1.4 Polar and NRZ-I schemes

Let us compare these two schemes based on the criteria we previously defined. Although baseline wandering is a problem for both variations, it is

twice as severe in NRZ-L. If there is a long sequence of 0s or 1s in NRZ-L, the average signal power becomes skewed. The receiver might have difficulty discerning the bit value. In NRZ-I this problem occurs only for a long sequence of 0s. If somehow, we can eliminate the long sequence of 0s, we can avoid baseline wandering. We will see shortly how this can be done.

The synchronization problem (sender and receiver clocks are not synchronized) also exists in both schemes. Again, this problem is more serious in NRZ-L than in NRZ-I. While a long sequence of 0s can cause a problem in both schemes, a long sequence of 1s affects only NRZ-L.

Another problem with NRZ-L occurs when there is a sudden change of polarity in the system. For example, if twisted-pair cable is the medium, a change in the polarity of the wire results in all 0s interpreted as 1s and all 1s interpreted as 0s. NRZ-I does not have this problem. Both schemes have an average signal rate of $N/2$ Bd.

3.1.5 Return to Zero (RZ)

The main problem with NRZ encoding occurs when the sender and receiver clocks are not synchronized. The receiver does not know when one bit has ended, and the next bit is starting. One solution is the return-to-zero (RZ) scheme, which uses three values: positive, negative, and zero. In RZ, the signal changes not between bits but during the bit. In Figure 3.7 we see that the signal goes to 0 in the middle of each bit. It remains there until the beginning of the next bit. The main disadvantage of RZ encoding is that it requires two signal changes to encode a bit and therefore occupies greater bandwidth. The same problem we mentioned, a sudden change of polarity resulting in all 0s interpreted as 1s and all 1s interpreted as 0s, still exist here, but there is no DC component problem. Another problem is the complexity: RZ uses three levels of voltage, which is more complex to create and discern. As a result of all these deficiencies, the scheme is not used today. Instead, it has been replaced by the better-performing Manchester and differential Manchester schemes (discussed next).

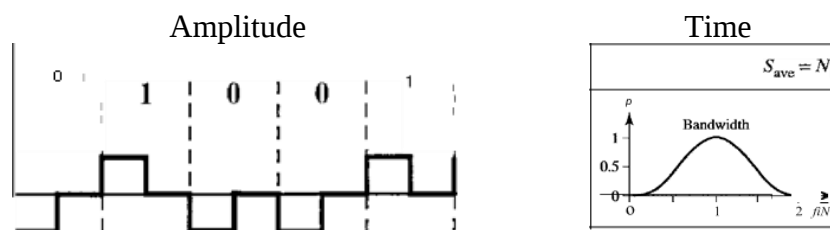


Figure 3.1.5 Polar RZ scheme

3.1.6 Biphasic

Manchester and Differential Manchester The idea of RZ (transition at the middle of the bit) and the idea of NRZ-L are combined into the Manchester scheme. In Manchester encoding, the duration of the bit is divided into two halves. The voltage remains at one level during the first half and moves to the other level in the second half. The transition at the middle of the bit provides synchronization. Differential Manchester, on the

other hand, combines the ideas of RZ and NRZ-I. There is always a transition at the middle of the bit, but the bit values are determined at the beginning of the bit. If the next bit is 0, there is a transition; if the next bit is 1, there is none. Figure 3.1.6 shows both Manchester and differential Manchester encoding.

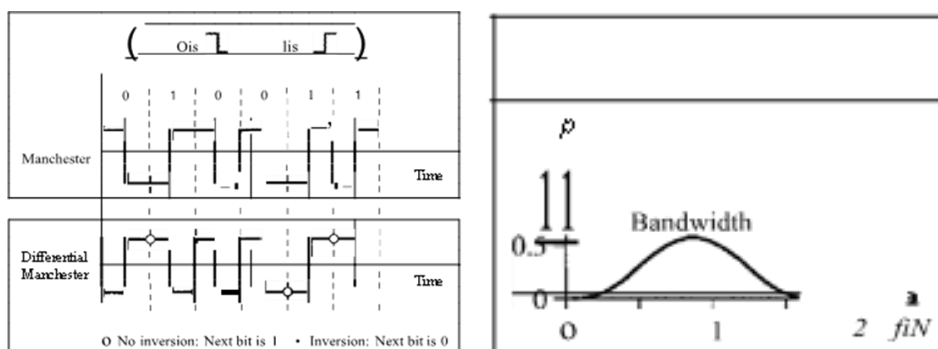


Figure 3.1.6 Polar biphase: Manchester and differential Manchester schemes

The Manchester scheme overcomes several problems associated with NRZ-L, and differential Manchester overcomes several problems associated with NRZ-I. First, there is no baseline wandering. There is no DC component because each bit has a positive and negative voltage contribution. The only drawback is the signal rate. The signal rate for Manchester and differential Manchester is double that for NRZ. The reason is that there is always one transition at the middle of the bit and maybe one transition at the end of each bit. Figure 3.8 shows both Manchester and differential Manchester encoding schemes. Note that Manchester and differential Manchester schemes are also called biphase schemes.

3.1.7 Bipolar Schemes

In bipolar encoding (sometimes called *multilevel binary*), there are three voltage levels: positive, negative, and zero. The voltage level for one data element is at zero, while the voltage level for the other element alternates between positive and negative.

AMI and Pseudoternary Figure 3.1.7 shows two variations of bipolar encoding: AMI and pseudoternary. A common bipolar encoding scheme is called bipolar alternate mark inversion (AMI). In the term *alternate mark inversion*, the word *mark* comes from telegraphy and means 1. So AMI means alternate 1 inversion. A neutral zero voltage represents binary 0. Binary 1s are represented by alternating positive and negative voltages. A variation of AMI encoding is called pseudoternary in which the 1 bit is encoded as a zero voltage and the 0 bit is encoded as alternating positive and negative voltages.

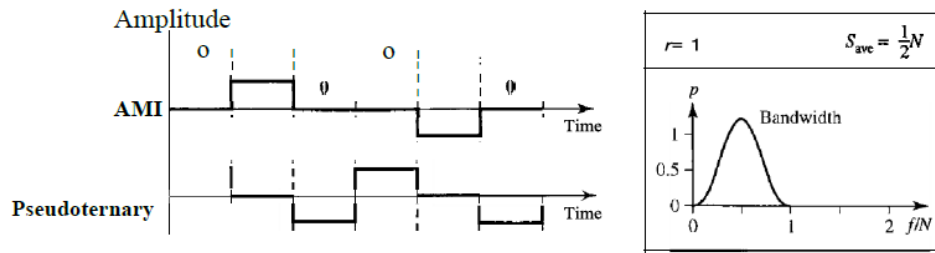


Figure 3.1.7 Bipolar schemes: AMI and pseudoternary

The bipolar scheme was developed as an alternative to NRZ. The bipolar scheme has the same signal rate as NRZ, but there is no DC component. The NRZ scheme has most of its energy concentrated near zero frequency, which makes it unsuitable for transmission over channels with poor performance around this frequency. The concentration of the energy in bipolar encoding is around frequency $N/2$. Figure 3.1.7 shows the typical energy concentration for a bipolar scheme.

One may ask why we do not have DC component in bipolar encoding. We can answer this question by using the Fourier transform, but we can also think about it intuitively. If we have a long sequence of 1s, the voltage level alternates between positive and negative; it is not constant. Therefore, there is no DC component. For a long sequence of 0s, the voltage remains constant, but its amplitude is zero, which is the same as having no DC component. In other words, a sequence that creates a constant zero voltage does not have a DC component.

AMI is commonly used for long-distance communication, but it has a synchronization problem when a long sequence of 0s is present in the data. Later in the chapter, we will see how a scrambling technique can solve this problem.

3.1.8 Multilevel Schemes

The desire to increase the data speed or decrease the required bandwidth has resulted in the creation of many schemes. The goal is to increase the number of bits per baud by encoding a pattern of m data elements into a pattern of n signal elements. We only have two types of data elements (0s and 1s), which means that a group of m data elements can produce a combination of 2^m data patterns. We can have different types of signal elements by allowing different signal levels. If we have L different levels, then we can produce Ln combinations of signal patterns. If $2^m = Ln$, then each data pattern is encoded into one signal pattern. If $2^m < Ln$, data patterns occupy only a subset of signal patterns. The subset can be carefully designed to prevent baseline wandering, to provide synchronization, and to detect errors that occurred during data transmission. Data encoding is not possible if $2^m > Ln$ because some of the data patterns cannot be encoded.

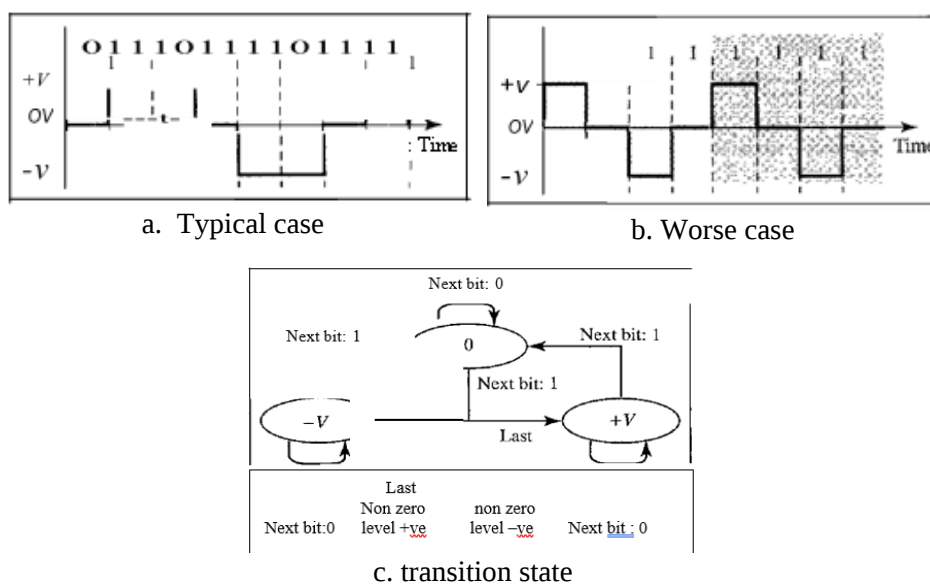
3.1.9 Multiline Transmissions: MLT-3

NRZ-I and differential Manchester are classified as differential encoding but use two transition rules to encode binary data (no inversion, inversion). If we have a signal with more than two levels, we can design a differential encoding scheme with more than two transition rules. MLT-3 is one of them. The multiline transmission, three level (MLT-3) scheme uses three levels ($+V$, 0 , and $-V$) and three transition rules to move between the levels.

1. If the next bit is 0, there is no transition.
2. If the next bit is 1 and the current level is not 0, the next level is 0.
3. If the next bit is 1 and the current level is 0, the next level is the opposite of the last nonzero level.

The behavior of MLT-3 can best be described by the state diagram shown in Figure 3.1.9. The three voltage levels ($-V$, 0 , and $+V$) are shown by three states (ovals). The transition from one state (level) to another is shown by the connecting lines. Figure 3.1.9 also shows two examples of an MLT-3 signal

Figure 3.1.9 Multitransition: MLT-3 scheme



3.1.10 Block Coding

We need redundancy to ensure synchronization and to provide some kind of inherent error detecting. Block coding can give us this redundancy and improve the performance of line coding. In general, block coding changes a block of m bits into a block of n bits, where n is larger than m . Block coding is referred to as an mB/nB encoding technique.

Block coding is normally referred to as mB/nB coding; it replaces each m -bit group with an n -bit group

The slash in block encoding (for example, 4B/5B) distinguishes block encoding from multilevel encoding (for example, 8B6T), which is written

without a slash. Block coding normally involves three steps: division, substitution, and combination. In the division step, a sequence of bits is divided into groups of m bits. For example, in 4B/5B encoding, the original bit sequence is divided into 4-bit groups. The heart of block coding is the substitution step. In this step, we substitute an m -bit group for an n -bit group. For example, in 4B/5B encoding we substitute a 4-bit code for a 5-bit group. Finally, the n -bit groups are combined to form a stream. The new stream has more bits than the original bits. Figure 3.14 shows the procedure.

Division of a stream into m -bit groups Combining n -bit groups into a stream

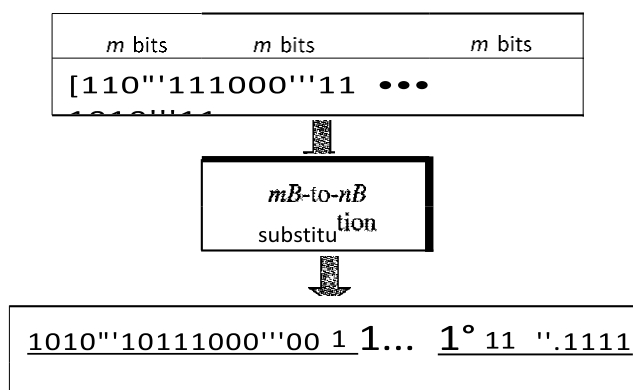


Figure 3.1.10 Block coding concept

3.1.11 4B/5B

The four binary/five binary (4B/5B) coding scheme was designed to be used in combination with NRZ-I. Recall that NRZ-I has a good signal rate, one-half that of the biphase, but it has a synchronization problem. A long sequence of as can make the receiver clock lose synchronization. One solution is to change the bit stream, prior to encoding with NRZ-I, so that it does not have a long stream of as. The 4B/5B scheme achieves this goal. The block-coded stream does not have more than three consecutive as, as we will see later. At the receiver, the NRZ-I encoded digital signal is first decoded into a stream of bits and then decoded to remove the redundancy. Figure 3.1.11 shows the idea.

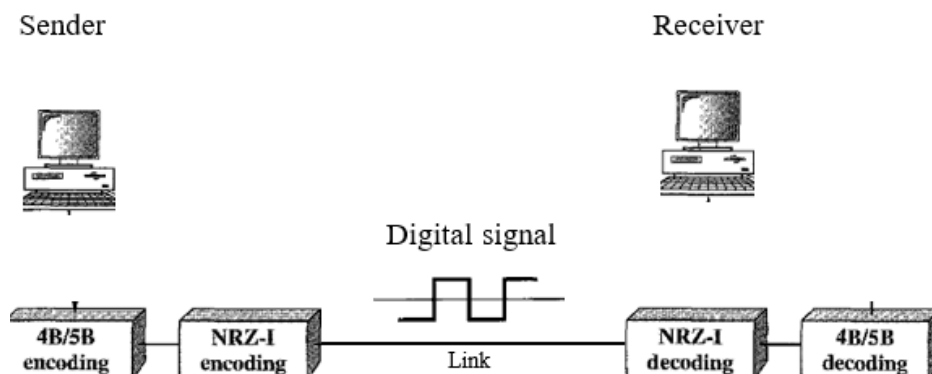


Figure 3.1.11 Using block coding 4B/5B with NRZ-I line coding scheme

In 4B/5B, the 5-bit output that replaces the 4-bit input has no more than one leading zero (left bit) and no more than two trailing zeros (right bits). So when different groups are combined to make a new sequence, there are never more than three consecutive *as*.

A group of 4 bits can have only 16 different combinations while a group of 5 bits can have 32 different combinations. This means that there are 16 groups that are not used for 4B/5B encoding. Some of these unused groups are used for control purposes; the others are not used at all. The latter provide a kind of error detection. If a 5-bit group arrives that belongs to the unused portion, the receiver knows that there is an error in the transmission.

3.1.12 Scrambling

Bi-phase schemes that are suitable for dedicated links between stations in a LAN are not suitable for long-distance communication because of their wide bandwidth requirement. The combination of block coding and NRZ line coding is not suitable for long-distance encoding either, because of the DC component. Bipolar AMI encoding, on the other hand, has a narrow bandwidth and does not create a DC component. However, a long sequence of 0s upsets the synchronization. If we can find a way to avoid a long sequence of 0s in the original stream, we can use bipolar AMI for long distances. We are looking for a technique that does not increase the number of bits and does provide synchronization. We are looking for a solution that substitutes long zero-level pulses with a combination of other levels to provide synchronization. One solution is called scrambling. We modify part of the AMI rule to include scrambling, as shown in Figure 3.1.12. Note that scrambling, as opposed to block coding, is done at the same time as encoding. The system needs to insert the required pulses based on the defined scrambling rules.

Two common scrambling techniques are as follows:

- B8ZS and HDB4.
- R8ZS

Bipolar with S-zero substitution (BSZS) is commonly used in North America. In this technique, eight consecutive zero-level voltages are replaced by the sequence.

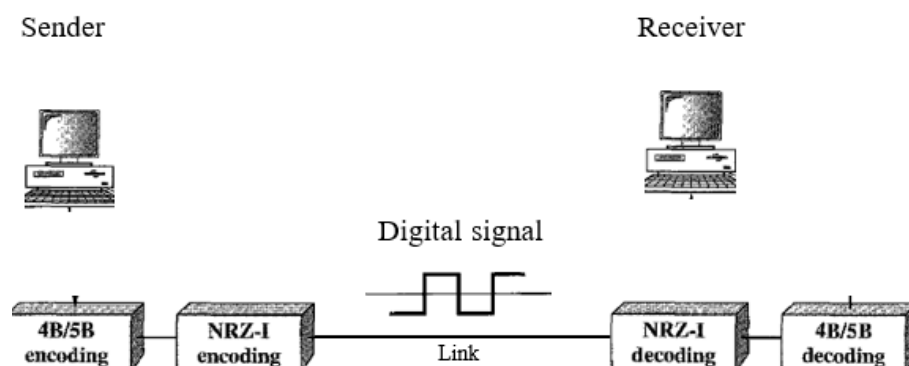
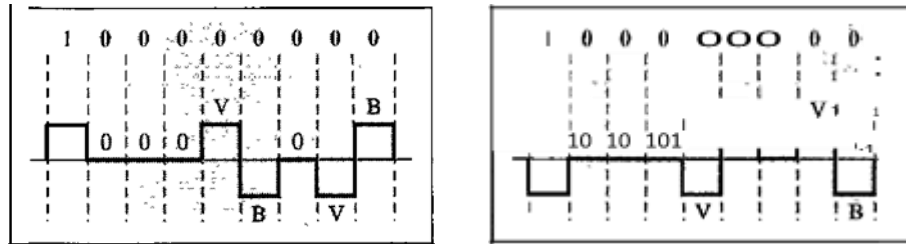


Figure 3.1.12 AMI used with scrambling

OOOVBOVB. The V in the sequence denotes violation; this is a nonzero voltage that breaks an AMI rule of encoding (opposite polarity from the previous). The B in the sequence denotes bipolar; which means a nonzero level voltage in accordance with the AMI rule. There are two cases, as shown below

Two cases of B8ZS scrambling technique



a. Previous level is positive.

b. Previous level is negative.

Note that the scrambling in this case does not change the bit rate. Also, the technique balances the positive and negative voltage levels (two positives and two negatives), which means that the DC balance is maintained. Note that the substitution may change the polarity of a 1 because, after the substitution, AMI needs to follow its rules.

One more point is worth mentioning. The letter V (violation) or B (bipolar) here is relative. The V means the same polarity as the polarity of the previous nonzero pulse; B means the polarity opposite to the polarity of the previous nonzero pulse.

3.1.13 HDB3

High-density bipolar 3-zero (HDB3) is commonly used outside of North America. In this technique, which is more conservative than B8ZS, four consecutive zero-level voltages are replaced with a sequence of OOOV or

The reason for two different substitutions is to maintain the even number of nonzero pulses after each substitution. The two rules can be stated as follows:

- If the number of nonzero pulses after the last substitution is odd, the substitution pattern will be OOOV, which makes the total number of nonzero pulses even.
- If the number of nonzero pulses after the last substitution is even, the substitution pattern will be BOOV, which makes the total number of nonzero pulses even.

There are several points we need to mention here. First, before the first substitution, the number of nonzero pulses is even, so the first substitution is BODY. After this substitution, the polarity of the 1 bit is changed because the AMI scheme, after each substitution, must follow its own rule. After this bit, we need another substitution, which is OOOV because we have only one nonzero pulse (odd) after the last substitution. The third substitution is BOOV because there are no nonzero pulses after the second substitution (even).

The techniques described in Section 3.1 convert digital data to digital signals. Sometimes, however, we have an analog signal such as one created by a microphone or camera. The tendency today is to change an analog signal to digital data. In this section we describe two techniques, pulse code modulation and delta modulation. After the digital data are created (digitization), we can use one of the techniques described in Section 3.1 to convert the digital data to a digital signal

3.2.1 Pulse Code Modulation (PCM)

The most common technique to change an analog signal to digital data (digitization) is called pulse code modulation (PCM). A PCM encoder has three processes, as shown in figure.

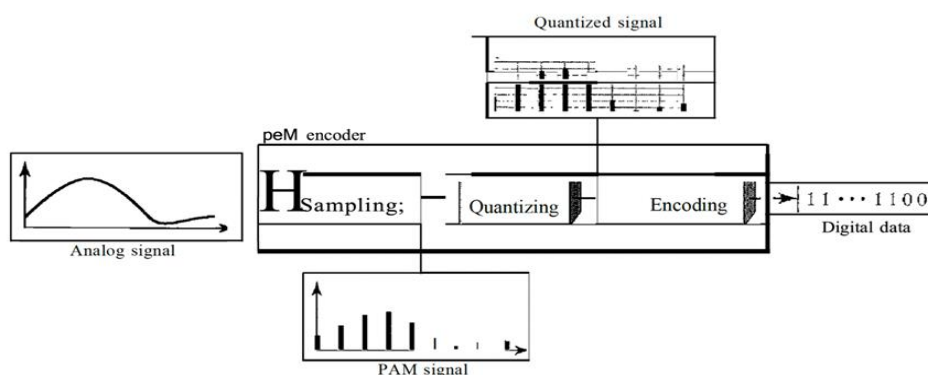


Figure 3.2.1 Components of PCM encoder

1. The analog signal is sampled.
2. The sampled signal is quantized.
3. The quantized values are encoded as streams of bits.

3.2.2 Sampling

The first step in PCM is sampling. The analog signal is sampled every T_s s, where T_s is the sample interval or period. The inverse of the sampling interval is called the sampling rate or sampling frequency and denoted by f_s , where $f_s = 1/T_s$. There are three sampling methods-ideal, natural, and flat-top-as shown in Figure 3.2.2.

In ideal sampling, pulses from the analog signal are sampled. This is an ideal sampling method and cannot be easily implemented. In natural sampling, a high-speed switch is turned on for only the small period of time when the sampling occurs. The result is a sequence of samples that retains the shape of the analog signal. The most common sampling method, called sample and hold, however, creates flat-top samples by using a circuit.

The sampling process is sometimes referred to as pulse amplitude modulation (PAM). We need to remember, however, that the result is still an analog signal with nonintegral values.

Sampling Rate, one important consideration is the sampling rate or frequency. What are the restrictions on T_s ? This question was elegantly answered by Nyquist. According to the Nyquist theorem, to reproduce the original analog signal, one necessary condition is that the sampling rate be at least twice the highest frequency in the original signal

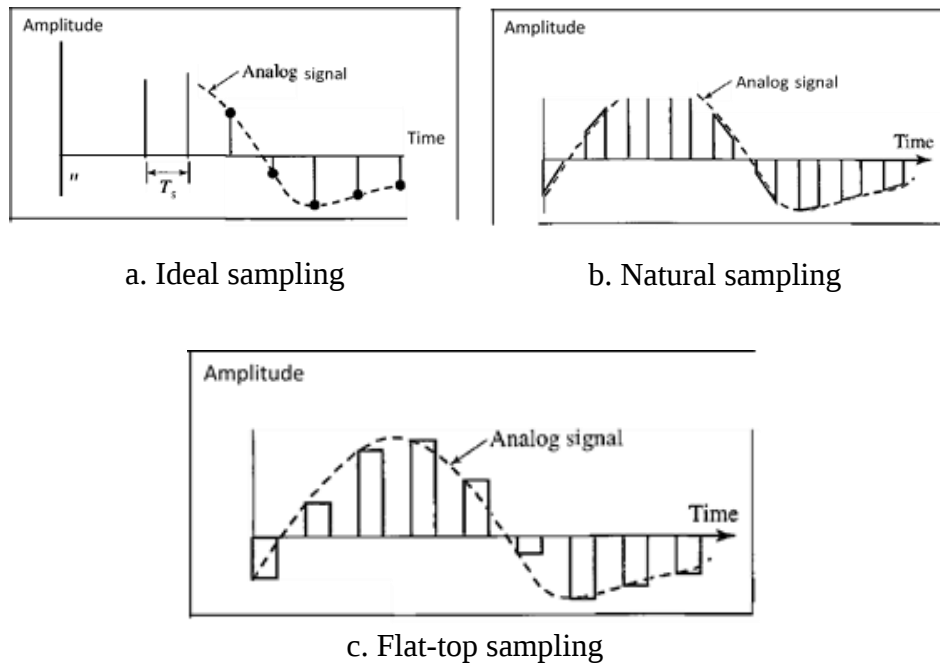


Figure 3.2.2 *Three different sampling methods for PCM*

According to the Nyquist theorem, the sampling rate must be at least 2 times the highest frequency contained in the signal.

3.2.3 Quantization

The result of sampling is a series of pulses with amplitude values between the maximum and minimum amplitudes of the signal. The set of amplitudes can be infinite with nonintegral values between the two limits. These values cannot be used in the encoding process. The following are the steps in quantization:

1. We assume that the original analog signal has instantaneous amplitudes between V_{min} and V_{max}
2. We divide the range into L zones, each of height Δ (delta).
3. We assign quantized values of 0 to $L - 1$ to the midpoint of each zone.
4. We approximate the value of the sample amplitude to the quantized values.

As a simple example, assume that we have a sampled signal and the sample amplitudes are between -20 and $+20$ V. We decide to have eight levels ($L = 8$). This means that $\Delta = 5$ V

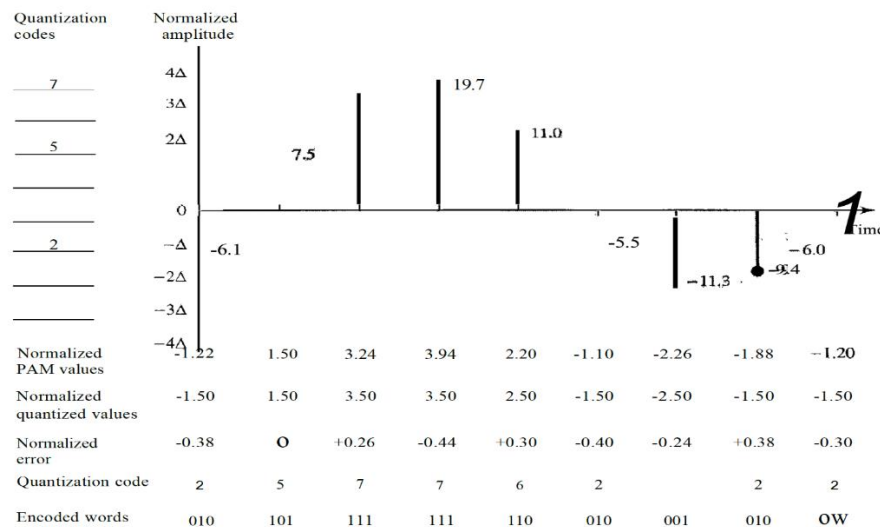


Figure 3.2.3 Quantization and encoding of a sampled signal

We have shown only nine samples using ideal sampling (for simplicity). The value at the top of each sample in the graph shows the actual amplitude. In the chart, the first row is the normalized value for each sample (actual amplitude/ Δ). The quantization process selects the quantization value from the middle of each zone. This means that the normalized quantized values (second row) are different from the normalized amplitudes. The difference is called the normalized error (third row). The fourth row is the quantization code for each sample based on the quantization levels at the left of the graph. The encoded words (fifth row) are the final products of the conversion.

Quantization Levels In the previous example, we showed eight quantization levels. The choice of L , the number of levels, depends on the range of the amplitudes of the analog signal and how accurately we need to recover the signal. If the amplitude of a signal fluctuates between two values only, we need only two levels; if the signal, like voice, has many amplitude values, we need more quantization levels. In audio digitizing, L is normally chosen to be 256; in video it is normally thousands. Choosing lower values of L increases the quantization error if there is a lot of fluctuation in the signal.

Quantization Error One important issue is the error created in the quantization process. (Later, we will see how this affects high-speed modems.) Quantization is an approximation process. The input values to the quantizer are the real values; the output values are the approximated values. The output values are chosen to be the middle value in the zone. If the input value is also at the middle of the zone, there is no quantization error; otherwise, there is an error. In $\Delta/2$, previous example, the normalized amplitude of the third sample is 4.24, but the normalized quantized value is 4.50. This means that there is an error of +0.26. The value of the error for any sample is less than

In other words, we have $-\Delta/2 \leq e \leq \Delta/2$.

The quantization error changes the signal-to-noise ratio of the signal, which in turn reduces the upper limit capacity according to Shannon.

It can be proven that the contribution of the quantization error to the SNR_{dB} of the signal depends on the number of quantization levels L , or the bits per sample nb' as shown in the following formula:

$$SNR_{dB} = 6.02nb + 1.76 \text{ dB}$$

3.2.4 Encoding

The last step in PCM is encoding. After each sample is quantized and the number of bits per sample is decided, each sample can be changed to an l lb-bit code word. In Figure 1.26 the encoded words are shown in the last row. A quantization code of 2 is encoded as 010; 5 is encoded as 101; and so on. Note that the number of bits for each sample is determined from the number of quantization levels. If the number of quantization levels is L , the number of bits is $lbb = \log_2 L$. In our example L is 8 and lbb is

therefore 4. The bit rate can be found from the formula.

Bit rate :::: sampling rate x number of bits per sample :::: $f_s \times nb$

3.3 DELTA MODULATION (DM)

PCM is a very complex technique. Other techniques have been developed to reduce the complexity of PCM. The simplest is delta modulation. PCM finds the value of the signal amplitude for each sample; DM finds the change from the previous sample. Figure 3.3 shows the process. Note that there are no code words here; bits are sent one after another.

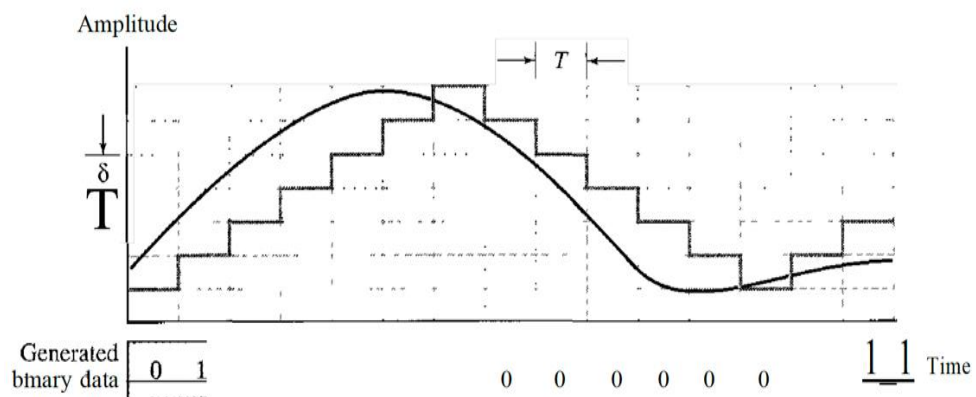


Figure 3.3 Process of Delta Modulation

3.3.1 Modulator

The modulator is used at the sender site to create a stream of bits from an analog signal. The process records the small positive or negative changes, called delta δ . If the delta is positive, the process records a 1; if it is negative, the process records a 0. However, the process needs a base against which the analog signal is compared. The modulator builds a second signal that resembles a staircase. Finding the change is then reduced to comparing the input signal with the gradually made staircase signal. Figure 3.3.1 shows a diagram of the process.

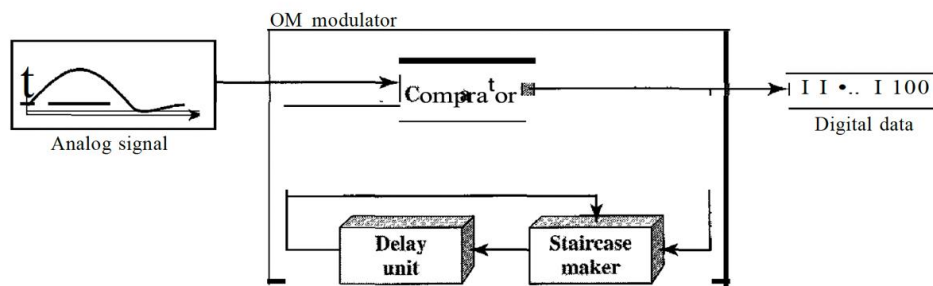


Figure 3.3.1 Delta modulation components

The modulator, at each sampling interval, compares the value of the analog signal with the last value of the staircase signal. If the amplitude of the analog signal is larger, the next bit in the digital data is 1; otherwise, it is 0. The output of the comparator, however, also makes the staircase itself. If the next bit is 1, the staircase maker moves the last point of the staircase signal δ up; if the next bit is 0, it moves it δ down. Note that we need a delay unit to hold the staircase function for a period between two comparisons.

3.3.2 Demodulator

The demodulator takes the digital data and, using the staircase maker and the delay unit, creates the analog signal. The created analog signal, however, needs to pass through a low-pass filter for smoothing. Figure 3.3.2 shows the schematic diagram.

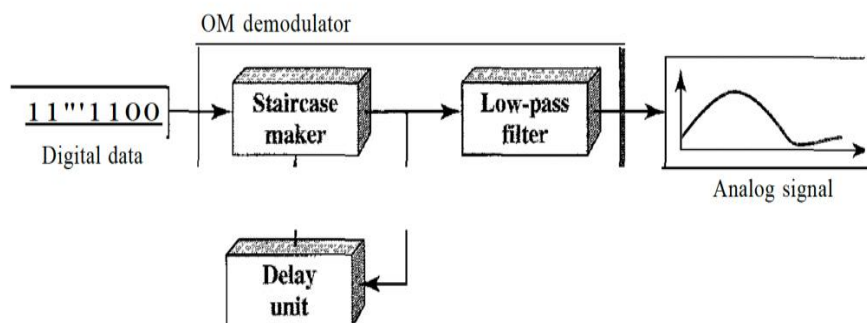


Figure 3.3.2 Delta modulation components

3.3.3 Adaptive DM

A better performance can be achieved if the value of δ is not fixed. In adaptive delta modulation, the value of δ changes according to the amplitude of the analog signal.

Quantization Error

It is obvious that DM is not perfect. Quantization error is always introduced in the process. The quantization error of DM, however, is much less than that for PCM.

3.4 SUMMARY

- Digital-to-digital conversion involves three techniques: line coding, block coding, and scrambling.
- Line coding is the process of converting digital data to a digital signal.
- We can roughly divide line coding schemes into five broad categories: unipolar polar, bipolar, multilevel, and multitransition.
- Block coding provides redundancy to ensure synchronization and inherent error detection. Block coding is normally referred to as mB/nB coding; it replaces each. m-bit group with an n-bit group.
- The most common technique to change an analog signal to digital data (digitization) is called pulse code modulation (PCM).

3.5 REVIEW QUESTIONS

- Q1. List three techniques of digital-to-digital conversion.
- Q2. Distinguish between a signal element and a data element.
- Q3. Define baseline wandering and its effect on digital transmission.
- Q4. Define a DC component and its effect on digital transmission.
- Q5. Define the characteristics of a self-synchronizing signal.
- Q6. List five line coding schemes discussed in this book.
- Q7. Define block coding and give its purpose.
- Q8. Define scrambling and give its purpose.
- Q9. Compare and contrast PCM and DM.

TRANSMISSION MODES

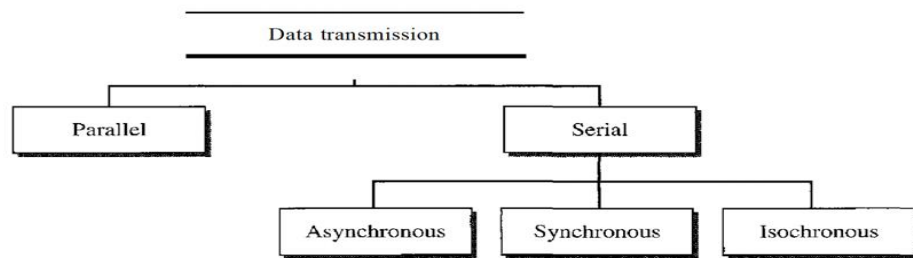
Unit Structure:

- 4.0 Introduction
- 4.1 Data Transmissions
 - 4.1.1 Parallel Transmission
 - 4.1.2 Serial Transmission
 - 4.1.3 Asynchronous transmission
 - 4.1.4 Synchronous transmission
 - 4.1.5 Isochronous transmissions
- 4.2 Analog Transmission
 - 4.2.1 DIGITAL-TO-ANALOG CONVERSION
 - 4.2.2 Aspects of Digital-to-Analog Conversion
 - 4.2.3 Amplitude Shift Keying (ASK)
 - 4.2.4 Binary ASK (BASK)
 - 4.2.5 Multilevel ASK
- 4.3 Frequency Shift Keying
 - 4.3.1 Binary FSK (BFSK)
 - 4.3.2 Multilevel FSK
- 4.4 Phase Shift Keying (PSK)
 - 4.4.1 Binary PSK (BPSK)
 - 4.4.2 Quadrature PSK (QPSK)
- 4.5 Constellation Diagram
 - 4.5.1 Concept of Constellation Diagram
- 4.6 Quadrature Amplitude Modulation
- 4.7 Review Questions

4.0 INTRODUCTION

Of primary concern when we are considering the transmission of data from one device to another is the wiring, and of primary concern when we are considering the wiring is the data stream. Do we send 1 bit at a time; or do we group bits into larger groups and, if so, how? The transmission of binary data across a link can be accomplished in either parallel or serial mode. In parallel mode, multiple bits are sent with each clock tick. In serial mode, 1 bit is sent with each clock tick. While there is only one way to send parallel

data, there are three subclasses of serial transmission: asynchronous, synchronous, and isochronous (see Figure below).



4.1 DATA TRANSMISSIONS

4.1.1 Parallel Transmission

Binary data, consisting of 1s and 0s, may be organized into groups of n bits each. Computers produce and consume data in groups of bits much as we conceive of and use spoken language in the form of words rather than letters. By grouping, we can send data n bits at a time instead of 1. This is called parallel transmission.

The mechanism for parallel transmission is a conceptually simple one: Use n wires to send n bits at one time. That way each bit has its own wire, and all n bits of one group can be transmitted with each clock tick from one device to another. Figure 4.1 shows how parallel transmission works for $n = 8$. Typically, the eight wires are bundled in a cable with a connector at each end.

The advantage of parallel transmission is speed. All else being equal, parallel transmission can increase the transfer speed by a factor of n over serial transmission.

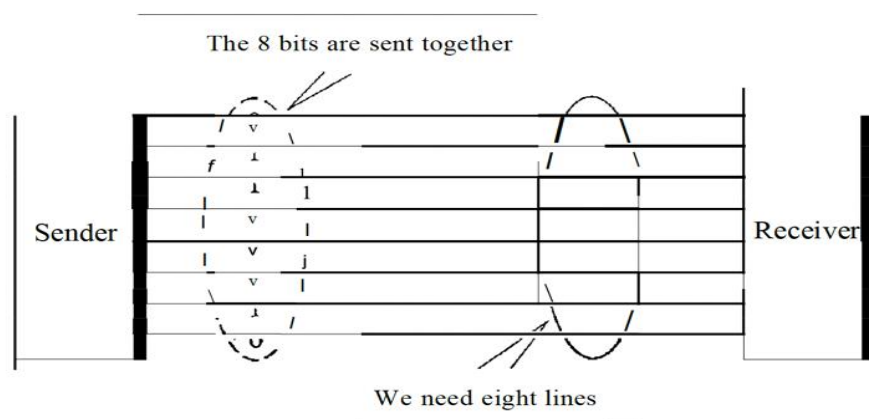


Figure 4.1.1 Parallel Transmission

But there is a significant disadvantage: cost. Parallel transmission requires n communication lines (wires in the example) just to transmit the data stream. Because this is expensive, parallel transmission is usually limited to short distances.

4.1.2 Serial Transmission

In serial transmission one bit follows another, so we need only one communication channel rather than n to transmit data between two communicating devices (see Figure 5.2).

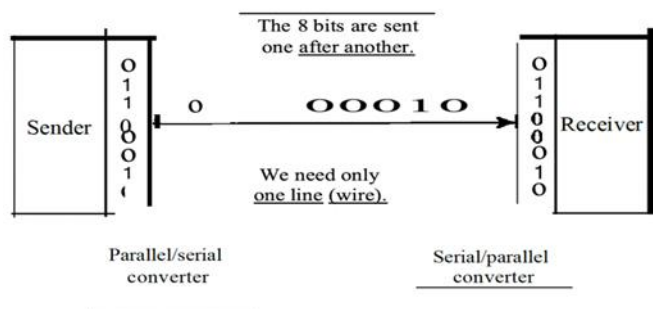


Figure 5.1.2 Serial transmission

The advantage of serial over parallel transmission is that with only one communication channel, serial transmission reduces the cost of transmission over parallel by roughly a factor of n .

Since communication within devices is parallel, conversion devices are required at the interface between the sender and the line (parallel-to-serial) and between the line and the receiver (serial-to-parallel).

Serial transmission occurs in one of three ways: asynchronous, synchronous, and isochronous.

4.1.3 Asynchronous Transmission

Asynchronous transmission is so named because the timing of a signal is unimportant. Instead, information is received and translated by agreed upon patterns. As long as those patterns are followed, the receiving device can retrieve the information without regard to the rhythm in which it is sent. Patterns are based on grouping the bit stream into bytes. Each group, usually 8 bits, is sent along the link as a unit. The sending system handles each group independently, relaying it to the link whenever ready, without regard to a timer.

Without synchronization, the receiver cannot use timing to predict when the next group will arrive. To alert the receiver to the arrival of a new group, therefore, an extra bit is added to the beginning of each byte. This bit, usually a 0, is called the start bit. To let the receiver know that the byte is finished, 1 or more additional bits are appended to the end of the byte. These bits, usually 1 s, are called stop bits. By this method, each byte is increased in size to at least 10 bits, of which 8 bits is information and 2 bits or more are signals to the receiver. In addition, the transmission of each byte may then be followed by a gap of varying duration. This gap can be represented either by an idle channel or by a stream of additional stop bits.

In asynchronous transmission, we send 1 start bit (0) at the beginning and 1 or more stop bits (1s) at the end of each byte. There may be a gap between each byte.

The start and stop bits and the gap alert the receiver to the beginning and end of each byte and allow it to synchronize with the data stream. This mechanism is called *asynchronous* because, at the byte level, the sender and receiver do not have to be synchronized. But within each byte, the receiver must still be synchronized with the incoming bit stream. That is, some synchronization is required, but only for the duration of a single byte. The receiving device resynchronizes at the onset of each new byte. When the receiver detects a start bit, it sets a timer and begins counting bits as they come in. After n bits, the receiver looks for a stop bit. As soon as it detects the stop bit, it waits until it detects the next start bit.

Asynchronous here means "asynchronous at the byte **level**," but the bits are still synchronized; their durations are the same.

Figure 4.1.3 is a schematic illustration of asynchronous transmission. In this example, the start bits are 0s, the stop bits are 1s, and the gap is represented by an idle line rather than by additional stop bits. The addition of stop and start bits and the insertion of gaps into the bit stream make asynchronous transmission slower than forms of transmission that can operate without the addition of control information. But it is cheap and effective, two advantages that make it an attractive choice for situations such as low-speed communication. For example, the connection of a keyboard to a computer is a natural application for asynchronous transmission. A user types only one character at a time, types extremely slowly in data processing terms, and leaves unpredictable gaps of time between each character.

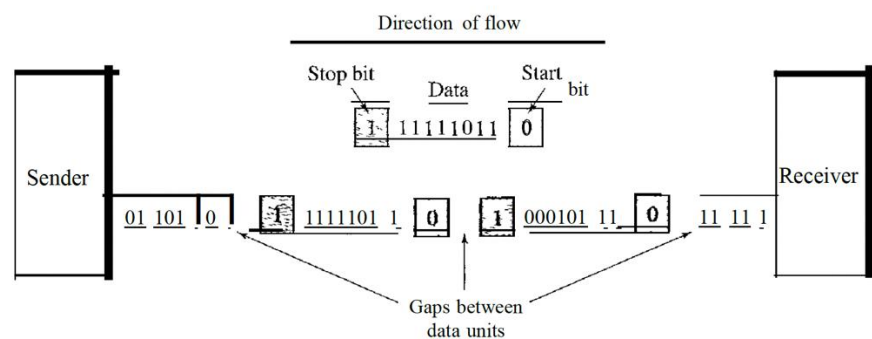


Figure 4.1.3 Asynchronous transmission

4.1.4 Synchronous Transmission

In synchronous transmission, the bit stream is combined into longer "frames," which may contain multiple bytes. Each byte, however, is introduced onto the transmission link without a gap between it and the next one. It is left to the receiver to separate the bit stream into bytes for decoding purposes. In other words, data are transmitted as an unbroken string of 1s

and Os, and the receiver separates that string into the bytes, or characters, it needs to reconstruct the information.

In synchronous transmission, we send bits one after another without start or stop bits or gaps. It is the responsibility of the receiver to group the bits.

Figure 4.1.4 gives a schematic illustration of synchronous transmission. We have drawn in the divisions between bytes. In reality, those divisions do not exist; the sender puts its data onto the line as one long string. If the sender wishes to send data in separate bursts, the gaps between bursts must be filled with a special sequence of Os and Is that means *idle*. The receiver counts the bits as they arrive and groups them in 8-bit units.

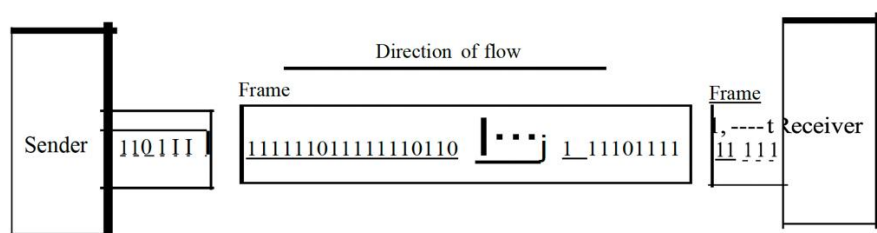


Figure 4.1.4 Synchronous transmission

Without gaps and start and stop bits, there is no built-in mechanism to help the receiving device adjust its bit synchronization midstream. Timing becomes very important, therefore, because the accuracy of the received information is completely dependent on the ability of the receiving device to keep an accurate count of the bits as they come in.

The advantage of synchronous transmission is speed. With no extra bits or gaps to introduce at the sending end and remove at the receiving end, and, by extension, with fewer bits to move across the link, synchronous transmission is faster than asynchronous transmission. For this reason, it is more useful for high-speed applications such as the transmission of data from one computer to another. Byte synchronization is accomplished in the data link layer.

We need to emphasize one point here. Although there is no gap between characters in synchronous serial transmission, there may be uneven gaps between frames.

4.1.5 Isochronous transmissions

In real-time audio and video, in which uneven delays between frames are not acceptable, synchronous transmission fails. For example, TV images are broadcast at the rate of 30 images per second; they must be viewed at the same rate. If each image is sent by using one or more frames, there should be no delays between frames. For this type of application, synchronization between characters is not enough; the entire stream of bits must be synchronized. The isochronous transmission guarantees that the data arrive at a fixed rate.

4.2 ANALOG TRANSMISSION

Converting digital data to a bandpass analog signal is traditionally called digital to-analog conversion. Converting a low-pass analog signal to a bandpass analog signal is traditionally called analog-to-analog conversion. In this chapter, we discuss these two types of conversions.

4.2.1 Digital-to-Analog Conversion

Digital-to-analog conversion is the process of changing one of the characteristics of an analog signal based on the information in digital data. Figure 4.6 shows the relationship between the digital information, the digital-to-analog modulating process, and the resultant analog signal.

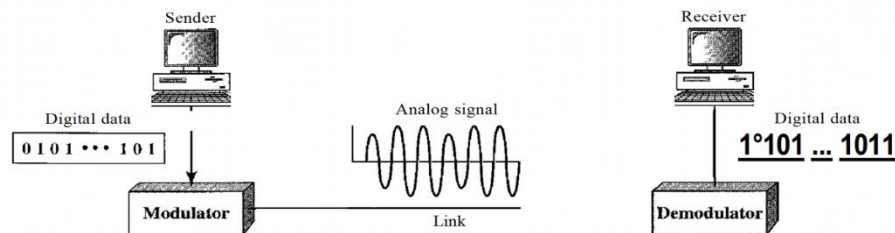
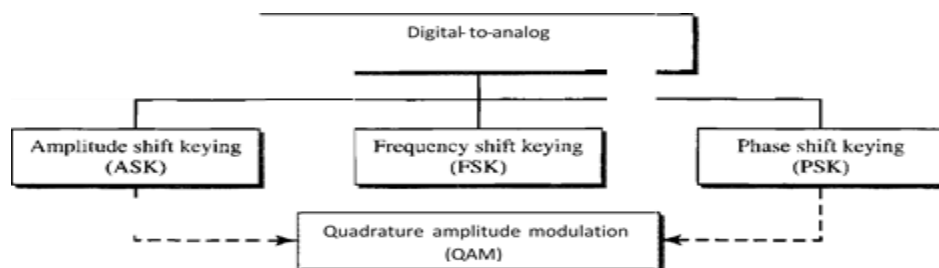


Figure 4.2.1 Digital to Analog transmission

A sine wave is defined by three characteristics: amplitude, frequency, and phase. When we vary anyone of these characteristics, we create a different version of that wave. So, by changing one characteristic of a simple electric signal, we can use it to represent digital data. Any of the three characteristics can be altered in this way, giving us at least three mechanisms for modulating digital data into an analog signal: amplitude shift keying (ASK), frequency shift keying (FSK), and phase shift keying (PSK). In addition, there is a fourth (and better) mechanism that combines changing both the amplitude and phase, called quadrature amplitude modulation (QAM). QAM is the most efficient of these options and is the mechanism commonly used today (see Figure below).



4.2.2 Aspects of Digital-to-Analog Conversion

- **Data Element Versus Signal Element**

We defined a data element as the smallest piece of information to be exchanged, the bit. We also defined a signal element as the smallest unit of a signal that is constant. Although we continue to use the same terms in this chapter, we will see that the nature of the signal element is a little bit different in analog transmission.

- **Data Rate Versus Signal Rate**

We can define the data rate (bit rate) and the signal rate (baud rate) as we did for digital transmission. The relationship between them is

$$S = N \times \frac{1}{R} \text{ baud}$$

where N is the data rate (bps) and r is the number of data elements carried in one signal element. The value of r in analog transmission is $r = \log_2 L$, where L is the type of signal element, not the level. The same nomenclature is used to simplify the comparisons.

The same analogy we used in Chapter 4 for bit rate and baud rate applies here. In transportation, a baud is analogous to a vehicle, and a bit is analogous to a passenger. We need to maximize the number of people per car to reduce the traffic.

- **Bandwidth**

The required bandwidth for analog transmission of digital data is proportional to the signal rate except for FSK, in which the difference between the carrier signals needs to be added. We discuss the bandwidth for each technique.

- **Carrier Signal**

In analog transmission, the sending device produces a high-frequency signal that acts as a base for the information signal. This base signal is called the carrier signal or carrier frequency.

The receiving device is tuned to the frequency of the carrier signal that it expects from the sender. Digital information then changes the carrier signal by modifying one or more of its characteristics (amplitude, frequency, or phase). This kind of modification is called modulation (shift keying).

4.2.3 Amplitude Shift Keying (ASK)

In amplitude shift keying, the amplitude of the carrier signal is varied to create signal elements. Both frequency and phase remain constant while the amplitude changes.

4.2.4 Binary ASK (BASK)

Although we can have several levels (kinds) of signal elements, each with a different amplitude, ASK is normally implemented using only two levels. This is referred to as binary amplitude shift keying or *on-off keying* (OOK). The peak amplitude of one signal level is 0; the other is the same as the amplitude of the carrier frequency. Figure 4.3 gives a conceptual view of binary ASK.

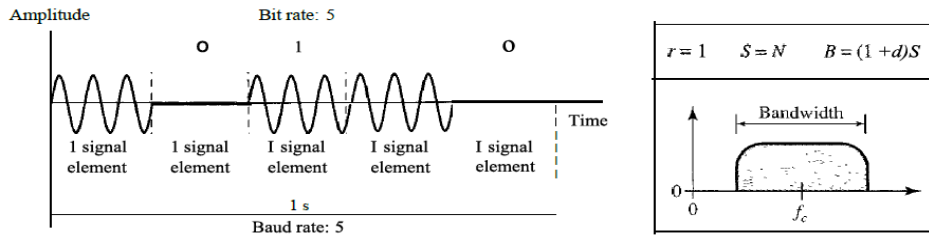


Figure 4.2.4 Binary amplitude shift keying

Bandwidth for ASK Figure 4.2.4 also shows the bandwidth for ASK. Although the carrier signal is only one simple sine wave, the process of modulation produces a nonperiodic composite signal. This signal, has a continuous set of frequencies. As we expect, the bandwidth is proportional to the signal rate (baud rate). However, there is normally another factor involved, called d , which depends on the modulation and filtering process. The value of d is between 0 and 1. This means that the bandwidth can be expressed as shown, where S is the signal rate and the B is the bandwidth.

$$B = (1 + d) \times S$$

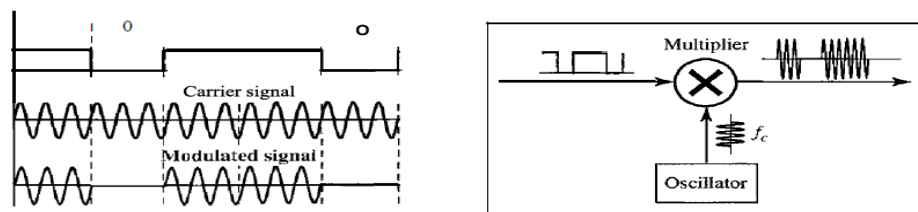
The formula shows that the required bandwidth has a minimum value of S and a maximum value of $2S$. The most important point here is the location of the bandwidth.

The middle of the bandwidth is where the carrier frequency, f_c , is located. This means if we have a bandpass channel available; we can choose our f_c so that the modulated signal occupies that bandwidth. This is in fact the most important advantage of digital-to-analog conversion. We can shift the resulting bandwidth to match what is available.

Implementation The complete discussion of ASK implementation is beyond the scope of this book. However, the simple ideas behind the implementation may help us to better understand the concept itself. Figure 4.4 shows how we can simply implement binary ASK.

If digital data are presented as a unipolar NRZ (see Chapter 4) digital signal with a high voltage of 1 V and a low voltage of 0 V, the implementation can be achieved by multiplying the NRZ digital signal by the carrier signal coming from an oscillator. When the amplitude of the NRZ signal is 1, the amplitude of the carrier frequency is

Implementation of binary ASK



held; when the amplitude of the NRZ signal is 0, the amplitude of the carrier frequency is zero.

The above discussion uses only two amplitude levels. We can have multilevel ASK in which there are more than two levels. We can use 4, 8, 16, or more different amplitudes for the signal and modulate the data using 2, 3, 4, or more bits at a time. In these cases,

$r = 2, r = 3, r = 4$, and so on. Although this is not implemented with pure ASK, it is implemented with QAM (as we will see later).

4.3 FREQUENCY SHIFT KEYING (FSK)

In frequency shift keying, the frequency of the carrier signal is varied to represent data. The frequency of the modulated signal is constant for the duration of one signal element, but changes for the next signal element if the data element changes. Both peak amplitude and phase remain constant for all signal elements.

4.3.1 Binary FSK (BFSK)

One way to think about binary FSK (or BFSK) is to consider two carrier frequencies. In figure 4.3.1 We have selected two carrier frequencies f_1 and f_2 . We use the first carrier if the data element is 0; we use the second if the data element is 1. However, note that this is an unrealistic example used only for demonstration purposes. Normally the carrier frequencies are very high, and the difference between them is very small.

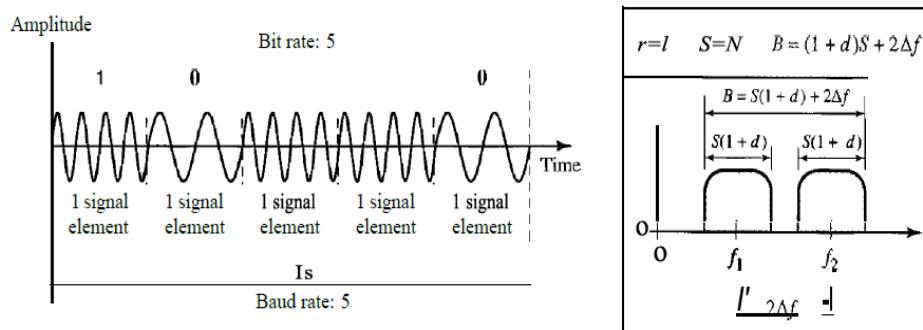


Figure 4.3.1: Binary frequency shift keying

As Figure 4.3.1 shows, the middle of one bandwidth is f_1 and the middle of the other is f_2 . both f_1 and f_2 are Δf apart from the midpoint between the two bands. The difference between the two frequencies is $2\Delta f$.

Bandwidth for BFSK also shows the bandwidth of FSK. Again the carrier signals are only simple sine waves, but the modulation creates a nonperiodic composite signal with continuous frequencies. We can think of FSK as two ASK signals, each with its own carrier frequency C_{il} or f_2). If the difference between the two frequencies is $2\Delta f$, then the required bandwidth is $B = (1 + d) \times S + 2\Delta f$

What should be the minimum value of $2\Delta f$? In Figure 4.6, we have chosen a value greater than $(1 + d)S$. It can be shown that the minimum value should be at least S for the proper operation of modulation and demodulation.

Implementation: There are two implementations of BFSK: non coherent and coherent. In noncoherent BFSK, there may be discontinuity in the phase when one signal element ends and the next begins. In coherent BFSK, the phase continues through the boundary of two signal elements. Noncoherent BFSK can be implemented by treating BFSK as two ASK modulations and using two carrier frequencies. Coherent BFSK can be implemented by using one *voltage-controlled oscillator* (VeO) that changes its frequency according to the input voltage.

4.3.2 Multilevel FSK

Multilevel modulation (MFSK) is not uncommon with the FSK method. We can use more than two frequencies. For example, we can use four different frequencies f_1, f_2, f_3 , and 14 to send 2 bits at a time. To send 3 bits at a time, we can use eight frequencies. And so on. However, we need to remember that the frequencies need to be $2\Delta f$ apart.

For the proper operation of the modulator and demodulator, it can be shown that the minimum value of $2\Delta f$ needs to be S . We can show that the bandwidth with $d = 0$ is

$$B = (1 + d) \times S + (L - 1)2\Delta f \rightarrow B = L \times S$$

4.4 PHASE SHIFT KEYING (PSK)

In phase shift keying, the phase of the carrier is varied to represent two or more different signal elements. Both peak amplitude and frequency remain constant as the phase changes.

Today, PSK is more common than ASK or FSK. However, we will see that QAM, which combines ASK and PSK, is the dominant method of digital to-analog modulation.

4.4.1 Binary PSK (BPSK)

The simplest PSK is binary PSK, in which we have only two signal elements, one with a phase of 0° , and the other with a phase of 180° . Figure 4.14 gives a conceptual view of PSK.

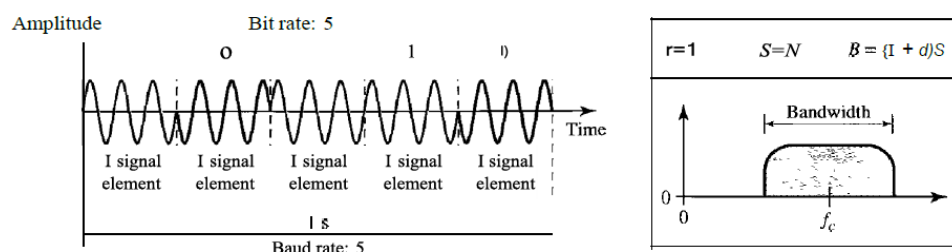


Figure 4.4.1: Binary phase shift keying

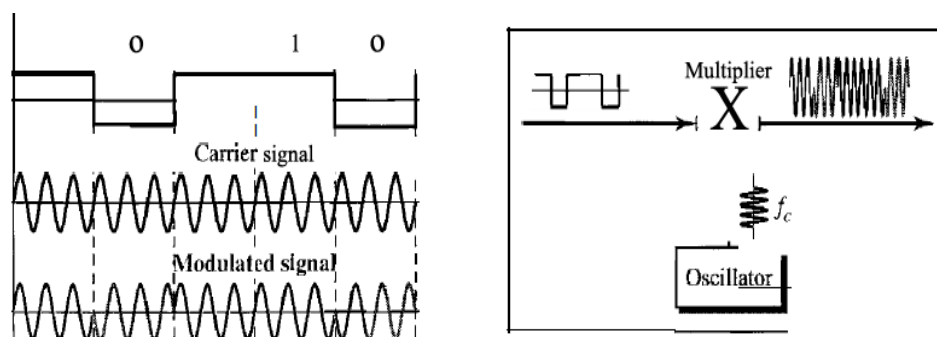
Binary PSK is as simple as binary ASK with one big advantage—it is less susceptible to noise. In ASK, the criterion for bit detection is the amplitude of the signal. In PSK, it is the phase. Noise can change the amplitude easier than it can change the phase. In other words, PSK is less susceptible to noise than ASK. PSK is superior to FSK because we do not need two carrier signals.

Bandwidth Figure 4.4.1 also shows the bandwidth for BPSK. The bandwidth is the same as that for binary ASK, but less than that for BFSK. No bandwidth is wasted for separating two carrier signals.

Implementation: The implementation of BPSK is as simple as that for ASK. The reason is that the signal element with phase 180° can be seen as the complement of the signal element with phase 0° . This gives us a clue on how to implement BPSK. We use the same idea we used for ASK but with a polar NRZ signal instead of a unipolar NRZ signal.

The polar NRZ signal is multiplied by the carrier frequency; the 1 bit (positive voltage) is represented by a phase starting at 0° ; the 0 bit (negative voltage) is represented by a phase starting at 180° .

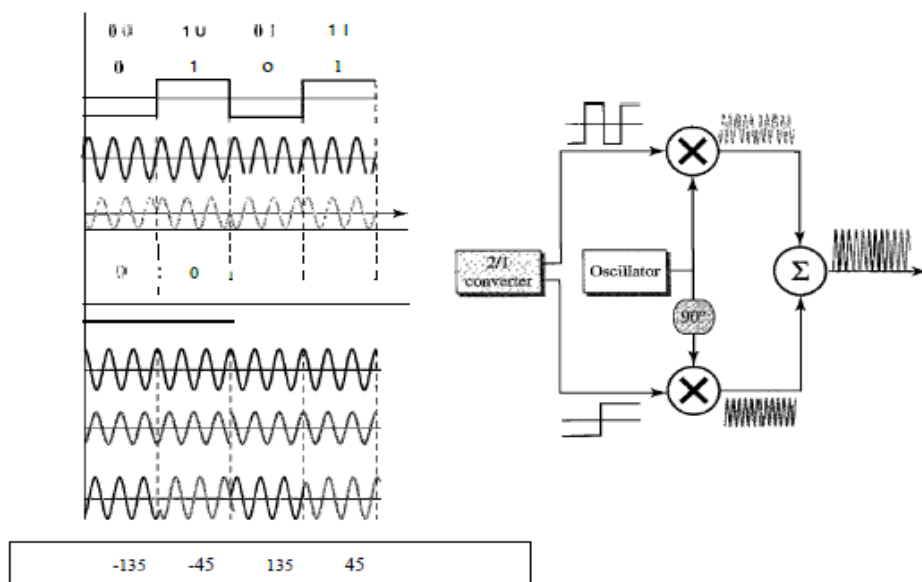
Implementation of BASK



4.4.2 Quadrature PSK (QPSK)

The simplicity of BPSK enticed designers to use 2 bits at a time in each signal element, thereby decreasing the baud rate and eventually the required bandwidth. The scheme is called quadrature PSK or QPSK because it uses two separate BPSK modulations; one is in-phase, the other quadrature (out-of-phase). The incoming bits are first passed through a serial-to-parallel conversion that sends one bit to one modulator and the next bit to the other modulator. If the duration of each bit in the incoming signal is T , the duration of each bit sent to the corresponding BPSK signal is $2T$. This means that the bit to each BPSK signal has one-half the frequency of the original signal. Figure 4.4.2 shows the idea.

The two composite signals created by each multiplier are sine waves with the same frequency, but different phases. When they are added, the result is another sine wave, with one of four possible phases: 45° , -45° , 135° , and -135° . There are four kinds of signal elements in the output signal ($L = 4$), so we can send 2 bits per signal element ($r = 2$).



4.5 CONSTELLATION DIAGRAM

A **constellation diagram** can help us define the amplitude and phase of a signal element, particularly when we are using two carriers (one in-phase and one quadrature). The diagram is useful when we are dealing with multilevel ASK, PSK, or QAM (see next section). In a constellation diagram, a signal element type is represented as a dot. The bit or combination of bits it can carry is often written next to it.

The diagram has two axes. The horizontal X axis is related to the in-phase carrier; the vertical Y axis is related to the quadrature carrier. For each point on the diagram, four pieces of information can be deduced. The projection of the point on the X axis defines the peak amplitude of the in-phase component; the projection of the point on the Y axis defines the peak amplitude of the quadrature component. The length of the line (vector) that connects the point to the origin is the peak amplitude of the signal element (combination of the X and Y components); the angle the line makes with the X axis is the phase of the signal element. All the information we need, can easily be found on a constellation diagram.

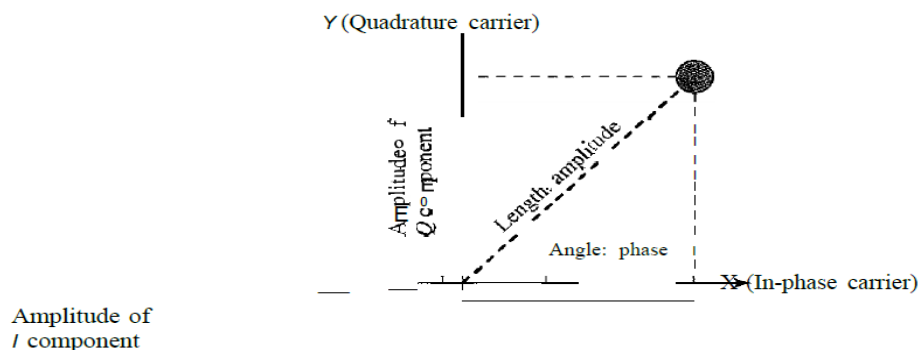


Figure 4.5.1 Concept of a constellation diagram

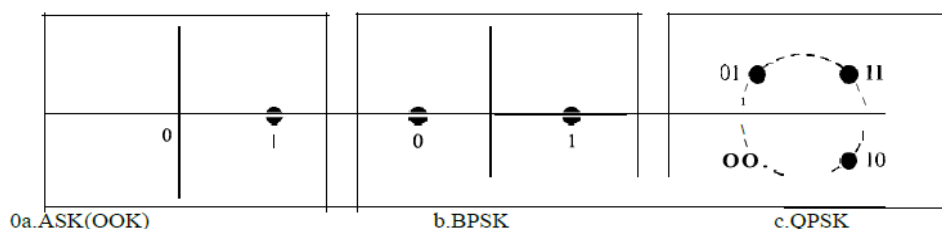
Example:

Show the constellation diagrams for an ASK (OOK), BPSK, and QPSK signals.

Solution

Figure below shows the three constellation diagrams.

Three constellation diagrams



Let us analyze each case separately:

- For ASK, we are using only an in-phase carrier. Therefore, the two points should be on the X axis. Binary 0 has an amplitude of 0 V; binary 1 has an amplitude of 1 V (for example). The points are located at the origin and at 1 unit.
- BPSK also uses only an in-phase carrier. However, we use a polar NRZ signal for modulation. It creates two types of signal elements, one with amplitude 1 and the other with amplitude -1. This can be stated in other words: BPSK creates two different signal elements, one with amplitude 1 V and in phase and the other with amplitude 1 V and out of phase.
- QPSK uses two carriers, one in-phase and the other quadrature. The point representing 11 is made of two combined signal elements, both with an amplitude of 1 V. One element is represented by an in-phase carrier, the other element by a quadrature carrier. The amplitude of the final signal element sent for this 2-bit data element is $2^{1/2}$, and the phase is 45° . The argument is similar for the other three points.

All signal elements have an amplitude of $2^{1/2}$, but their phases are different (45° , 135° , -135° , and -45°). Of course, we could have chosen the amplitude of the carrier to be $1/(2^{1/2})$ to make the final amplitudes 1 V.

4.6 QUADRATURE AMPLITUDE MODULATION

PSK is limited by the ability of the equipment to distinguish small differences in phase. This factor limits its potential bit rate. So far, we have been altering only one of the three characteristics of a sine wave at a time; but what if we alter two? Why not combine ASK and PSK? The idea of using two carriers, one in-phase and the other quadrature, with different amplitude levels for each carrier is the concept behind quadrature amplitude modulation (QAM).

Quadrature amplitude modulation is a combination of ASK and PSK.

The possible variations of QAM are numerous. Figure 4.6 shows some of these schemes. Figure 4.6a shows the simplest 4-QAM scheme (four different signal element types) using a unipolar NRZ signal to modulate each carrier. This is the same mechanism we used for ASK (OOK). Part b shows another 4-QAM using polar NRZ, but this is exactly the same as QPSK. Part c shows another QAM-4 in which we used a signal with two positive levels to modulate each of the two carriers. Finally, Figure 4.6d shows a 16-QAM constellation of a signal with eight levels, four positive and four negative.

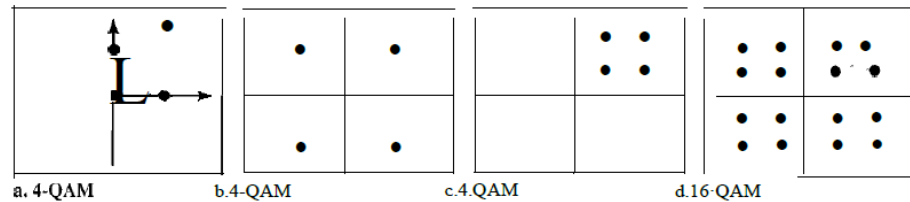


Figure 4.6 Constellation diagrams for some QAMs

Bandwidth for QAM

The minimum bandwidth required for QAM transmission is the same as that required for ASK and PSK transmission. QAM has the same advantages as PSK over ASK.

4.7 REVIEW QUESTION

1. Define analog transmission.
2. Define carrier signal and its role in analog transmission.
3. Define digital-to-analog conversion.
4. Which characteristics of an analog signal are changed to represent the digital signal in each of the following digital-to-analog conversion?
 - a. ASK
 - b. FSK
 - c. PSK
 - d. QAM
5. Which of the four digital-to-analog conversion techniques (ASK, FSK, PSK or QAM) is the most susceptible to noise? Defend your answer.
6. Define constellation diagram and its role in analog transmission.



ANALOG TO ANALOG CONVERSION AND MULTIPLEXING

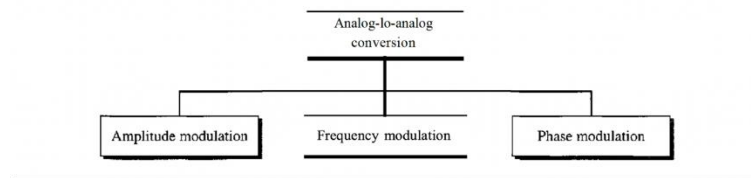
Unit Structure:

- 5.0 Introduction
- 5.1 Analog-to-Analog conversion
 - 5.1.1 Amplitude Modulation
 - 5.1.2 AM Bandwidth
 - 5.1.3 Standard bandwidth Allocation for AM Radio
 - 5.1.4 Frequency Modulation
 - 5.1.5 FM Bandwidth
 - 5.1.6 Standard Bandwidth allocation for FM Radio
 - 5.1.7 Phase Modulation
 - 5.1.7 PM Bandwidth
- 5.2 Bandwidth Utilization: Multiplexing and Spreading
 - 5.2.1 Multiplexing
 - 5.2.2 Frequency-Division Multiplexing
 - 5.2.3 Multiplexing Process
 - 5.2.4 Demultiplexing Process
 - 5.2.5 Analog Carrier System
 - 5.2.6 Wavelength-Division Multiplexing
 - 5.2.7 Synchronous Time-Division Multiplexing
 - 5.2.8 Time Slots and Frames
 - 5.2.9 Interleaving
 - 5.2.10 Empty Slots
 - 5.2.11 Statistical Time-Division Multiplexing
- 5.3 Spread Spectrum
 - 5.3.1 Frequency Hopping Spread Spectrum (FHSS)
 - 5.3.2 Bandwidth Sharing
 - 5.3.3 Direct Sequence Spread Spectrum (DSSS)
- 5.4 Review Questions

5.0 INTRODUCTION

Analog-to-analog conversion, or analog modulation, is the representation of analog information by an analog signal. One may ask why we need to modulate an analog signal; it is already analog. Modulation is needed if the medium is bandpass in nature or if only a bandpass channel is available to us. An example is radio. The government assigns a narrow bandwidth to each radio station. The analog signal produced by each station is a low-pass signal, all in the same range. To be able to listen to different stations, the low-pass signals need to be shifted, each to a different range.

Analog-to-analog conversion can be accomplished in three ways: amplitude modulation (AM), frequency modulation (FM), and phase modulation (PM). FM and PM are usually categorized together. See Figure below



5.1 ANALOG-TO-ANALOG CONVERSION

5.1.1 Amplitude Modulation

In AM transmission, the carrier signal is modulated so that its amplitude varies with the changing amplitudes of the modulating signal. The frequency and phase of the carrier remain the same; only the amplitude changes to follow variations in the information. Below figure shows how this concept works. The modulating signal is the envelope of the carrier.

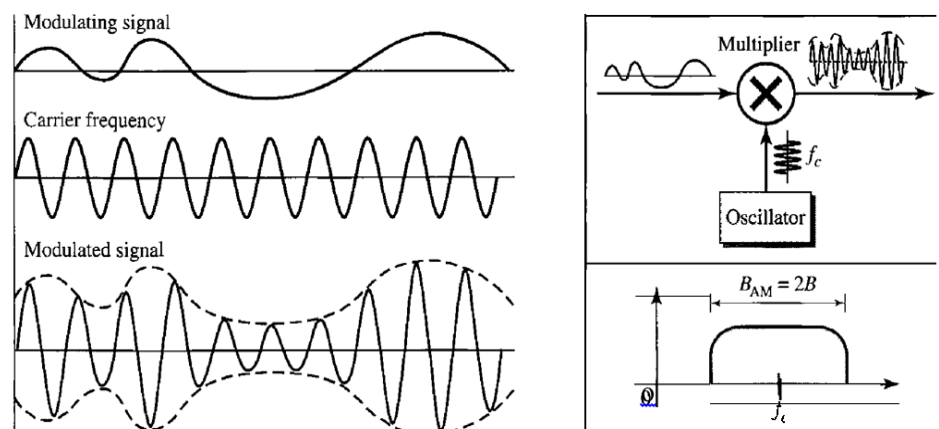


Figure 5.1.1 Amplitude modulation

As Figure 5.1.1 shows, AM is normally implemented by using a simple multiplier because the amplitude of the carrier signal needs to be changed according to the amplitude of the modulating signal.

5.1.2 AM Bandwidth

Figure 5.1.2 also shows the bandwidth of an AM signal. The modulation creates a bandwidth that is twice the bandwidth of the modulating signal and covers a range centered on the carrier frequency. However, the signal components above and below the carrier frequency carry the same information. For this reason, some implementations discard one-half of the signals and cut the bandwidth in half.

The total bandwidth required for AM can be determined from the bandwidth of the audio signal: $B_{AM} = 2B$.

5.1.3 Standard bandwidth Allocation for AM Radio

The bandwidth of an audio signal (speech and music) is usually 5 kHz. Therefore, an AM radio station needs a bandwidth of 10kHz. In fact, the Federal Communications Commission (FCC) allows 10 kHz for each AM station.

AM stations are allowed carrier frequencies anywhere between 530 and 1700 kHz (1.7 MHz). However, each station's carrier frequency must be separated from those on either side of it by at least 10 kHz (one AM bandwidth) to avoid interference. If one station uses a carrier frequency of 1100 kHz, the next station's carrier frequency cannot be lower than 1110 kHz (see Figure 5.1.2).

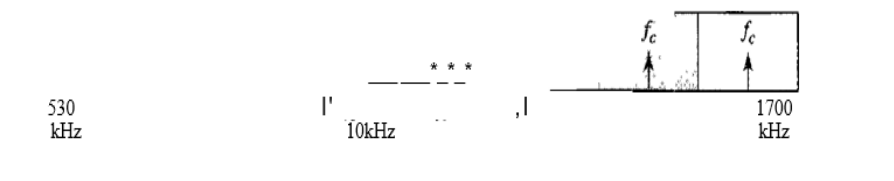


Figure 5.1.2 AM band allocation

5.1.4 Frequency Modulation

In FM transmission, the frequency of the carrier signal is modulated to follow the changing voltage level (amplitude) of the modulating signal. The peak amplitude and phase of the carrier signal remain constant, but as the amplitude of the information signal changes, the frequency of the carrier changes correspondingly.

FM is implemented by using a voltage-controlled oscillator as with FSK. The frequency of the oscillator changes according to the input voltage which is the amplitude of the modulating signal.

5.1.5 FM Bandwidth

Figure 5.1.3 also shows the bandwidth of an FM signal. The actual bandwidth is difficult to determine exactly, but it can be shown empirically that it is several times that of the analog signal or $2(1 + \beta)B$ where β is a factor depends on modulation technique with a common value of 4.

5.1.6 Standard bandwidth Allocation for FM Radio

The bandwidth of an audio signal (speech and music) broadcast in stereo is almost 15 kHz. The FCC allows 200 kHz (0.2 MHz) for each station. This means $\beta = 4$ with some extra guard band. FM stations are allowed carrier frequencies anywhere between

88 and 108 MHz. Stations must be separated by at least 200 kHz to keep their bandwidths from overlapping. To create even more privacy, the FCC requires that in each area, only alternate bandwidth allocations may be used. The others remain unused to prevent any possibility of two stations interfering with each other. Given 88 to 108 MHz as a range, there are 100 potential PM bandwidths in an area, of which 50 can operate at any one time.

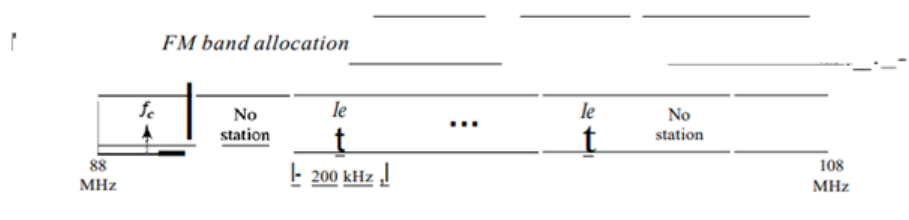


Figure 5.1.3

5.1.7 Phase Modulation

In PM transmission, the phase of the carrier signal is modulated to follow the changing voltage level (amplitude) of the modulating signal. The peak amplitude and frequency of the carrier signal remain constant, but as the amplitude of the information signal changes, the phase of the carrier changes correspondingly. It can be proved mathematically (see Appendix C) that PM is the same as FM with one difference.

In FM, the instantaneous change in the carrier frequency is proportional to the amplitude of the modulating signal. In PM the instantaneous change in the carrier frequency is proportional to the derivative of the amplitude of the modulating signal. Below figure shows the relationships of the modulating signal, the carrier signal, and the resultant PM signal.

PM is normally implemented by using a voltage-controlled oscillator along with a derivative. The frequency of the oscillator changes according to the derivative of the input voltage which is the amplitude of the modulating signal.

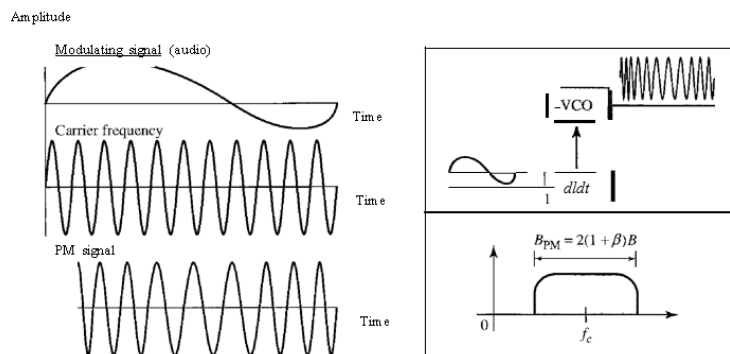


Figure 5.1.4 Phase Modulations

5.1.8 PM Bandwidth

Figure 5.1.4 also shows the bandwidth of a PM signal. The actual bandwidth is difficult to determine exactly, but it can be shown empirically that it is several times that of the analog signal. Although, the formula shows the same bandwidth for FM and PM, the value of β is lower in the case of PM (around 1 for narrowband and 3 for wideband).

5.2 BANDWIDTH UTILIZATION: MULTIPLEXING AND SPREADING

In real life, we have links with limited bandwidths. The wise use of these bandwidths has been, and will be, one of the main challenges of electronic communications. However, the meaning of *wise* may depend on the application. Sometimes we need to combine several low-bandwidth channels to make use of one channel with a larger bandwidth. Sometimes we need to expand the bandwidth of a channel to achieve goals such as privacy and antijamming. In this chapter, we explore these two broad categories of bandwidth utilization: multiplexing and spreading. In multiplexing, our goal is efficiency; we combine several channels into one. In spreading, our goals are privacy and antijamming; we expand the bandwidth of a channel to insert redundancy, which is necessary to achieve these goals.

5.2.1 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. We can accommodate this increase by continuing to add individual links each time a new channel is needed; or we can install higher-bandwidth links and use each to carry multiple signals. Today's technology includes high-bandwidth media such as optical fiber and terrestrial and satellite microwaves. Each has a bandwidth far more than that needed for the average transmission

signal. If the bandwidth of a link is greater than the bandwidth needs of the devices connected to it, the bandwidth is wasted. An efficient system maximizes the utilization of all resources; bandwidth is one of the most precious resources we have in data communications.

In a multiplexed system, n lines share the bandwidth of one link. Figure 5.2.1 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.

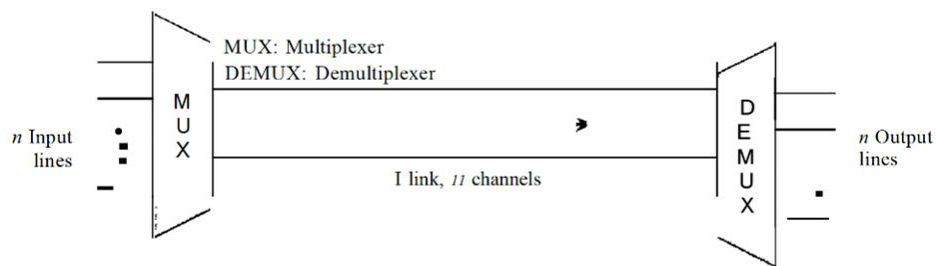
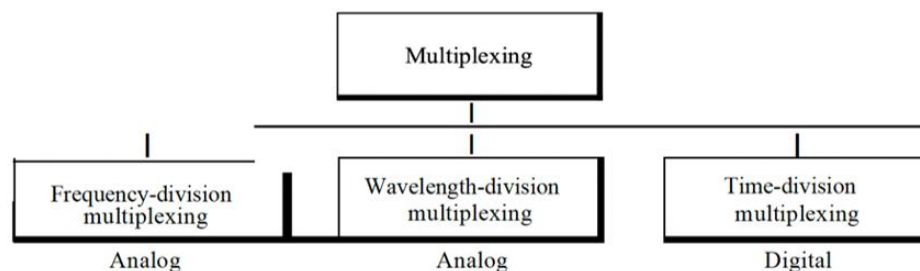


Figure 5.2.1 dividing a link into channels

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 (see below figure).



Categories of multiplexing

5.2.2 Frequency-Division Multiplexing

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 below 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.

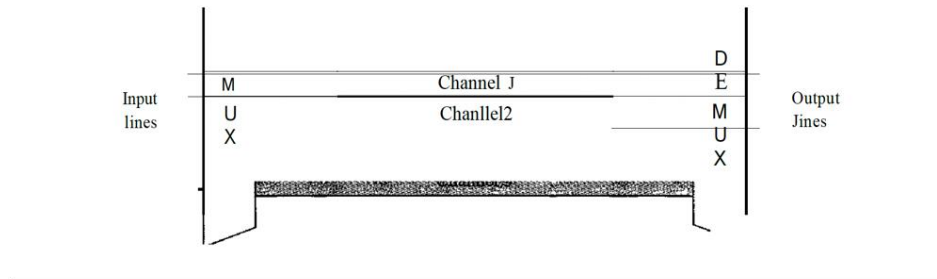


Figure 5.2.2 frequency-division multiplexing

We consider FDM to be an analog multiplexing technique; however, this does not mean that FDM cannot be used to combine sources sending digital signals. A digital signal can be converted to an analog signal before FDM is used to multiplex them.

5.2.3 Multiplexing Process

Figure 5.2.3 is a conceptual illustration of the multiplexing process. Each source generates a signal of a similar frequency range. Inside the multiplexer, these similar signals modulate different carrier frequencies (f_1 , f_2 , and f_3). 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.

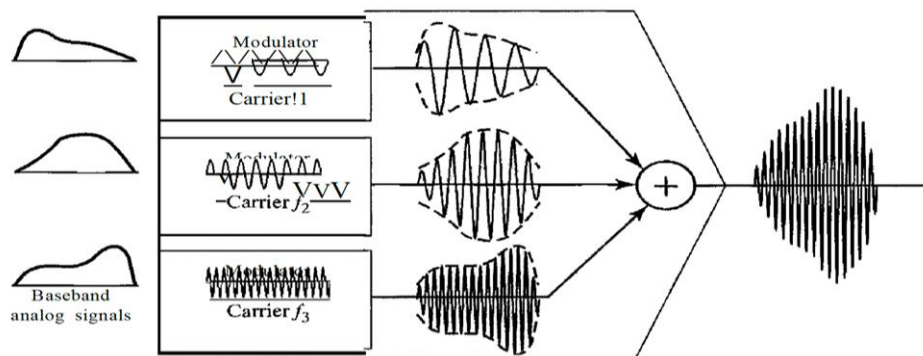


Figure 5.2.3 FDM Process

5.2.4 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. Figure 5.2.4 is a conceptual illustration of demultiplexing process.

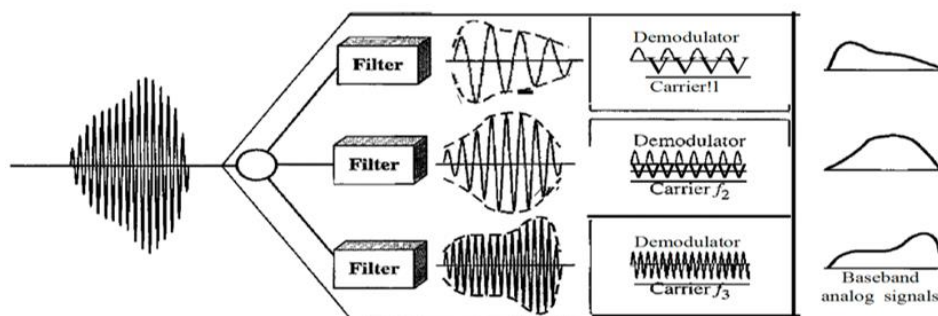
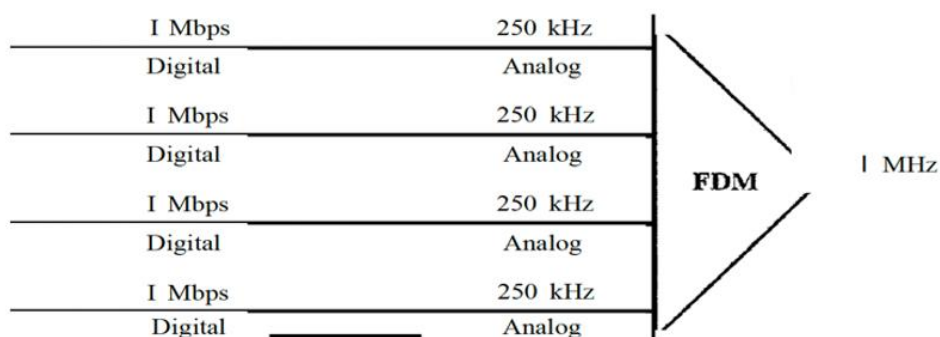


Figure 5.2.4 FDM Demultiplexing

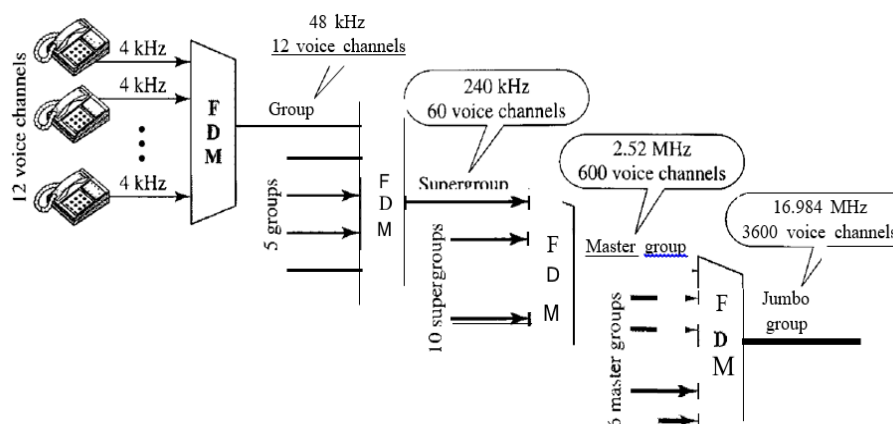
5.2.5 The Analog Carrier System

To maximize the efficiency of their infrastructure, telephone companies have traditionally multiplexed signals from lower-bandwidth lines onto higher-bandwidth lines. In this way, many switched or leased lines can be combined into fewer but bigger channels. For analog lines, FDM is used.



One of these hierarchical systems used by AT&T is made up of groups, supergroups, master groups, and jumbo groups.

Figure 5.2.5 Analog hierarchy



In this analog hierarchy, 12 voice channels are multiplexed onto a higher-bandwidth line to create a group. A group has 48 kHz of bandwidth and supports 12 voice channels.

At the next level, up to five groups can be multiplexed to create a composite signal called a supergroup. A supergroup has a bandwidth of 240 kHz and supports up to 60 voice channels. Supergroups can be made up of either five groups or 60 independent voice channels.

At the next level, 10 supergroups are multiplexed to create a master group. A master group must have 2.40 MHz of bandwidth, but the need for guard bands between the supergroups increases the necessary bandwidth to 2.52 MHz. Master groups support up to 600 voice channels.

Finally, six master groups can be combined into a jumbo group. A jumbo group must have 15.12 MHz (6×2.52 MHz) but is augmented to 16.984 MHz to allow for guard bands between the master groups.

Other Applications of FDM

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. As discussed in Chapter 5, each AM station needs 10 kHz 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. However, we need to know that there is physical multiplexer or demultiplexer here. As we will see in Chapter 12 multiplexing is done at the data link layer.

The situation is similar in FM broadcasting. However, FM has a wider band of 88 to 108 MHz 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.

The first generation of cellular telephones (still in operation) also uses FDM. Each user is assigned two 30-kHz channels, one for sending voice and the other for receiving. The voice signal, which has a bandwidth of 3 kHz (from 300 to 3300 Hz), is modulated by using FM. Remember that an FM signal has a bandwidth 10 times that of the modulating signal, which means each channel has 30 kHz (10×3) of bandwidth. Therefore, each user is given, by the base station, a 60-kHz bandwidth in a range available at the time of the call.

Implementation

FDM can be implemented very easily. In many cases, such as radio and television broadcasting, there is no need for a physical multiplexer or demultiplexer. If the stations agree to send their broadcasts to the air using different carrier frequencies, multiplexing is achieved. In other cases, such as the cellular telephone system, a base station needs to assign a carrier frequency to the telephone user. There is not enough bandwidth in a cell to permanently assign a bandwidth range to every

telephone user. When a user hangs up, her or his bandwidth is assigned to another caller.

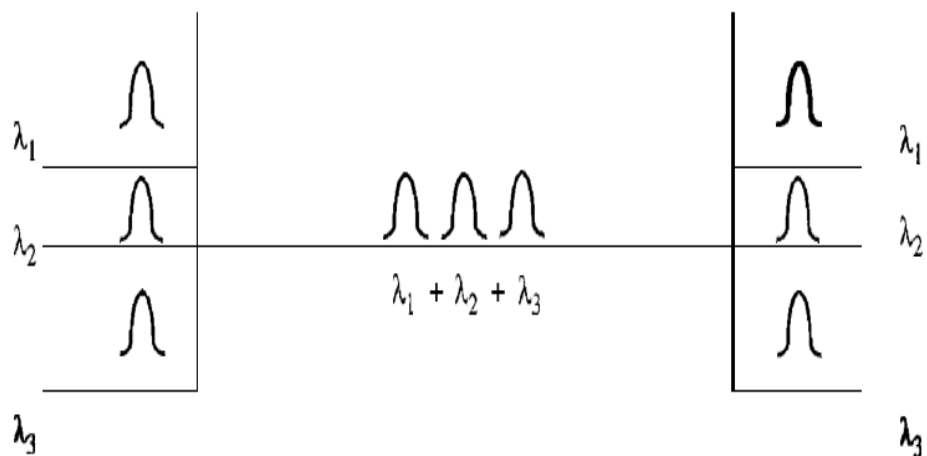
5.2.6 Wavelength-Division Multiplexing

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 one.

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 5.2.6 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.

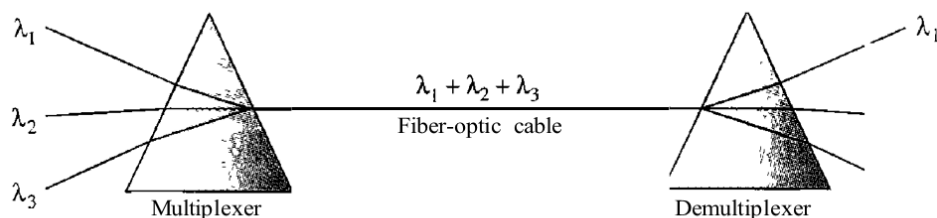
Figure 5.2.6 Wavelength-division multiplexing



WDM is an analog multiplexing technique to combine optical signals.

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. Figure below shows the concept.

Prisms in wavelength-division multiplexing and demultiplexing



One application of WDM is the SONET network in which multiple optical fiber lines are multiplexed and demultiplexed.

A new method, called dense WDM (DWDM), can multiplex a very large number of channels by spacing channels very close to one another. It achieves even greater efficiency.

5.2.7 Synchronous Time-Division Multiplexing

Time division multiplexing (TDM) is a digital process that allows several connections to share the high bandwidth of a link. Instead of sharing a portion of the bandwidth as in shared. Each connection occupies a portion of time in the link. Figure 5.17 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.

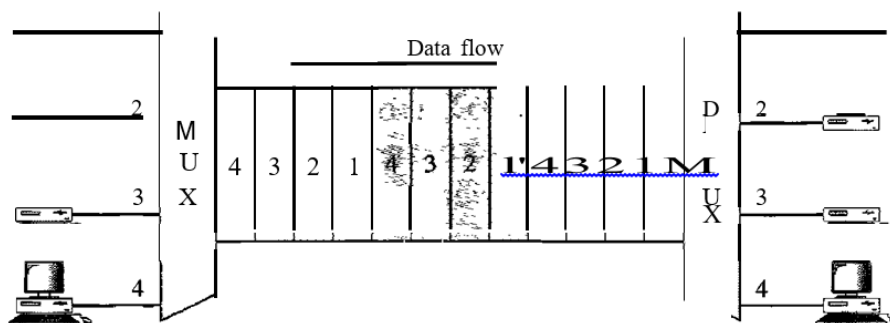


Figure 5.2.7 TDM

Note that in Figure 6.2.7 we are concerned with only multiplexing, not switching. This means that all the data in a message from source 1 always go to one specific destination, be it 1, 2, 3, or 4. The delivery is fixed and unvarying, unlike switching.

We also need to remember that TDM is, in principle, a digital multiplexing technique. 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.

TDM is a digital multiplexing technique for combining several low-rate channels into one high-rate one.

We can divide TDM into two different schemes: synchronous and statistical. We first discuss synchronous TDM and then show how statistical TDM differs. In synchronous TDM, each input connection has an allotment in the output even if it is not sending data.

5.2.8 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 5.18 shows an example of synchronous TDM where n is 6.

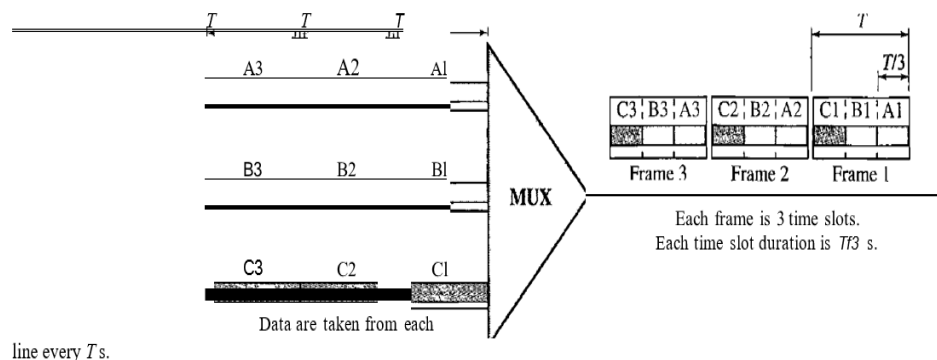


Figure 5.2.8 Synchronous time-division multiplexing

In synchronous TDM, a round of data units from each input connection is collected into a frame (we will see the reason for this shortly). 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 (unless a frame carries some other information, as we will see shortly).

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.

In synchronous TDM, the data rate of the link is n times faster, and the unit duration is n times shorter.

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.

5.2.9 Interleaving

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

Figure 5.2.9 shows the interleaving process for the connection shown in Figure 5.2.9. In this figure, we assume that no switching is involved and that the data from the first connection at the multiplexer site go to the first connection at the demultiplexer.

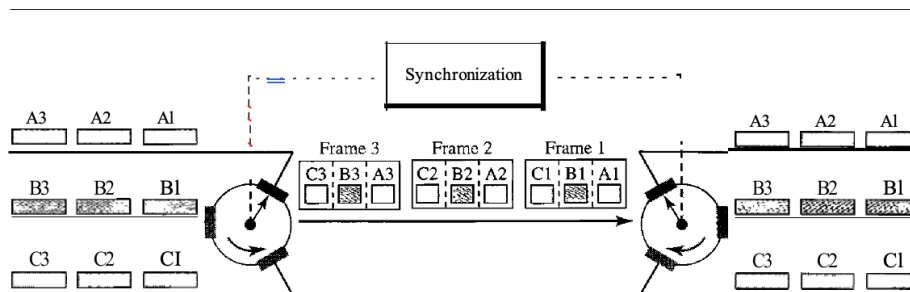


Figure 5.2.9 Interleaving

5.2.10 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. Figure 5.2.10 shows a case in which one of the input lines has no data to send and one slot in another input line has discontinuous data.

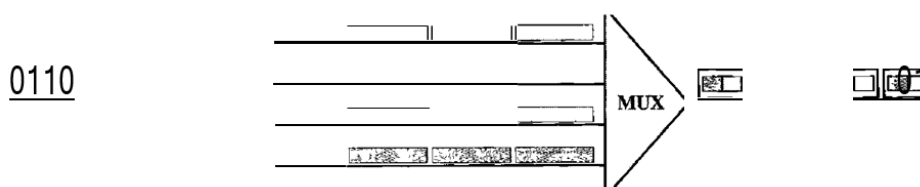


Figure 5.2.10 Empty slots

The first output frame has three slots filled, the second frame has two slots filled, and the third frame has three slots filled. No frame is full. We learn in the next section that 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. In all our discussion so far, we assumed that the data rates of all input lines were the same. However, if data rates are not the same, three strategies, or a combination of them, can be used. We call these three

strategies **multilevel multiplexing, multiple-slot allocation, and pulse stuffing.**

5.2.11 Statistical Time-Division Multiplexing

As we saw in the previous section, 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 5.2.11 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 if there are data to be sent by any input line.

Addressing

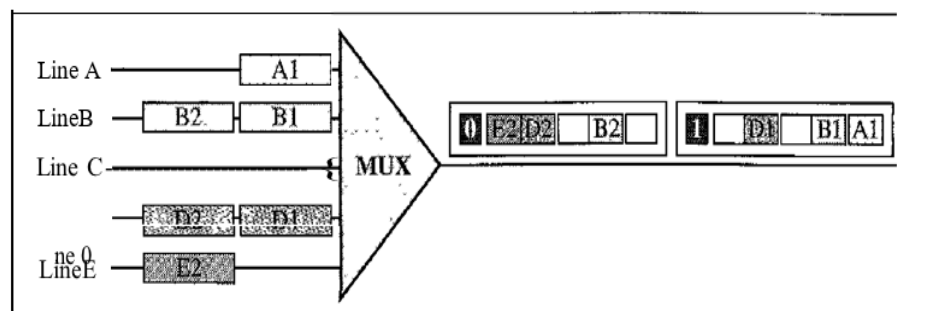
Figure 5.2.11 also shows a major difference between slots in synchronous TDM and statistical TDM. An output slot in synchronous TDM is totally occupied by data; in statistical TDM, a slot needs to carry data as well as the address of the destination. In synchronous TDM, there is no need for addressing; synchronization and preassigned relationships between the inputs and outputs serve as an address.

We know, for example, that input 1 always goes to output 2. If the multiplexer and the demultiplexer are synchronized, this is guaranteed. In statistical multiplexing, there is no fixed relationship between the inputs and outputs because there are no preassigned or reserved slots.

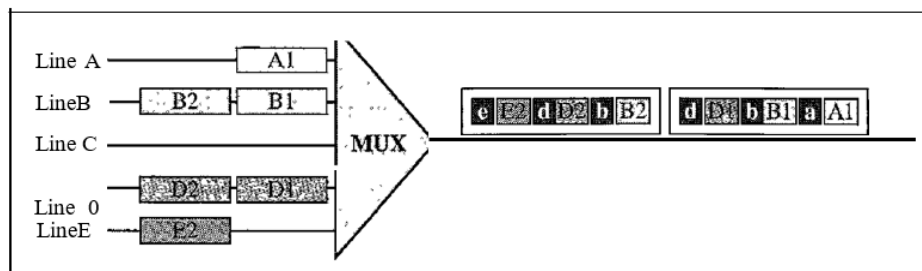
We need to include the address of the receiver inside each slot to show where it is to be delivered. The addressing in its simplest form can be n bits to define N different output lines with $n = \log_2 N$. For example, for eight different output lines, we need a 3-bit address.

Figure 5.2.11 TDM slot comparison

a. Synchronous TDM



b. Statistical TDM

**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

There is another difference between synchronous and statistical TDM, but this time it is at the frame level. The frames in statistical TDM need not be synchronized, so we do not need synchronization bits.

Bandwidth

In statistical TDM, the capacity of the link is normally less than the sum of the capacities of each channel. The designers of statistical TDM define the capacity of the link based on the statistics of the load for each channel. If on average only x percent of the input slots are filled, the capacity of the link reflects this. Of course, during peak times, some slots need to wait.

5.3 SPREAD SPECTRUM

Multiplexing combines signals from several sources to achieve bandwidth efficiency; the available bandwidth of a link is divided between the sources. In spread spectrum we also combine signals from different sources to fit into a larger bandwidth, but our goals are somewhat different. Spread spectrum is designed to use in wireless applications (LANs and WANs). In these types of applications, we have some concerns that outweigh bandwidth efficiency. In wireless applications, all stations use air (or a vacuum) as the medium for communication. Stations must be able to share this medium without interception by an eavesdropper and without being subject to jamming from a malicious intruder (in military operations, for example).

To achieve these goals, spread spectrum techniques add redundancy; they spread the original spectrum needed for each station. If the required bandwidth for each station is B , spread spectrum expands it to B_{SS} such

that $B_{SS} \gg B$. The expanded bandwidth allows the source to wrap its message in a protective envelope for a more secure transmission. An analogy is the sending of a delicate, expensive gift. We can insert the gift in a special box to prevent it from being damaged during transportation, and we can use a superior delivery service to guarantee the safety of the package.

Figure 5.3 shows the idea of spread spectrum. Spread spectrum achieves its goals through two principles:

The bandwidth allocated to each station needs to be, by far, larger than what is needed. This allows redundancy.

The expanding of the original bandwidth B to the bandwidth B_{SS} must be done by a process that is independent of the original signal. In other words, the spreading process occurs after the signal is created by the source.

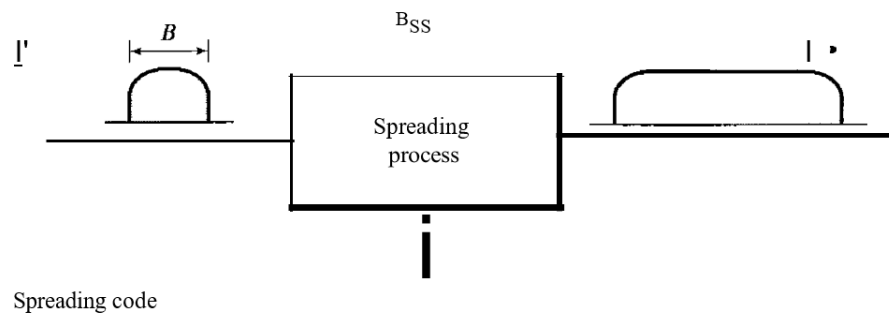


Figure 5.3 Spread spectrum.

After the signal is created by the source, the spreading process uses a spreading code and spreads the bandwidth. The figure shows the original bandwidth B and the spreaded bandwidth B_{SS} . The spreading code is a series of numbers that look random, but are a pattern.

There are two techniques to spread the bandwidth: frequency hopping spread spectrum (FHSS) and direct sequence spread spectrum (DSSS).

5.3.1 Frequency Hopping Spread Spectrum (FHSS)

The frequency hopping spread spectrum (FHSS) technique uses M different carrier frequencies that are modulated by the source signal. At one moment, the signal modulates one carrier frequency; at the next moment, the signal modulates another carrier frequency. Although the modulation is done using one carrier frequency at a time, M frequencies are used in the long run. The bandwidth occupied by a source after spreading is $B_{FHSS} \gg B$.

Figure 5.3.1 shows the general layout for FHSS. A pseudorandom code generator called pseudorandom noise (PN), creates a k -bit pattern for every hopping period T_h .

The frequency table uses the pattern to find the frequency to be used for this hopping period and passes it to the frequency synthesizer. The frequency synthesizer creates a carrier signal of that frequency, and the source signal modulates the carrier signal.

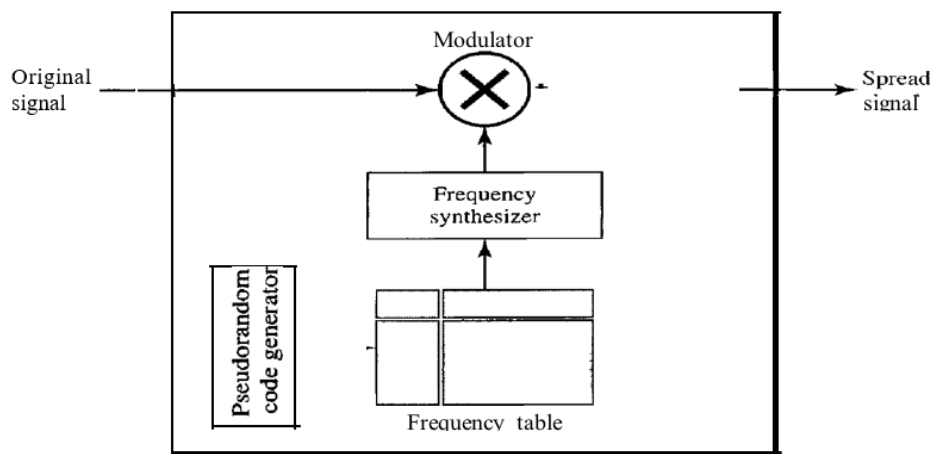
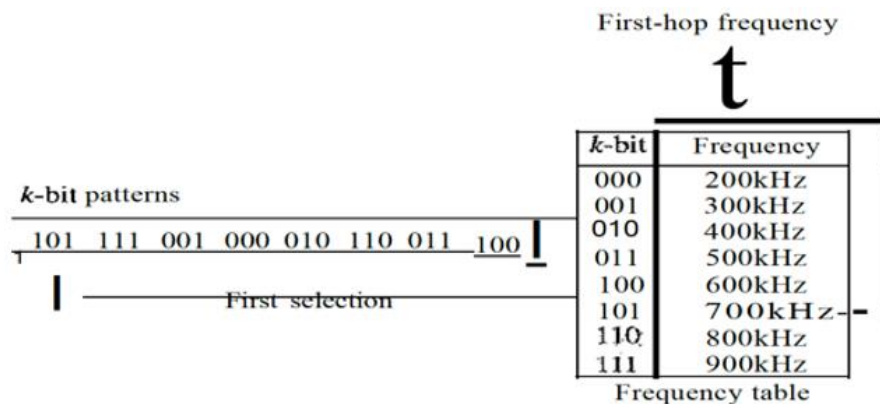


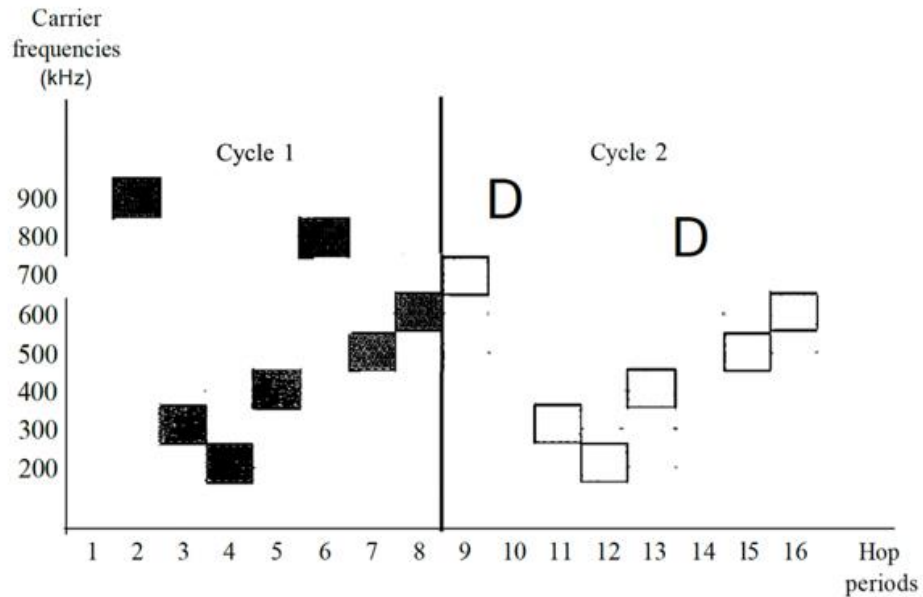
Figure 5.3.1 Frequency hopping spread spectrum (FHSS)

Suppose we have decided to have eight hopping frequencies. This is extremely low for real applications and is just for illustration. In this case, M is 8 and k is 6. The pseudo random code generator will create eight different 3-bit patterns. These are mapped to eight different frequencies in the frequency table (see Figure below)



Frequency selection in FHSS

The pattern for this station is 101, 111, 001, 000, 010, all, 100. Note that the pattern is pseudorandom it is repeated after eight hopping. This means that at hopping period 1, the pattern is 101. The frequency selected is 700 kHz; the source signal modulates this carrier frequency. The second k-bit pattern selected is 111, which selects the 900-kHz carrier; the eighth pattern is 100, the frequency is 600 kHz. After eight hopping, the pattern repeats, starting from 101 again. Figure below shows how the signal hops around from carrier to carrier. We assume the required bandwidth of the original signal is 100 kHz.



It can be shown that this scheme can accomplish the previously mentioned goals. If there are many k -bit patterns and the hopping period is short, a sender and receiver can have privacy. If an intruder tries to intercept the transmitted signal, she can only access a small piece of data because she does not know the spreading sequence to quickly adapt herself to the next hop. The scheme has also an antijamming effect. A malicious sender may be able to send noise to jam the signal for one hopping period (randomly), but not for the whole period.

5.3.2 Bandwidth Sharing

If the number of hopping frequencies is M , we can multiplex M channels into one by using the same B_{SS} bandwidth. This is possible because a station uses just one frequency in each hopping period; $M - 1$ other frequencies can be used by other $M - 1$ stations. In other

words, M different stations can use the same B_{SS} if an appropriate modulation technique such as multiple FSK (MFSK) is used. FHSS is like FDM, as shown in Figure 5.3.2

Figure 5.3.2 shows an example of four channels using FDM and four channels using FHSS. In FDM, each station uses $1/M$ of the bandwidth, but the allocation is fixed; in FHSS, each station uses $1/M$ of the bandwidth, but the allocation changes hop to hop.

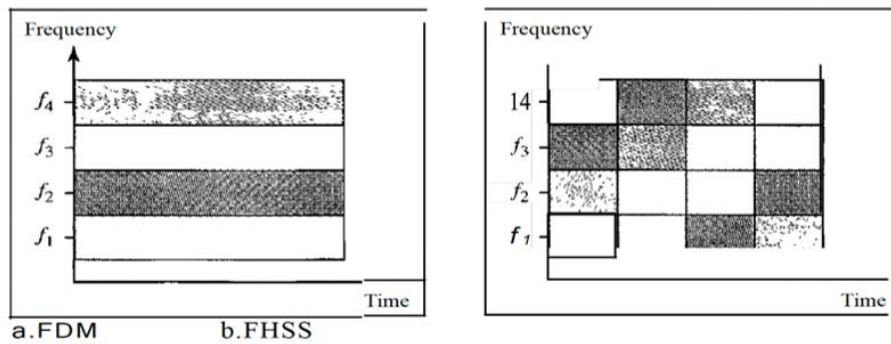


Figure 5.3.2 Bandwidth Sharing

5.3.3 Direct Sequence Spread Spectrum (DSSS)

The direct sequence spread spectrum (DSSS) technique also expands the bandwidth of the original signal, but the process is different. In DSSS, we replace each data bit with bits using a spreading code. In other words, each bit is assigned a code of bits, called chips, where the chip rate is times that of the data bit. Figure 5.3.3 shows the concept of DSSS.

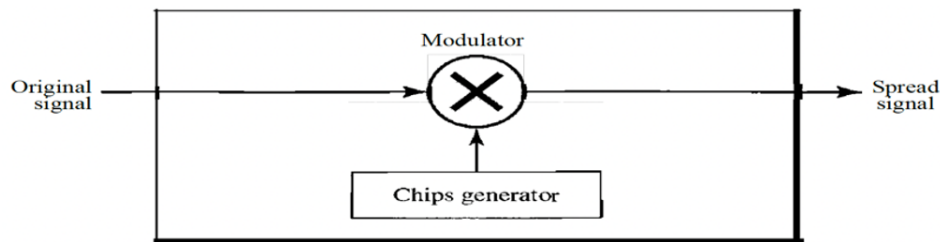


Figure 5.3.3 DSSS

As an example, let us consider the sequence used in a wireless LAN, the famous Barker sequence where is 11. We assume that the original signal and the chips in the chip generator use polar NRZ encoding. Figure 5.3.4 shows the chips and the result of multiplying the original data by the chips to get the spread signal.

In Figure 5.3.4, the spreading code is 11 chips having the pattern 10110111000 (in this case). If the original signal rate is N , the rate of the spread signal is $11N$. This means that the required bandwidth for the spread signal is 11 times larger than the bandwidth of the original signal. The spread signal can provide privacy if the intruder does not know the code. It can also provide immunity against interference if each station uses a different code.

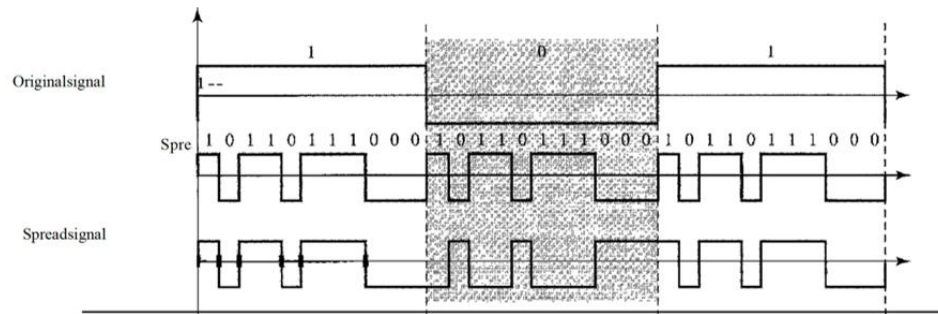


Figure 5.3.4 DSSS Example

Review Questions

1. Describe the goals of multiplexing.
2. List three main multiplexing techniques mentioned in this chapter.
3. Distinguish between a link and a channel in multiplexing.
4. Which of the three multiplexing techniques is (are) used to combine analog signals? Which of the three multiplexing techniques is (are) used to combine digital signals?
5. Define the analog hierarchy used by telephone companies and list different levels of the hierarchy.
6. Define the digital hierarchy used by telephone companies and list different levels of the hierarchy.
7. Which of the three multiplexing techniques is common for fiber optic links? Explain the reason.
8. Distinguish between synchronous and statistical TDM.
9. Define spread spectrum and its goal. List the two spread spectrum techniques discussed in this chapter.
10. Define FHSS and explain how it achieves bandwidth spreading.
11. Define DSSS and explain how it achieves bandwidth spreading.
12. Which of the three analog-to-analog conversion techniques (AM, FM, or PM) is the most susceptible to noise?

TRANSMISSION MEDIA AND SWITCHING

Unit Structure:

- 6.0 Introduction
- 6.1 GUIDED MEDIA
 - 6.1.1 Twisted-Pair Cable
 - 6.1.2 Unshielded versus Shielded Twisted-pair cable
 - 6.1.3 Coaxial Cable
 - 6.1.4 Coaxial Cable Connectors
 - 6.1.5 Fiber-Optic Cable
 - 6.1.6 Propagation Modes
- 6.2 UNGUIDED MEDIA: WIRELESS
 - 6.2.1 Radio Waves
 - 6.2.2 Microwaves
 - 6.2.3 Infrared
 - 6.2.4 Switching
- 6.3 CIRCUIT-SWITCHED NETWORKS
 - 6.3.1 Three Phases
 - 6.3.2 Efficiency
 - 6.3.3 Delay
 - 6.3.4 Circuit-Switched Technology in Telephone Networks
- 6.4 Packet Switching
 - 6.4.1 Datagram Packet Switching
 - 6.4.2 Routing Table
 - 6.4.3 Efficiency
 - 6.4.4 Delay
 - 6.4.5 Advantages of Package Switching
 - 6.4.6 Disadvantages of Package Switching
- 6.5 Message Switching
 - 6.5.1 Advantages of Message Switching
 - 6.5.2 Disadvantages of Message Switching

6.0 INTRODUCTION

Transmission media are actually located below the physical layer and are directly controlled by the physical layer. You could say that transmission media belong to layer zero. Figure 6.0.1 shows the position of transmission media in relation to the physical layer.

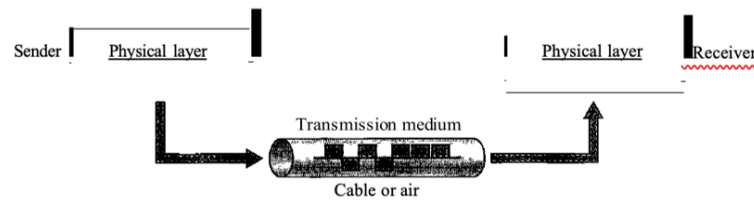


Figure 6.0.1 Transmission media and physical layer

A transmission **medium** can be broadly defined as anything that can carry information from a source to a destination. For example, the transmission medium for two people having a dinner conversation is the air. The air can also be used to convey the message in a smoke signal or semaphore. For a written message, the transmission medium might be a mail carrier, a truck, or an airplane.

In data communications the definition of the information and the transmission medium is more specific. The transmission medium is usually free space, metallic cable, or fiber-optic cable. The information is usually a signal that is the result of a conversion of data from another form.

The use of long-distance communication using electric signals started with the invention of the telegraph by Morse in the 19th century. Communication by telegraph was slow and dependent on a metallic medium.

Extending the range of the human voice became possible when the telephone was invented in 1869. Telephone communication at that time also needed a metallic medium to carry the electric signals that were the result of a conversion from the human voice.

The communication was, however, unreliable due to the poor quality of the wires. The lines were often noisy, and the technology was unsophisticated.

Wireless communication started in 1895 when Hertz was able to send high frequency signals. Later, Marconi devised a method to send telegraph-type messages over the Atlantic Ocean.

We have come a long way. Better metallic media have been invented (twisted pair and coaxial cables, for example). The use of optical fibers has increased the data rate incredibly. Free space (air, vacuum, and water) is used more efficiently, in part due to the technologies (such as modulation and multiplexing) discussed in the previous chapters.

Computers and other telecommunication devices use signals to represent

data. These signals are transmitted from one device to another in the form of electromagnetic energy, which is propagated through transmission media.

Electromagnetic energy, a combination of electric and magnetic fields vibrating in relation to each other, includes power, radio waves, infrared light, visible light, ultraviolet light, and X, gamma, and cosmic rays. Each of these constitutes a portion of the electromagnetic spectrum. Not all portions of the spectrum are currently usable for tele-communications, however. The media to harness those that are usable are also limited to a few types.

In telecommunications, 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. Figure 6.0.2 shows this taxonomy.

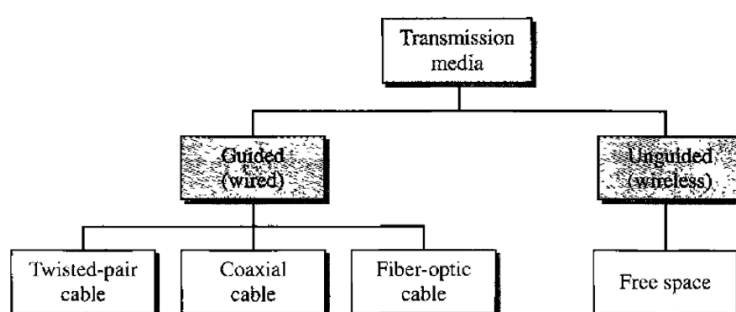


Figure 6.0.2 classes of transmission media

6.1 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.

6.1.1 Twisted-Pair Cable

A twisted pair consists of two conductors (normally copper), each with its own plastic insulation, twisted together, as shown in Figure 6.1.1.

Figure 6.1.1 *Twisted-pair cable*



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. For example, suppose in one twist, one wire is closer to the noise source and the other is farther; in the next twist, the reverse is true. Twisting makes it probable that both wires are equally affected by external influences (noise or crosstalk). 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.

6.1.2 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 6.1.2 shows the difference between UTP and STP. Our discussion focuses primarily on UTP because STP is seldom used outside of IBM.

Categories

The Electronic Industries Association (EIA) has developed standards to classify unshielded twisted-pair cable into seven categories. Categories are determined by cable quality, with 1 as the lowest and 7 as the highest. Each EIA category is suitable for specific uses. Below table show these categories.

Connectors

The most common UTP connector is RJ45 (RJ stands for registered jack). The RJ45 is a keyed connector, meaning the connector can be inserted in only one way.

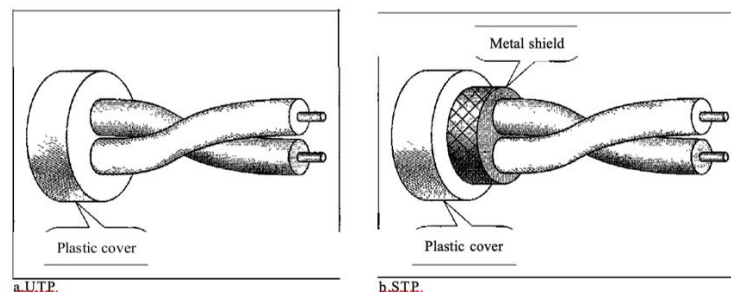


Figure 6.1.2 UTP and STP Cables

Table Categories of unshielded twisted-pair cables

<i>Category</i>	<i>Specification</i>	<i>Data Rate (Mbps)</i>	<i>Use</i>
1	Unshielded twisted-pair used in telephone	< 0.1	Telephone
2	Unshielded twisted-pair originally used in T-lines	2	T-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
SE	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

Performance

One way to measure the performance of twisted-pair cable is to compare attenuation versus frequency and distance. A twisted-pair cable can pass a wide range of frequencies.

Applications

Twisted-pair cables are used in telephone lines to provide voice and data channels. The local loop—the line that connects subscribers to the central telephone office—commonly consists of unshielded twisted-pair cables.

The DSL lines that are used by the telephone companies to provide high-data-rate connections also use the high-bandwidth capability of unshielded twisted-pair cables.

Local-area networks, such as 10Base-T and 100Base-T, also use twisted-pair cables.

6.1.3 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 (see Figure 6.1.3).

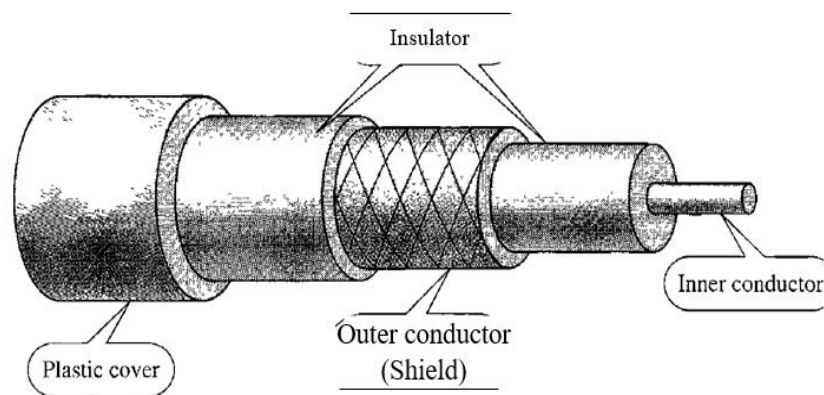


Figure 6.1.3 *Coaxial cable*

Coaxial Cable Standards

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, as shown in Table 6.1.1.

Table 6.1.1 *categories of coaxial cables*

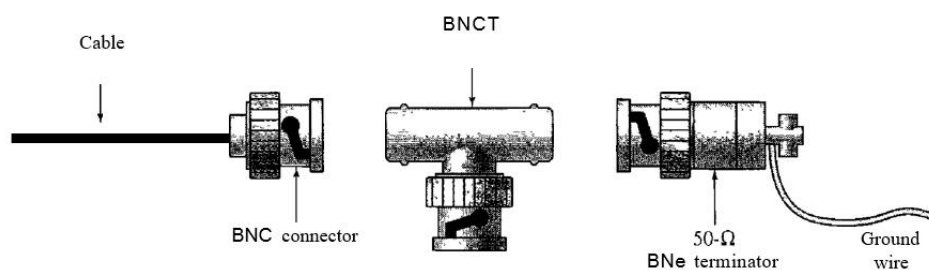
<i>Category</i>	<i>Impedance</i>	<i>Use</i>
RG-59	75 Ω	Cable TV
RG-58	50 Ω	Thin Ethernet
RG-11	50 Ω	Thick Ethernet

6.1.4 Coaxial Cable Connectors

To connect coaxial cable to devices, we need coaxial connectors. The most common type of connector used today is the Bayone-Neill-Concelman (BNC), connector. Figure 6.8 shows three popular types of these connectors: the BNC connector, the BNC T connector, and the BNC terminator.

The BNC connector is used to connect the end of the cable to a device, such as a TV set. The BNC T connector is used in Ethernet networks to branch out to a connection to a computer or other device. The BNC terminator is used at the end of the cable to prevent the reflection of the signal.

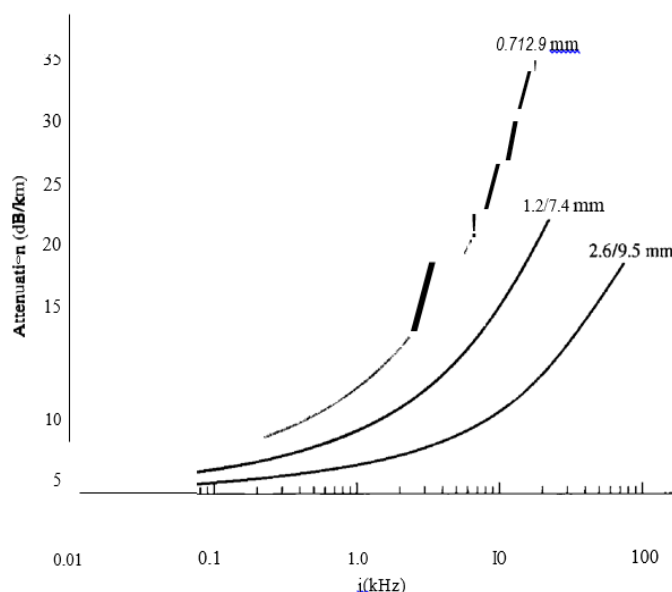
Figure 6.1.4 BNC connectors



Performance

As we did with twisted-pair cables, we can measure the performance of a coaxial cable. We notice in the figure given below, that the attenuation is much higher in coaxial cables than in twisted-pair cable. In other words, although coaxial cable has a much higher bandwidth, the signal weakens rapidly and requires the frequent use of repeaters.

Coaxial cable performance



Coaxial cable was widely used in analog telephone networks where a single coaxial network could carry 10,000 voice signals. Later it was used in digital telephone networks where a single coaxial cable could carry digital data up to 600 Mbps. However, coaxial cable in telephone networks has largely been replaced today with fiber-optic cable.

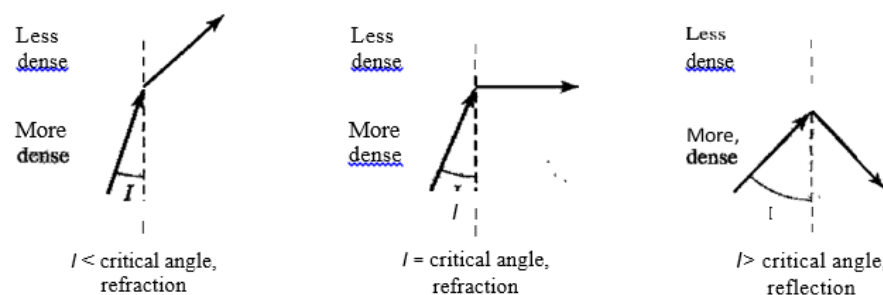
Cable TV networks also use coaxial cables. In the traditional cable TV network, the entire network used coaxial cable. Later, however, cable TV providers replaced most of the media with fiber-optic cable; hybrid networks use coaxial cable only at the network boundaries, near the consumer premises. Cable TV uses RG-59 coaxial cable.

Another common application of coaxial cable is in traditional Ethernet LANs. Because of its high bandwidth, and consequently high data rate, coaxial cable was chosen for digital transmission in early Ethernet LANs. The 10Base-2, or Thin Ethernet, uses RG-58 coaxial cable with BNC connectors to transmit data at 10 Mbps with a range of 185 m. The 10Base5, or Thick Ethernet, uses RG-11 (thick coaxial cable) to transmit 10 Mbps with a range of 5000 m. Thick Ethernet has specialized connectors.

6.1.5 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 if 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 direction. Figure 6.1.5 shows how a ray of light changes direction when going from a denser to a less dense substance.

Figure 6.1.5 bending of ray



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.

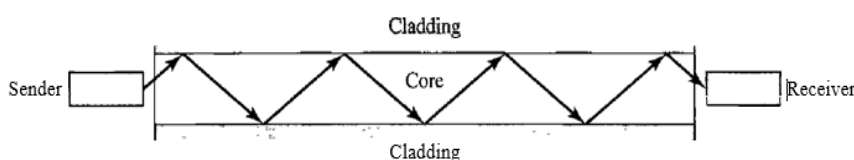
Note that the critical angle is a property of the substance, and its value differs from one substance to another.

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.

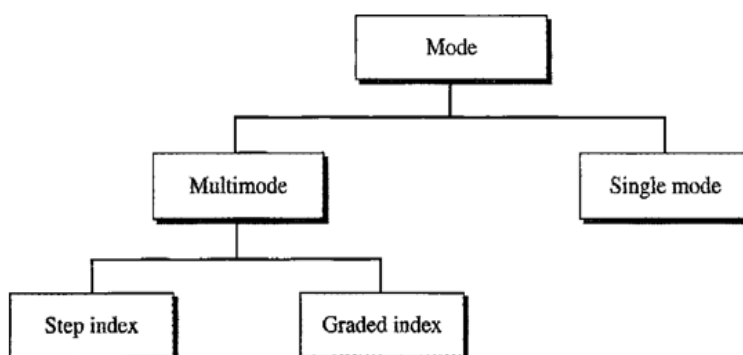
6.1.6 Propagation Modes

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.

Optical fiber



Propagation modes



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. (See Figure 6.1.6)

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.

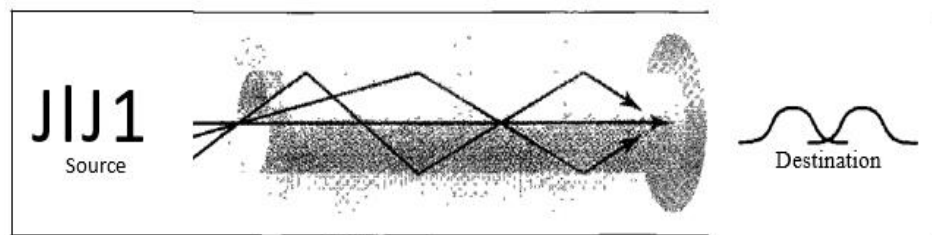
A second type of fiber, called multimode graded-index fiber, decreases this distortion of the signal through the cable. The word *index* here refers to the index of refraction. As we saw above, the index of refraction is related to density. A graded-index fiber, therefore, is one with varying densities. Density is highest at the center of the core and decreases gradually to its lowest at the edge. Figure 6.1.6 shows the impact of this variable density on the propagation of light beams.

Single-Mode

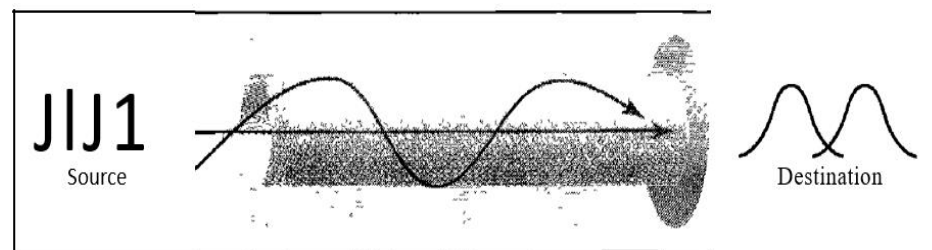
Single-mode uses step-index fiber and a highly focused source of light that limits beams to a small range of angles, all close to the horizontal. The single-mode fiber itself is manufactured with a much smaller diameter than that of multimode fiber, and with substantially lower density (index of refraction). 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 (see Figure 6.1.6).

Figure 6.1.6 *Modes*

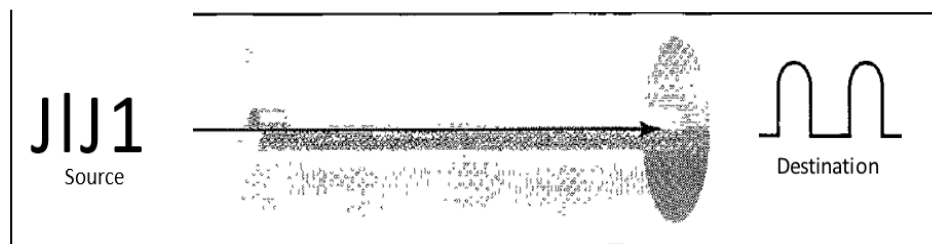
a. Multimode, step index



b. Multimode, graded index



c. Single mode



Optical fibers are defined by the ratio of the diameter of their core to the diameter of their cladding, both expressed in micrometers. The common sizes are shown in Table 6.1.2. Note that the last size listed is for single-mode only.

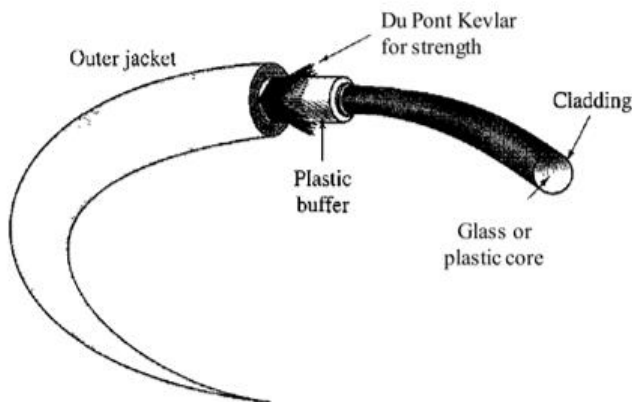
Table 6.1.2 Fiber types

Type	Core (μm)	Cladding (μm)	Mode
50/125	50.0	125	Multimode, graded index
62.5/125	62.5	125	Multimode, graded index
100/125	100.0	125	Multimode, graded index
7/125	7.0	125	Single mode

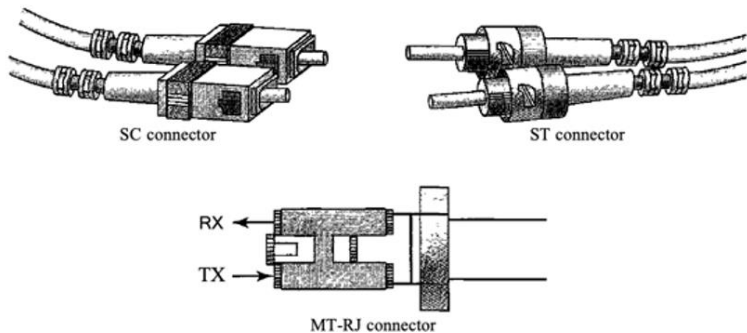
Fiber-Optic Cable Connectors

There are three types of connectors for fiber-optic cables, as shown in Figure.

Fiber Construction



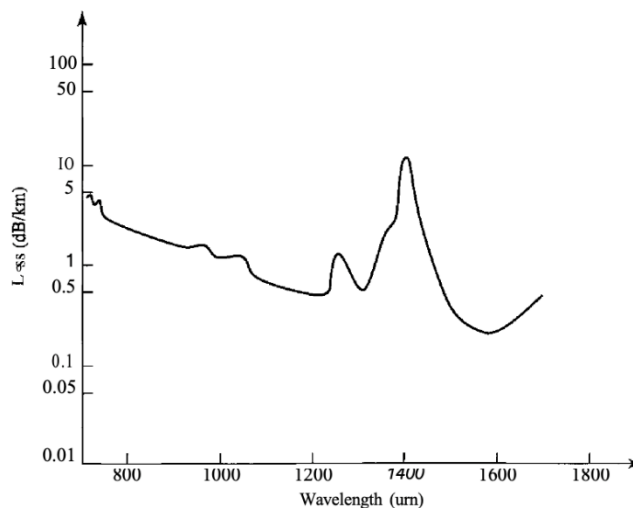
Fiber Optic Cable Connectors



The **subscriber channel (SC) connector** is used for cable TV. It uses a push/pull locking system. The **straight-tip (ST) connector** is used for connecting cable to networking devices. It uses a bayonet locking system and is more reliable than SC. **MT-RJ** is a connector that is the same size as RJ45.

Performance

The plot of attenuation versus wavelength in Figure below shows a very interesting phenomenon in fiber-optic cable. Attenuation is flatter than in the case of twisted-pair cable and coaxial cable. The performance is such that we need fewer (actually 10 times less) repeaters when we use fiber-optic cable.



Applications

Fiber-optic cable is often found in backbone networks because its wide bandwidth is cost-effective. Today, with wavelength-division multiplexing (WDM), we can transfer data at a rate of 1600 Gbps. The SONET network that we discuss in Chapter 17 provides such a backbone.

Some cable TV companies use a combination of optical fiber and coaxial cable, thus creating a hybrid network. Optical fiber provides the backbone structure while coaxial cable provides the connection to the user premises. This is a cost-effective configuration since the narrow bandwidth requirement at the user end does not justify the use of optical fiber. Local-area networks such as 100Base-FX network (Fast Ethernet) and 1000Base-X also use fiber-optic cable.

Advantages and Disadvantages of Optical Fiber

Advantages

- 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

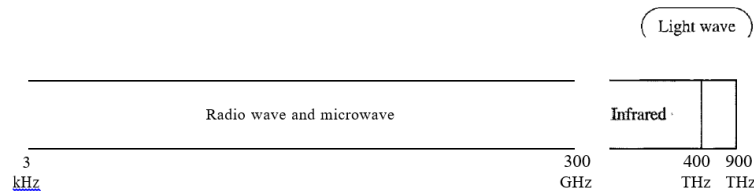
- Installation and maintenance. Fiber-optic cable is a relatively new technology. Its installation and maintenance require expertise that is not yet available everywhere.
- Unidirectional light propagation. Propagation of light is unidirectional. If we need bidirectional communication, two fibers are 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.

6.2 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 6.2.0 shows the part of the electromagnetic spectrum, ranging from 3 kHz to 900 THz, used for wireless communication.

Figure 6.2.0 Electromagnetic spectrum 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 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 low frequency (VLF) to extremely high frequency (EHF). Table 6.2.1 lists these bands, their ranges, propagation methods, and some applications.

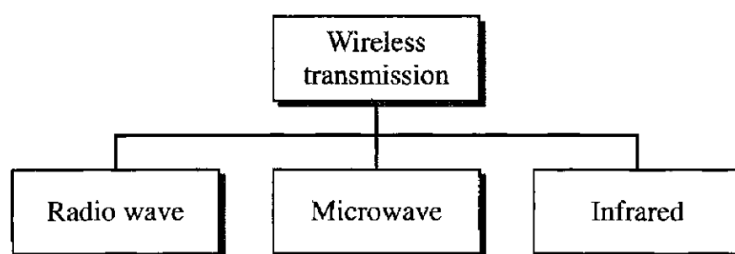
Table 6.2.1 Bands

<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 (super high frequency)	3-30 GHz	Line-of-sight	Satellite communication
EHF (extremely high frequency)	30-300 GHz	Line-of-sight	Radar, satellite

We can divide wireless transmission into three broad groups: radio waves, micro waves, and infrared waves. See Figure below

Wireless transmission waves



6.2.1 Radio Waves

Although there is no clear-cut demarcation between radio waves and microwaves, electromagnetic waves ranging in frequencies between 3 kHz and 1 GHz are normally called radio waves; waves ranging in frequencies between 1 and 300 GHz are called micro waves. However, the behavior of the waves, rather than the frequencies, is a better criterion for classification.

Radio waves, for the most part, 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, particularly those 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. The radio wave band is relatively narrow, just under 1 GHz, compared to the microwave band. When this band is divided into sub bands, the sub bands are also narrow, leading to a low data rate for digital communications.

Almost the entire band is regulated by authorities (e.g., the FCC in the United States). Using any part of the band requires permission from the authorities.

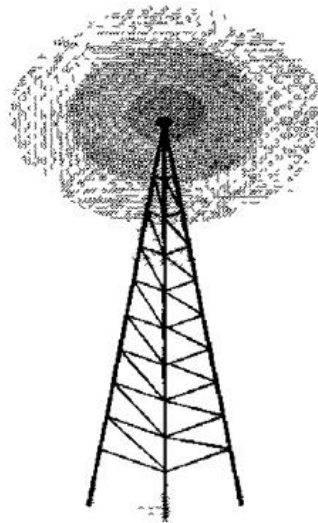
Omnidirectional Antenna

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 6.2.2 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.

Figure 6.2.2 omnidirectional antenna



Radio waves are used for multicast communications, such as radio and television, and paging systems.

6.2.2 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:

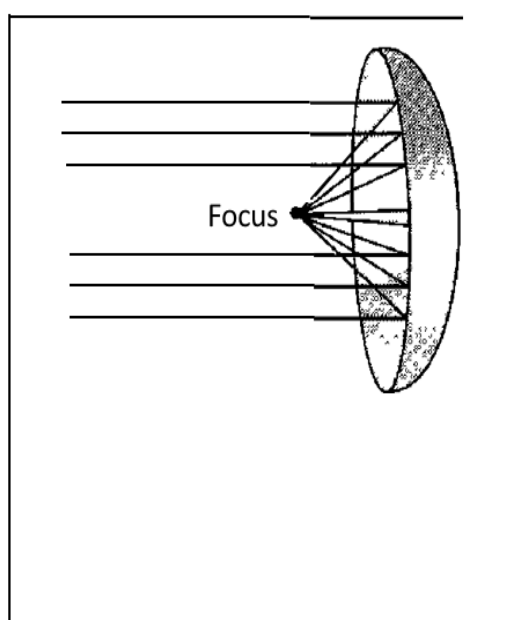
- Microwave propagation is line-of-sight. Since the towers with the mounted antennas need to be in direct sight of each other, towers that are far apart need to be very tall. The curvature of the earth as well as other blocking obstacles do not allow two short towers to communicate by using microwaves. 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.
- The microwave band is relatively wide, almost 299 GHz. Therefore, wider sub bands can be assigned, and a high data rate is possible.
- Use of certain portions of the band requires permission from authorities.

Unidirectional Antenna

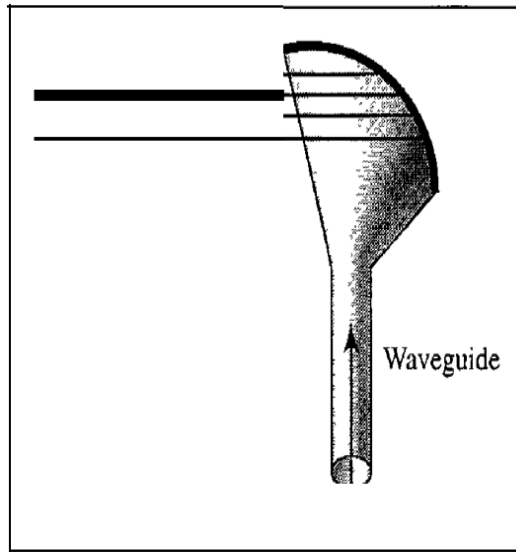
Microwaves need unidirectional antennas that send out signals in one direction. Two types of antennas are used for microwave communications: the parabolic dish and the horn (see below figure).

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,

a. Dish antenna



b. Horn antenna



Catching a wide range of waves and directing them to a common point. In this way, more of the signal is recovered than would be possible with a single-point receiver.

Outgoing transmissions are broadcast through a horn aimed at the dish. The micro waves 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 (resembling a handle) 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 like the parabolic dish, and are deflected down into the stem.

Applications

Microwaves, due to their unidirectional properties, are very useful when unicast (one-to-one) communication is needed between the sender and the receiver. They are used in cellular phones (Chapter 16), satellite networks (Chapter 16), and wireless LANs (Chapter 14).

Microwaves are used for unicast communication such as cellular telephones, satellite networks, and wireless LANs.

6.2.3 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 band, almost 400 THz, has an excellent potential for data transmission. Such a wide bandwidth can be used to transmit digital data with a very high data rate. 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. For example, some manufacturers provide a special port called the IrDA port that allows a wireless keyboard to communicate with a PC. The standard originally defined a data rate of 75 kbps for a distance up to 8 m. The recent standard defines a data rate of 4 Mbps.

Infrared signals defined by IrDA transmit through line of sight; the IrDA port on the keyboard needs to point to the PC for transmission to occur.

Infrared signals can be used for short-range communication in a closed area using line-of-sight propagation.

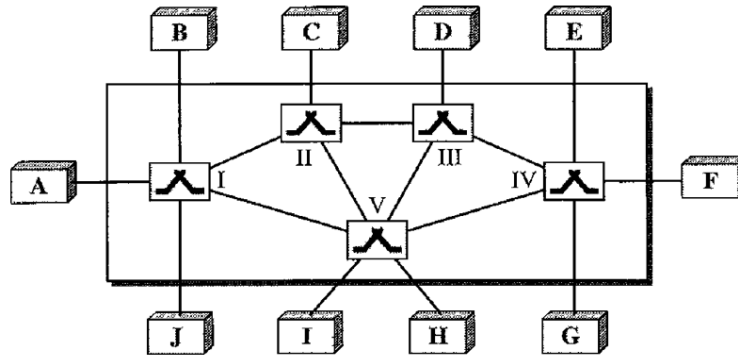
6.2.4 Switching

A network is a set of connected devices. Whenever we have multiple devices, we have the problem of how to connect them to make one-to-one communication possible. One solution is to make a point-to-point connection between each pair of devices (a mesh topology) or between a central device and every other device (a star topology). These methods, however, are impractical and wasteful when applied to very large networks. The number and length of the links require too much infrastructure to be cost-efficient, and most of those links would be idle most of the time. Other topologies employing multipoint connections, such as a bus, are ruled out because the distances between devices and the total number of devices increase beyond the capacities of the media and equipment.

A better solution is switching. A switched network consists of a series of interlinked

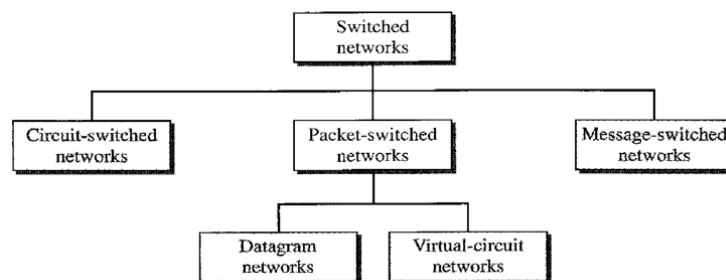
Nodes, called switches. Switches are devices capable of creating temporary connections between two or more devices linked to the switch. In a switched network, some of these nodes are connected to the end systems (computers or telephones, for example). Others are used only for routing. Figure below shows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

Switched network



The end systems (communicating devices) are labeled A, B, C, D, and so on, and the switches are labeled I, II, III, IV, and V. Each switch is connected to multiple links.

Traditionally, three methods of switching have been important: circuit switching, packet switching, and message switching. The first two are commonly used today. The third has been phased out in general communications but still has networking applications. We can then divide today's networks into three broad categories: circuit-switched networks, packet-switched networks, and message-switched. Packet-switched networks can further be divided into two subcategories-virtual-circuit networks and datagram networks as shown below:

Taxonomy of switched networks

We can say that the virtual-circuit networks have some common characteristics with circuit-switched and datagram networks. Thus, we first discuss circuit-switched networks, then datagram networks, and finally virtual-circuit networks.

Today the tendency in packet switching is to combine datagram networks and virtual-circuit networks. Networks route the first packet based on the datagram addressing idea, but then create a virtual-circuit network for the rest of the packets coming from the same source and going to the same destination. We will see some of these networks in future chapters.

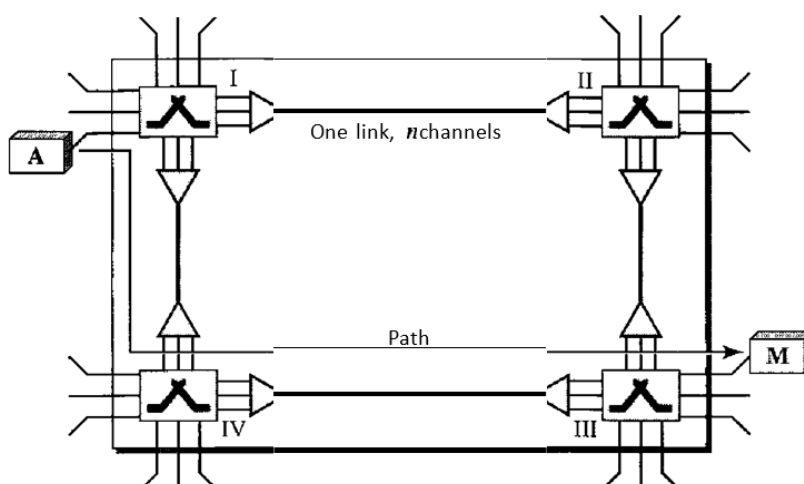
In message switching, each switch stores the whole message and forwards it to the next switch. Although, we don't see message switching at lower layers, it is still used in some applications like electronic mail (e-mail). We will not discuss this topic in this book.

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 as discussed in Chapter 6.

A circuit-switched network is made of a set of switches connected by physical links, in which each link is divided into n channels.

Figure 6.3.0 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.

Figure 6.3.0 a trivial circuit-switched network



We have explicitly shown the multiplexing symbols to emphasize the division of the link into channels even though multiplexing can be implicitly included in the switch fabric.

The end systems, such as computers or telephones, are directly connected to a switch. We have shown only two end systems for simplicity. 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.

We need to emphasize several points here:

- 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 (bandwidth in FDM and time slots in TDM), switch buffers, switch processing time, and switch input/output ports, must remain dedicated during the entire duration of data transfer until the teardown phase.

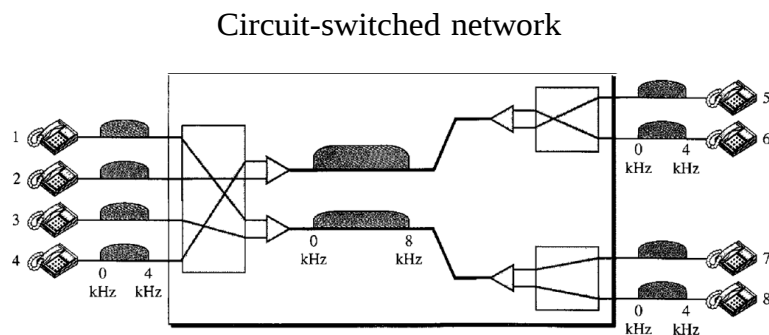
- Data transferred between the two stations are not packetized (physical layer transfer of the signal). The data are a continuous flow sent by the source station and received by the destination station, although there may be periods of silence.
- There is no addressing involved during data transfer. The switches route the data based on their occupied band (FDM) or time slot (TDM). Of course, there is end-to-end addressing used during the setup phase, as we will see shortly.

In circuit switching, the resources need to be reserved during the setup phase; the resources remain dedicated for the entire duration of data transfer until the teardown phase.

Example 1

As a trivial example, let us use a circuit-switched network to connect eight telephones in a small area. Communication is through 4-kHz voice channels. We assume that each link uses FDM to connect a maximum of two voice channels. The bandwidth of each link is then 8 kHz. Figure 8.4 shows the situation. Telephone 1 is connected to telephone 7; 2 to 5; 3 to 8; and 4 to 6. Of course, the situation may change when new connections are made. The switch controls the connections.

Circuit-switched network used in Example 1

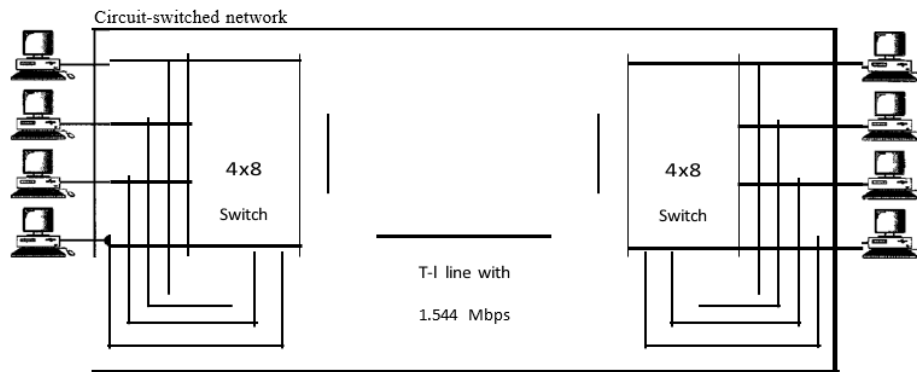


Example 2

As another example, consider a circuit-switched network that connects computers in two remote offices of a private company. The offices are connected using a T-1 line leased from a communication service provider. There are two 4 X 8 (4 inputs and 8 outputs) switches in this network. For each switch, four output ports are folded into the input ports to allow communication between computers in the same office. Four other output ports allow communication between the two offices. Figure 6.2 shows the situation.

Circuit-switched network used in Example 2

Circuit-switched network



6.3.1 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 (or multiple parties in a conference call) can communicate, a dedicated circuit (combination of channels in links) 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, in below figure 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 send 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 currently.

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.

Note that end-to-end addressing is required for creating a connection between the two end systems. These can be, for example, the addresses of the computers assigned by the administrator in a TDM network, or telephone numbers in an FDM network.

Data Transfer Phase

After the establishment other 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.

6.3.2 Efficiency

The resources remain dedicated as long as the connection is alive. Due to allocation of resources during the entire duration of the connection, the efficiency of circuit switched networks is lower than other two types of switching.

6.3.3 Delay

Even if the efficiency is low, delay is very small in circuit switched networks. During data transfer, data is not delayed at any switch because there is no waiting time involved.

6.3.4 Circuit-Switched Technology in Telephone Networks

The telephone companies previously used the circuit switching technology for switching and routing a call. This was a physical layer technology.

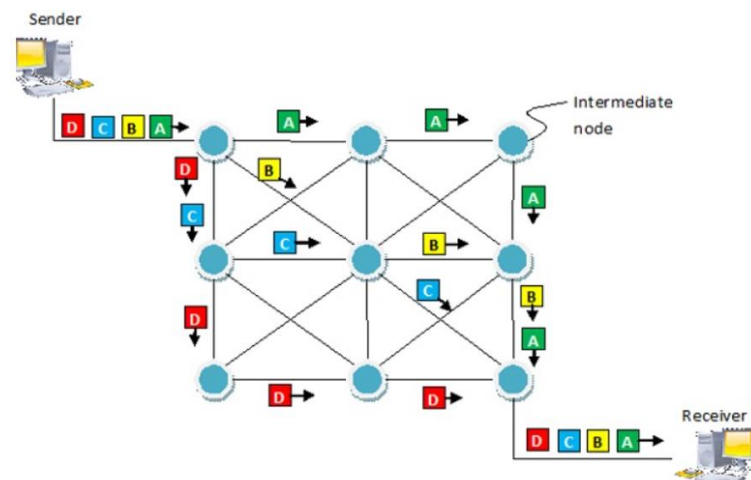
However, today the tendency is to use other switching technologies. For example the telephone number is used as the global address and a signaling system is used for creating and disconnecting connections

6.4 PACKET SWITCHING

Packet switching is a connectionless network switching technique. Here, the message is divided and grouped into a number of units called packets that are individually routed from the source to the destination. There is no need to establish a dedicated circuit for communication.

A packet is transmitted as soon as it is available in a node, based upon its header information. The packets of a message are not routed via the same path. So, the packets in the message arrives in the destination out of order. It is the responsibility of the destination to reorder the packets in order to retrieve the original message.

The process is diagrammatically represented in the following figure. Here the message comprises of four packets, A, B, C and D, which may follow different routes from the sender to the receiver.



6.4.1 Datagram Packet Switching

In this method, a message is divided into a stream of packets. Each packet has its individually included address and treated as an independent unit with its own control instructions.

The switching devices would route each packet independently through the network. Each intermediate node will determine the packet's next route segment.

Before transmission starts, the sequence of packets and their destinations are communicated by exchanging control information between the sending terminal, the network and the receiving terminal.

The datagram packet switching generally corresponds to the network layer. The packets are called as datagrams. The switches in the datagram network are called as routers.

The datagram networks are called as connectionless networks because the switch does not keep any information about the connection state.

6.4.2 Routing table

In packet switched networks, each packet switch has a routing table. This table contains the destination address. This table is dynamic and their information updates on a periodic basis.

Destination Address	Output Port
1323	1
4360	2
9140	3
6436	4

6.4.3 Efficiency

As resources are allocated only when the packets are to be transferred, the efficiency of datagram network is higher than that of the circuit switched network.

6.4.4 Delay

There are no set up or tear down phases in datagram switching but each packet may have to wait at a switch before getting forwarded.

All the packets in a message take different paths. Hence the delay associated with each packet is different.

6.4.5 Advantages of Packet Switching:

Advantages

- Delay in delivery of packets is less, since packets are sent as soon as they are available.

- Switching devices don't require massive storage, since they don't have to store the entire messages before forwarding them to the next node.
- Data delivery can continue even if some parts of the network faces link failure. Packets can be routed via other paths.
- It allows simultaneous usage of the same channel by multiple users.
- It ensures better bandwidth usage as a number of packets from multiple sources can be transferred via the same link.

6.4.6 Disadvantages of Package Switching

Disadvantages

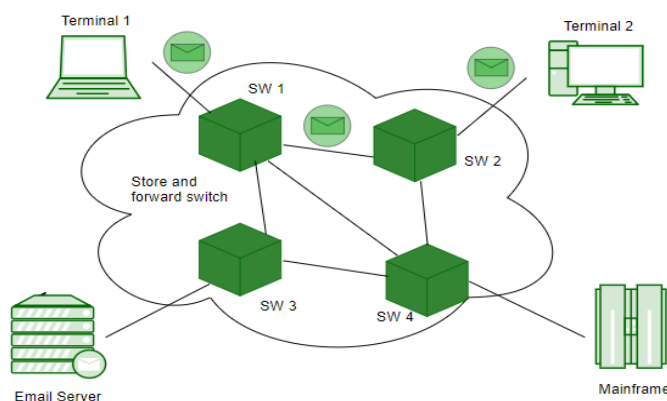
- They are unsuitable for applications that cannot afford delays in communication like high quality voice calls.
- Packet switching high installation costs.
- They require complex protocols for delivery.
- Network problems may introduce errors in packets, delay in delivery of packets or loss of packets. If not properly handled, this may lead to loss of critical information.

6.5 MESSAGE SWITCHING

In message switching, end-users communicate by sending and receiving *messages* that included the entire data to be shared. Messages are the smallest individual unit.

Also, the sender and receiver are not directly connected. There are a number of intermediate nodes that transfer data and ensure that the message reaches its destination.

Message switches can be programmed with information about the most efficient routes as well as information regarding neighbouring switches that can be used to forward messages to their ultimate destination



6.5.1 Advantages of Message Switching –

1. As message switching is able to store the message for which communication channel is not available, it helps in reducing the traffic congestion in the network.
2. In message switching, the data channels are shared by the network devices. It makes traffic management efficient by assigning priorities to the messages.
3. Because the messages are delivered via a store and forward method, it is possible to include priority in them.
4. It allows for infinite message lengths.
5. Unlike circuit switching, it does not necessitate the actual connection of source and destination devices.

6.5.2 Disadvantages of Message Switching –

1. Message switching cannot be used for real-time applications as storing messages causes delay.
2. In message switching, the message has to be stored for which every intermediate device in the network requires a large storing capacity
3. Because the system is so intricate, people are frequently unaware of whether or not messages are correctly conveyed. This could cause problems in social relationships.
4. The type of message switching does not create a dedicated path between the devices. It is not dependable communication because there is no direct relationship between sender and receiver.

Review Questions

1. Name the two major categories of transmission media.
2. How do guided media differ from unguided media?
3. What are the three major classes of guided media?
4. What is the significance of the twisting in twisted-pair cable?
5. Name the advantages of optical fiber over twisted-pair and coaxial cable.
6. Explain concept of datagram packet switching.
7. State the advantages and disadvantages of datagram packet switching
8. Explain the three switching methods.



INTRODUCTION TO DATA LINK LAYER

Unit Structure:

- 7.0 Objectives
- 7.1 Introduction to Data-Link Layer
- 7.2 Nodes and Links
- 7.3 Services
- 7.4 Two Sub-layers
- 7.5 Three Types of addresses
- 7.6 Address Resolution Protocol (ARP)
- 7.7 Error Detection and Correction: Introduction
- 7.8 Types of Errors
- 7.9 Redundancy
- 7.10 Detection versus Correction
- 7.11 Summary
- 7.12 List of References
- 7.13 Unit End Exercises

7.0 OBJECTIVES

- To understand the features, services and concept of data link layer
- To get familiar with data correction and detection mechanism

7.1 INTRODUCTION TO DATA-LINK LAYER

The physical layer or data-link layer protocols are not defined by the TCP/IP protocol suite. The Internet is made up of two levels, which are the domains of many networks. These networks, whether wired or wireless, offer services to the top three TCP/IP layers. This gives us a hint that there are a number of industry-standard protocols available right now.

A collection of networks is connected by connecting devices to form the Internet (routers or switches). A packet must pass through various networks in order to move from one host to another. The data-link layer's communication is depicted in Figure 1. Five distinct logical connections between the data-link layers in the path make up communication at the data-link layer.

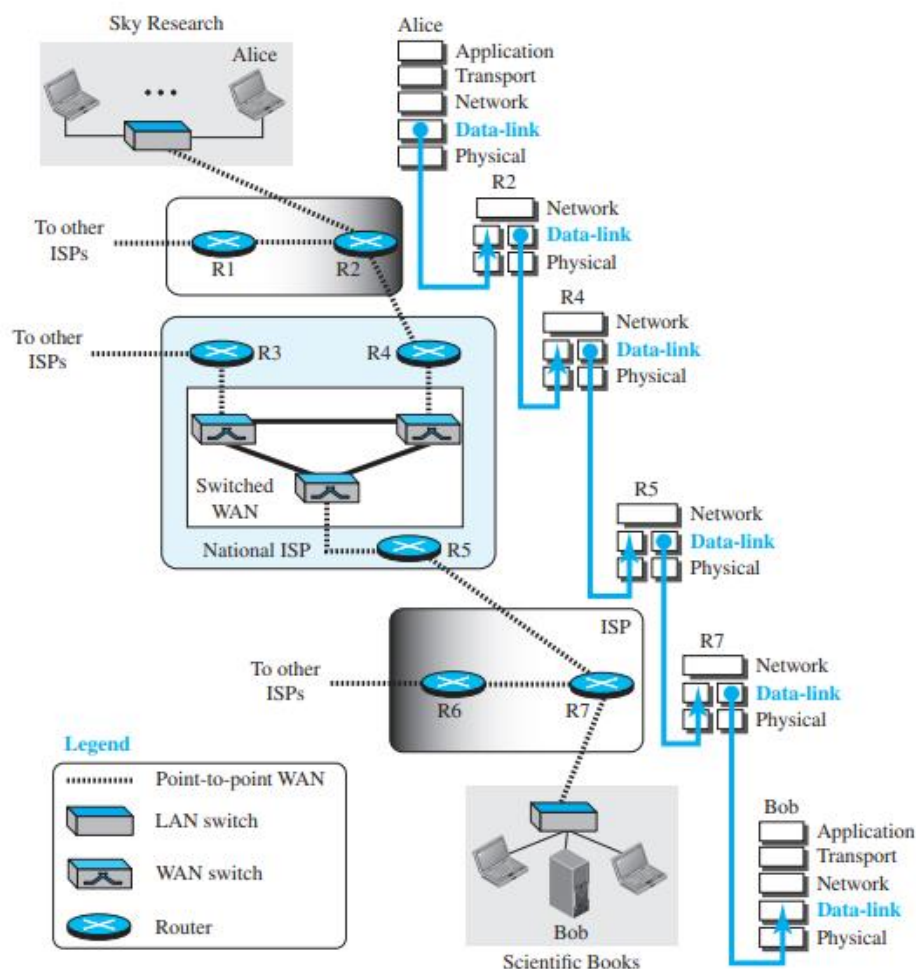


Figure 1: Communication at data link layer

Alice's PC and router R2 can connect with each other via the data-link layer. Communication takes place between the data-link layer at router R2 and the data-link layer at router R4, and so on. Finally, Bob's computer and router R7's data-link layer are in contact with one another. At the source or the destination, only one data-link layer is involved; however, two data-link levels are involved at each router. Since each router receives input from one network and delivers output to a different network, even though Alice and Bob's computers are both linked to the same network. Although switches are also engaged in communication at the data-link layer, we have omitted them from the figure for simplicity.

7.2 NODES AND LINKS

Node-to-node communication takes place at the data-link layer. To get from one location on the Internet to another, a data unit must travel over numerous networks (LANs and WANs). Routers link these LANs and WANs together. The two end hosts, the routers, and the networks in between are commonly referred to as nodes. When the data unit's path contains only six nodes, Figure 2 shows links and nodes simply.

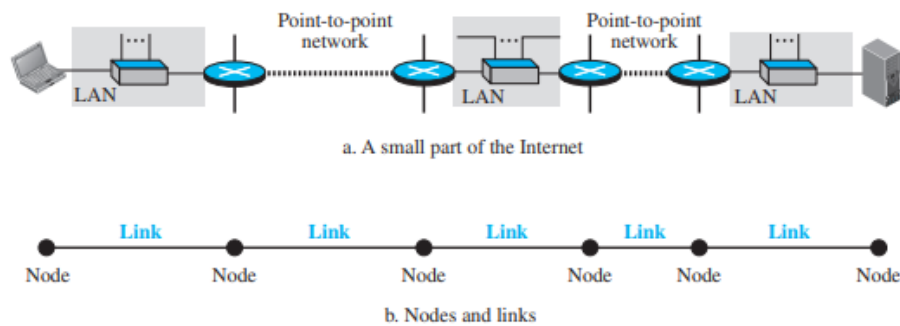


Figure 2: Nodes and Links

The source host is on the first node, and the destination host is on the last node. Four routers make up the remaining four nodes. The two WANs are represented by the second and fourth links, while the three LANs are represented by the first, third, and fifth links.

7.3 SERVICES

Between the physical and network levels is the data-link layer. In addition to receiving services from the physical layer, the data-link layer also offers services to the network layer. Let's talk about the services the data-link layer offers. The data-link layer's node-to-node duty scope. The data-link layer of a node (host or router) is in charge of sending a datagram to the following node along the path when a packet is travelling over the Internet. To accomplish this, the sending node's data-link layer must encapsulate the datagram it received from the network in a frame, and the receiving node's data-link layer must decapsulate the datagram from the frame. In other words, each intermediary node must do both encapsulation and decapsulation, whereas the data-link layer of the source host only needs to encapsulate and the data-link layer of the destination host only needs to decapsulate. Why encapsulation and decapsulation are required at each intermediate node may be a question. Because each link can be running a separate protocol with a distinct frame format, this is the cause. Encapsulation and decapsulation are required because link-layer addresses are typically distinct even when two links are utilising the same protocol. In this scenario, an analogy might be useful. Let's say someone has to get from their house to a friend's house in a different city. Three modes of transportation are available to the traveller. She can take a taxi to the railway station in her own city, take the train from there to the city where her friend lives, and then take another taxi to get to her friend's house. A source node, a destination node, and two intermediary nodes are present here. At the source node, the traveller must board a taxi; at the first intermediate node (a train station in the city where she lives), she must exit the taxi and board a train; at the second intermediate node (a train station in the city where her friend lives); and at the final intermediate node (her destination), she must exit the taxi. The source node experiences some form of encapsulation, the intermediate nodes experience encapsulation and decapsulation, and the destination node experiences decapsulation. Although our traveller is the same, she reaches her destination through three conveying devices.

The encapsulation and decapsulation at the data-link layer are depicted in Figure 3. We've assumed, for the sake of simplicity, that there is just one router in the path from source to destination. A frame is used to encapsulate the datagram that the source host's data-link layer has received. From the source host to the router, the frame is logically moved. In the router's data-link layer, the frame is decapsulated and then re-encapsulated. The destination host receives the new frame logically from the router. Due to its connections to three physical links, the router truly has three data-link layers, despite the fact that we have only depicted two in our diagram.

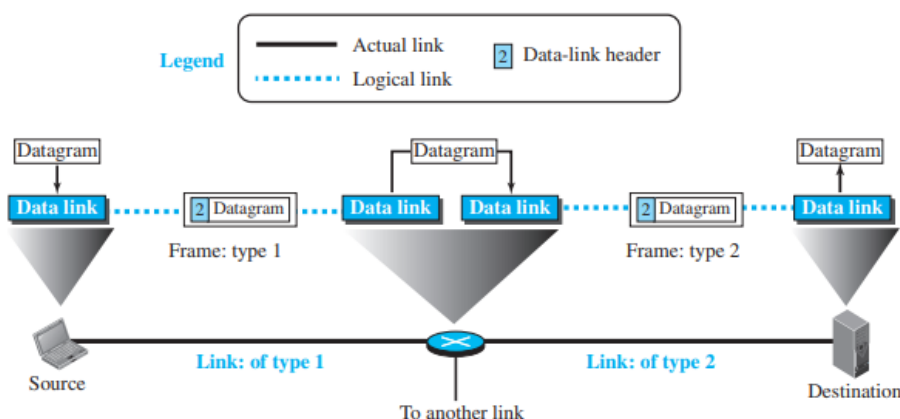


Figure 3: A communication with only three nodes

Following are the services provided by the data link layer:

- **Framing**

Without a doubt, framing is the first service the data-link layer offers. Before delivering the datagram (packet received from the network layer) to the following node, the data-link layer at each node must encapsulate it in a frame. Also, the datagram from the frame that was received on the logical channel must be decapsulated by the node. We have just showed the header for a frame, but in later chapters we will see that a frame can have both a header and a trailer. The framing formats for different data-link layers vary.

A packet at the data-link layer is normally called a frame.

- **Flow Control**

We must consider flow control whenever there is a producer and a consumer. The build-up of objects happens when the producer produces goods that cannot be used. The receiving data-link layer is a consumer whereas the sending data-link layer at one end of the link is a producer of frames. Frames at the receiving end must be buffered while awaiting consumption if the rate of produced frames is higher than the rate of consumed frames (processed). Without a doubt, the receiving side's buffer size cannot be infinite. Two options are available. If the receiving data-link layer's buffer is full, the first option is to let it drop the frames. The second option is to allow the

receiving data-link layer to instruct the transmitting data-link layer to halt or slow down by sending feedback to it. Various data-link-layer protocols employ various flow control techniques.

- **Error control**

A frame in a data-link layer needs to be converted to bits at the sending node before being converted to electromagnetic signals and sent over the transmission media. Electromagnetic signals are gathered at the receiving node, converted to bits, and assembled into frames. A frame is mistake-prone because electromagnetic impulses are prone to inaccuracy. First, the error must be found. It must either be fixed at the receiving node after discovery or rejected and sent again by the sender node.

- **Congestion control**

Most data-link-layer protocols do not explicitly use a congestion control to relieve congestion, however some wide-area networks do. A connection may be congested with frames, which may cause frame loss. Due to its end-to-end nature, congestion control is typically viewed as a problem in the network layer or the transport layer.

7.4 TWO SUB-LAYERS

We may break down the data-link layer into two sublayers, data link control (DLC) and media access control, to better understand its functionality and the services it offers (MAC). Because LAN protocols actually employ the same technique, this is not unusual. The media access control sublayer solely handles problems particular to broadcast links; the data link control sublayer handles all problems shared by point-to-point and broadcast links. In other words, as depicted in Figure 4, we distinguish between these two categories of links at the data-link layer.

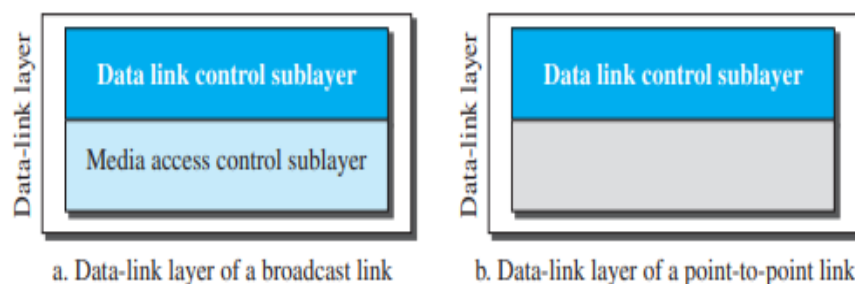


Figure 4: Dividing the data-link layer into two sublayers

7.5 THREE TYPES OF ADDRESSES

Some link-layer protocols define three types of addresses: unicast, multicast, and broadcast.

- **Unicast Address**

Each host or each interface of a router is assigned a unicast address. Unicasting means one-to-one communication. A frame with a unicast address destination is destined only for one entity in the link.

Example 1: The unicast link-layer addresses in the most common LAN, Ethernet, are 48 bits (six bytes) that are presented as 12 hexadecimal digits separated by colons; for example, the following is a link-layer address of a computer

A3:34:45:11:92:F1

- **Multicast Address**

Some link-layer protocols define multicast addresses. Multicasting means one-to-many communication. However, the jurisdiction is local (inside the link).

Example 2: The multicast link-layer addresses in the most common LAN, Ethernet, are 48 bits (six bytes) that are presented as 12 hexadecimal digits separated by colons. The second digit, however, needs to be an even number in hexadecimal. The following shows a multicast address:

A2:34:45:11:92:F1

- **Broadcast Address**

Some link-layer protocols define a broadcast address. Broadcasting means one-to-all communication. A frame with a destination broadcast address is sent to all entities in the link.

Example 3: The broadcast link-layer addresses in the most common LAN, Ethernet, are 48 bits, all 1s, that are presented as 12 hexadecimal digits separated by colons. The following shows a broadcast address:

FF:FF:FF:FF:FF:FF

7.6 ADDRESS RESOLUTION PROTOCOL (ARP)

Every time a node in a link needs to send an IP datagram to another node, it knows the IP address of the recipient node. The originating host is aware of the default router's IP address. The IP address of the next router is obtained by each router along the path, save for the last one, using its forwarding table. The last router is aware of the destination host's IP address. To move a frame through a link, we require the link-layer address of the next node; the following node's IP address is useless. The Address Resolution Protocol (ARP) comes in handy at this point. One of the auxiliary protocols defined in the network layer is the ARP protocol, as depicted in Figure 6. It belongs to the network layer, but since it converts an IP address to a logical-link address, we talk about it in this chapter. An

IP address from the IP protocol is accepted by ARP, which then converts it into the relevant link-layer address and sends it to the data-link layer.

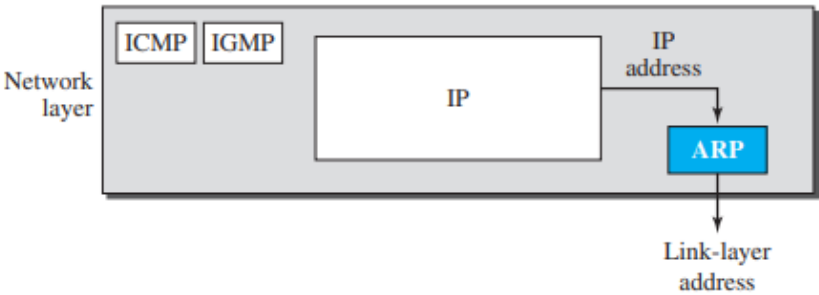


Figure 6: Position of ARP in TCP/IP protocol suite

ARP request packets are sent whenever a host or router needs to determine the link-layer address of another host or router in its network. The link-layer and IP addresses of the sender and the receiver are included in the packet. The query is broadcast over the link using the link-layer broadcast address because the sender does not know the receiver's link-layer address; we go over each protocol's link-layer broadcast address later (see Figure 7).

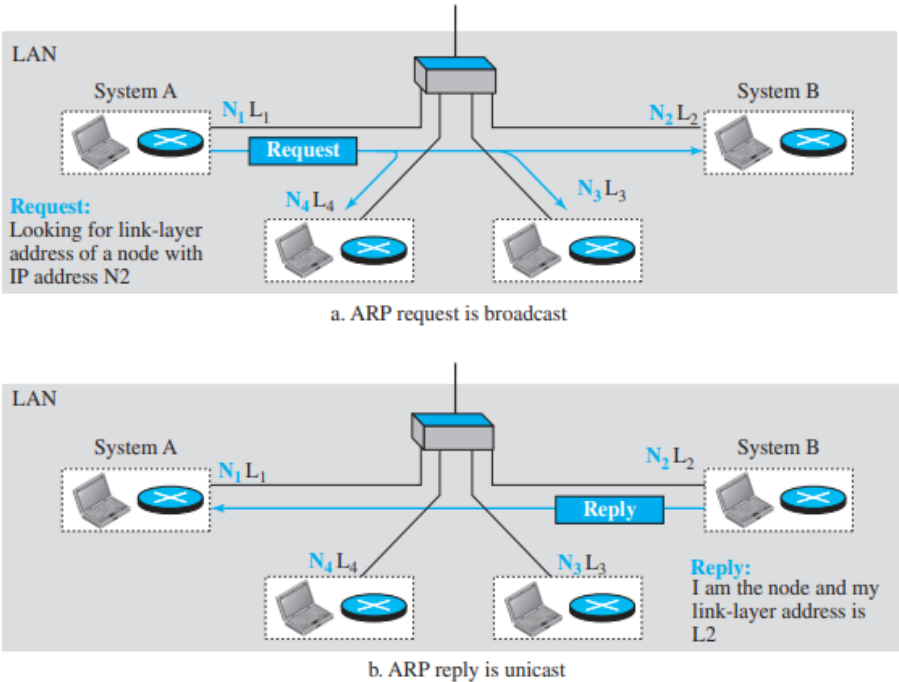


Figure 7: ARP operation

Only the intended recipient can identify its IP address and respond with an ARP response packet once every host or router on the network has received and processed the ARP request packet. The IP and link-layer addresses of the recipient are included in the response packet. The request packet's source node receives the packet through unicast.

In Figure 7a, the system on the left (A) needs to send a packet to the system on the right (B) with the IP address N2. System A does not know the

recipient's physical address; therefore, it must transfer the packet to its data-link layer for real delivery. By instructing the ARP protocol to send a broadcast ARP request packet to enquire about the physical address of a system with an IP address of N2, it makes advantage of the services of ARP.

Every system on the physical network receives this packet, but only system B will respond, as seen in Figure 7b. The physical address of System B is included in the ARP reply packet that it sends. System A can now use the physical address it was given to send all of the packets it has for this destination.

7.7 ERROR DETECTION AND CORRECTION: INTRODUCTION

Networks must be capable of accurately transferring data from one device to another. A system must ensure that the data received and those transferred are same for the majority of applications. Data might become corrupted while being sent from one node to the next. A message's components can change due to a variety of events. A technique for mistake detection and correction is necessary for some applications.

Some algorithms are tolerant of very little errors. For instance, random errors in audio or video transmissions may be acceptable, but we anticipate very high levels of accuracy when we communicate text.

Before moving on to additional nodes at the data-link layer, a frame that was corrupted between two nodes needs to be fixed. However, the majority of link-layer protocols merely toss the frame aside and let the upper-layer protocols handle the frame's retransmission. Nonetheless, some multimedia programmes attempt to fix the damaged frame.

7.8 TYPES OF ERRORS

Every time bits go from one point to another; interference might produce unexpected changes. The signal's form may alter as a result of this interference. When a single bit of a given data unit (such a byte, character, or packet) is altered from 1 to 0 or from 0 to 1, it is referred to as a single-bit error. When two or more bits in a data unit change from one to zero or from zero to one, it is referred to as a burst error. The impact of a single-bit and a burst error on a data unit is depicted in Figure 8.

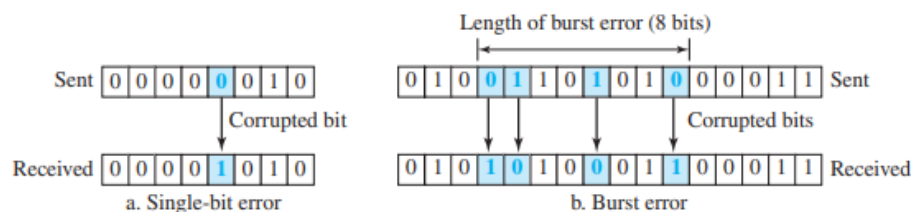


Figure 8: Single-bit and burst error

Because the noise signal typically lasts longer than one bit, which implies that when noise affects data, it impacts a set of bits, burst errors are more likely to happen than single-bit errors. The number of bits that are impacted depends on the noise duration and data rate. For instance, a noise of 1/100 second can damage 10 bits of data sent at 1 kbps, while the same noise can affect 10,000 bits of data sent at 1 Mbps.

7.9 REDUNDANCY

Redundancy is the key idea in error detection or correction. We need to provide a few extra bits along with our data in order to be able to recognise or fix problems. The sender adds and the recipient subtracts these extraneous bits. They enable the receiver to identify or fix damaged bits.

7.10 DETECTION VERSUS CORRECTION

Error rectification is more challenging than error detection. In error detection, our primary concern is determining whether an error has actually happened. Simple yes or no answers are provided. The quantity of corrupted bits doesn't even interest us. For us, a burst error is the same as a single-bit error. We need to know the precise number of corrupted bits and, more critically, where they are located inside the message in order to do error correction. Important considerations are the quantity of errors and the size of the message. When correcting a single error in an 8-bit data unit, we must take into account eight potential error locations; when correcting two errors in the same size data unit, we must take into account 28 (permutation of 8 by 2) possibilities. You may imagine how challenging it would be for the receiver to detect 10 faults in a data unit with 1000 bits.

7.11 SUMMARY

Many hosts, networks, and connecting elements like routers make up the Internet. The networks are referred to as links, while the hosts and associated hardware are known as nodes. A packet should travel through a set of nodes and links known as a path in the Internet from a source host to a destination host.

A frame must be created and delivered to another node along the link by the data-link layer. Along the link, it is in charge of packetizing (framing), flow control, error control, and congestion control. A frame is delivered from one node to the next via two data-link layers at the two endpoints of a link.

We require two different types of addressing, just like with any delivery between a source and destination where there are numerous pathways. The link-layer addressing specifies the addresses of the nodes that the packet should pass through, whereas the end-to-end addressing specifies the source and destination. The Address Resolution Protocol (ARP) was developed to map an IP address to its associated link-layer address in order to avoid including the link-layer addresses of all of these nodes in the frame. The forwarding table determines the IP address of the subsequent node and ARP

determines its link-layer address when a packet is at one node and ready to be transmitted to the next.

Data corruption during transmission is possible. For some applications, error detection and correction are necessary. Only one bit in the data unit has altered in a single-bit mistake. When there is a burst error, the data unit's two or more bits have changed. We must send extra (redundant) bits with data in order to detect or correct problems. Forward error correction and rectification by retransmission are the two basic types of error correction.

7.12 LIST OF REFERENCES

- 1] Data Communications and Networking, Behrouz A. Forouzan, Fifth Edition, TMH, 2013.
- 2] Computer Network, Andrew S. Tanenbaum, David J. Wetherall, Fifth Edition, Pearson Education, 2011.
- 3] Computer Network, Bhushan Trivedi, Oxford University Press.
- 4] Data and Computer Communication, William Stallings, PHI.

7.13 UNIT END EXERCISES

- 1] What is Nodes and Links?
- 2] Explain the services of data link layer.
- 3] Describe the two Sub-layers.
- 4] Write a note on three types of addresses.
- 5] Describe the Address Resolution Protocol (ARP).
- 6] What is Error Detection and Correction?
- 7] Discuss the types of errors.
- 8] Explain Redundancy.
- 9] Describe Detection versus Correction.

MEDIA ACCESS CONTROL (MAC)

Unit Structure:

- 8.0 Objectives
- 8.1 Introduction
- 8.2 Random access
- 8.3 CSMA
- 8.4 CSMA/CD
- 8.5 CSMA/CA
- 8.6 Controlled access
 - 8.6.1 Reservation
 - 8.6.2 Polling
 - 8.6.3 Token Passing
- 8.7 Channelization
 - 8.7.1 FDMA
 - 8.7.2 TDMA
 - 8.7.3 CDMA

Summary

List of References

Unit End Exercises

8.0 OBJECTIVES

- To understand the taxonomy and topologies involved in media access control protocols
- To get familiar with the fundamentals of control access
- To get acquainted with the time, code and frequency division schemas associated with channelization

8.1 INTRODUCTION

We require a multiple-access protocol to manage access to the link when nodes or stations are linked and utilise a common link, often known as a multipoint or broadcast link. The issue of regulating access to the media is comparable to the guidelines for public speaking. The protocols make sure that everyone's right to speak is respected and that nobody speaks twice, interrupts anyone else, dominates the conversation, etc. In order to manage access to a shared link, numerous protocols have been developed. All of

these protocols are a part of the media access control sublayer of the data-link layer (MAC). They are divided into three groups, as depicted in Figure 1.

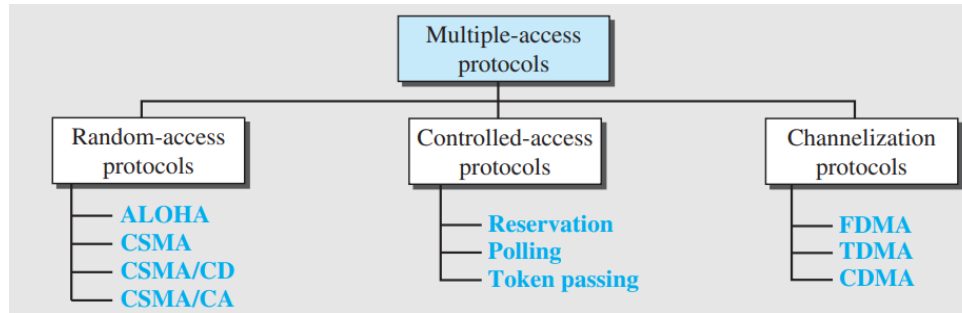


Figure 1: Taxonomy of multiple-access protocols

8.2 RANDOM ACCESS

No station has superiority over another station in random-access or contention methods, and no station is given control over another. Each time, a station with data to send decides whether or not to send using a process outlined in the protocol. The status of the medium will determine this choice (idle or busy). In other words, each station can transmit whenever it wants as long as it adheres to the established protocol, which includes determining the state of the medium.

Its name is a combination of two characteristics. The first issue is that there is no set time for a station to transmit. The stations' transmission is distributed at random. These techniques are known as random access because of this. Second, no station is designated as the next to send in the rules. The availability to the medium is a point of competition between stations. Because of this, these techniques are often known as contention approaches.

Each station in a random-access method is free to use the medium without interference from any other stations. However, there will be an access conflict (collision) and the frames will either be lost or altered if more than one station tries to communicate. Each station adheres to a process that responds to the following questions in order to prevent access conflicts or to resolve them when they do arise:

- Can the station access the media at any time?
- What actions can the station take if the medium is crowded?
- How would the station know whether the transmission was successful or unsuccessful?
- What options does the station have in the event of an access dispute?

The random-access techniques we examine in this chapter developed from the fascinating ALOHA protocol, which employed a relatively straightforward technique called multiple access (MA). The mechanism that makes the station feel the medium before transmitting was added to the

method to make it better. Carrier sense multiple access was used for this (CSMA). Later, this approach split into two parallel approaches: carrier sense multiple access with collision detection (CSMA/CD), which directs the station's response when a collision is detected, and carrier sense multiple access with collision avoidance (CSMA/CA), which makes an effort to avoid the collision.

8.3 CSMA

The CSMA method was created to reduce the possibility of collision and, as a result, improve performance. If a station perceives the medium first before attempting to use it, the likelihood of collision can be decreased. Each station must listen to the medium (or check the condition of the media) before sending using carrier sensing multiple access (CSMA). In other words, the foundation of CSMA is the idea of "listening before talking" or "sensing before transmitting."

While CSMA can lessen the likelihood of collision, it cannot completely prevent it. Figure 2, which depicts a space and time model of a CSMA network, explains why this is the case. A common channel connects stations (usually a dedicated medium).

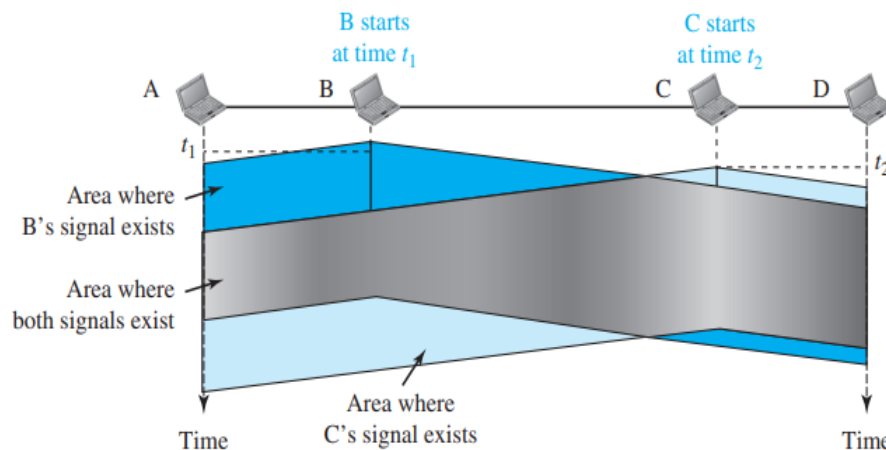


Figure 2: Space/time model of a collision in CSMA

Due to propagation delay, when a station sends a frame, it still requires some time (even though it is very little) for the initial bit to reach every station and for every station to detect it. In other words, a station could detect the medium and discover it to be empty just because it hasn't yet picked up the first bit delivered by another station.

At time t_1 , station B senses the medium and finds it idle, so it sends a frame. At time t_2 ($t_2 > t_1$), station C senses the medium and finds it idle because, at this time, the first bits from station B have not reached station C. Station C also sends a frame. The two signals collide and both frames are destroyed.

Vulnerable time

The propagation time T_p is when CSMA is most vulnerable. A signal needs this amount of time to go from one end of the medium to the other. A collision occurs when a station sends a frame and any other station tries to send a frame at the same time. However, if the frame's initial bit surpasses the limit of the medium, every station will have already heard it and will stop sending. The worst case is depicted in Figure 3. A frame is sent from the leftmost station, A, at time t_1 , and it reaches the rightmost station, D, at time $t_1 + T_p$. The grey region displays the area that is most vulnerable in both time and space.

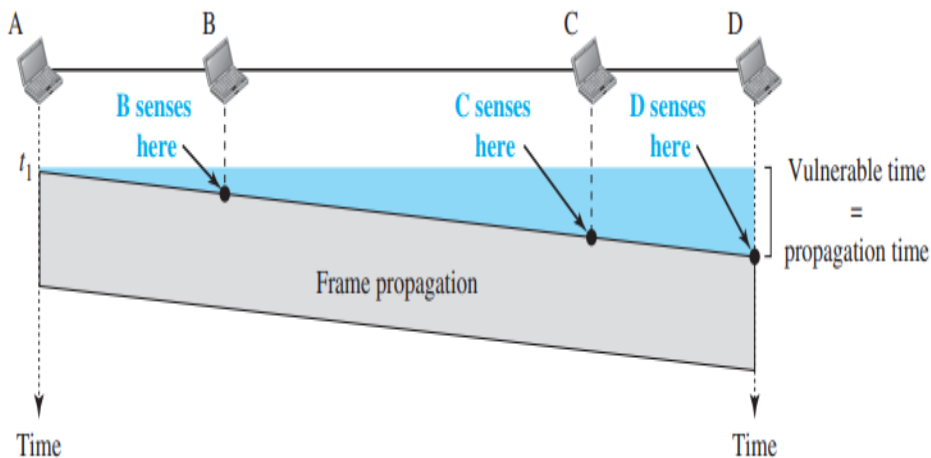


Figure 3: Vulnerable time in CSMA

Persistence methods

What ought a station to do in a busy channel? When the channel is empty, what should a station do? The 1-persistent approach, the nonpersistent method, and the p-persistent method have all been developed to provide solutions to these problems. Figure 4 depicts how three persistence techniques behave when a station detects a congested channel.

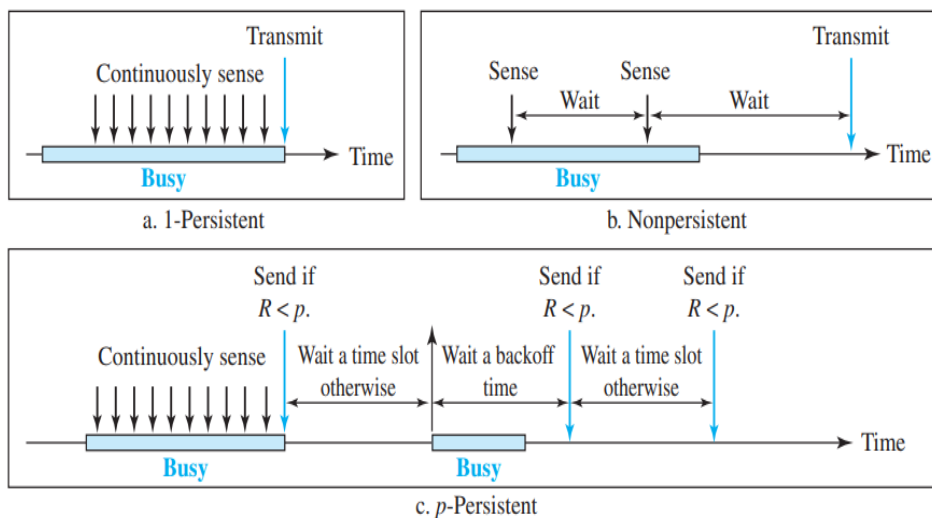


Figure 4: Behaviour of three persistence methods

Figure 5 shows the flow diagrams for these methods.

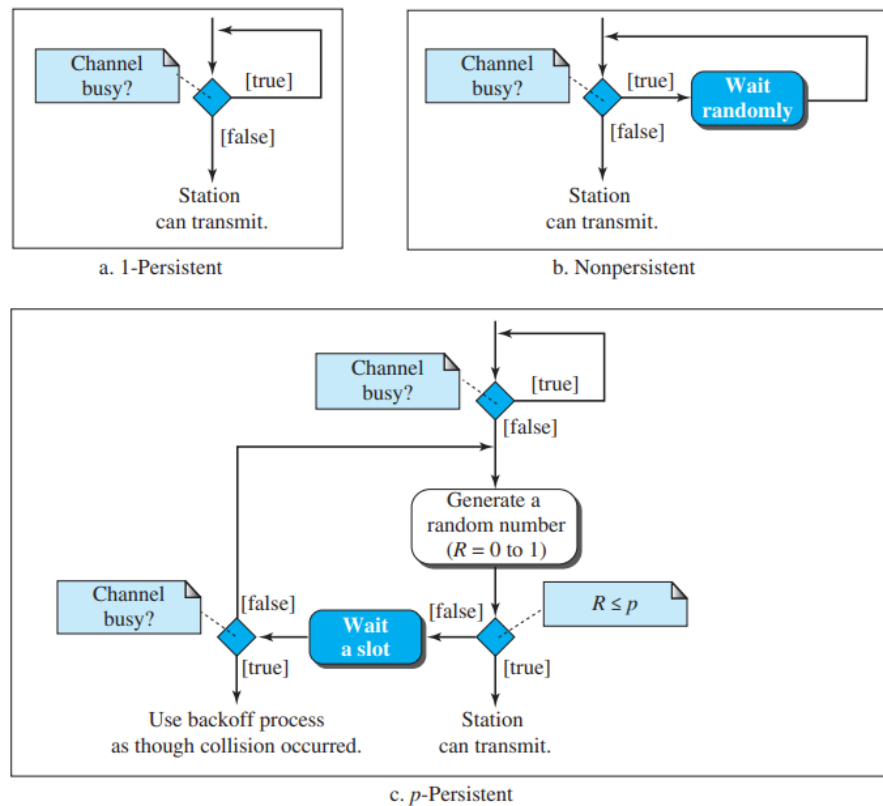


Figure 5: Flow diagram for three persistence methods

1-Persistent

The 1-persistent approach is easy to understand and use. With this technique, as soon as the station detects an empty line, it sends its frame (with probability 1). Due to the possibility of two or more stations finding the line idle and sending their frames right away, this strategy has the highest likelihood of a collision. Later, as we shall see, Ethernet makes use of this technique.

Nonpersistent

In the nonpersistent technique, the line is sensed by a station that has a frame to send. If the line is not in use, it dispatches right away. If the line is not idle, it waits an arbitrary period of time before sensing the line once more. Because it is rare that two or more stations will wait the same amount of time and retry to send at the same time, the nonpersistent technique lowers the likelihood of collision. However, because the medium is idle while there may be stations with frames to convey, this strategy lowers the network's efficiency.

p-persistent

If the channel has time slots with slot durations equal to or more than the maximum propagation time, the p-persistent technique is employed. The benefits of the other two systems are combined in the p-persistent strategy. It increases effectiveness and lowers the likelihood of contact. The following procedures are taken in this technique after the station determines the line is empty.

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.

8.4 CSMA/CD

The steps to be taken after a collision are not specified by the CSMA technique. The algorithm is improved by carrier sense multiple access with collision detection (CSMA/CD) to handle the collision.

In this technique, a station sends a frame and then checks the medium to verify if the transmission was successful. In that case, the station is complete. The frame is transmitted again if there is a collision, though.

Let's examine the first bits sent by the two stations involved in the collision to better comprehend CSMA/CD. We demonstrate what occurs when the first bits clash, despite the fact that each station keeps sending bits until it notices the collision. The collision in Figure 6 involves stations A and C.

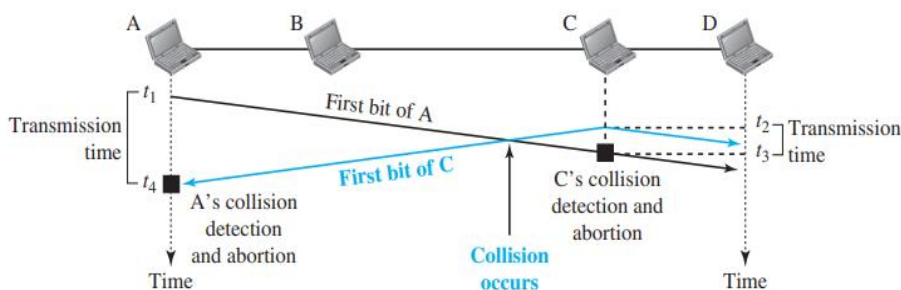


Figure 6: Collision of the first bits in CSMA/CD

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. Looking at the figure, we see that A

transmits for the duration $t_4 - t_1$; C transmits for the duration $t_3 - t_2$. Now that we know the time durations for the two transmissions, we can show a more complete graph in Figure 7.

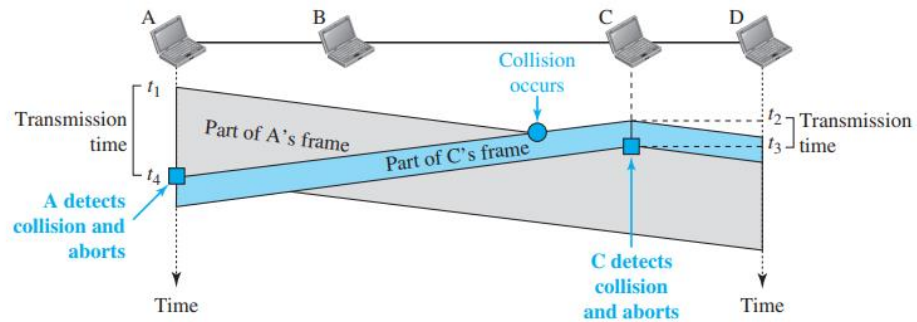


Figure 7: Collision and abortion in CSMA/CD

Minimum frame size

We require a frame size constraint for CSMA/CD to function. The sending station must identify any collisions and stop the transmission before sending the final bit of the frame. This is true because the station does not retain a copy of the frame after the complete frame has been transmitted and does not check the line for collision detection. Therefore, the frame transmission time T_{fr} must be at least two times the maximum propagation time T_p . To understand the reason, let us think about the worst-case scenario. If the two stations involved in a collision are the maximum distance apart, the signal from the first takes time T_p to reach the second, and the effect of the collision takes another time T_p to reach the first. So, the requirement is that the first station must still be transmitting after $2T_p$.

Energy level

We can say that a channel's energy level can be either zero, normal, or abnormal. The channel is inactive at zero. A station has successfully caught the channel and is delivering its frame at the normal level. When there is a collision at the abnormal level, the energy level is double what it is at the usual level. The energy level must be monitored by a station that has a frame to send or is sending a frame in order to know if the channel is free, busy, or in collision mode. Figure 8 depicts the current state.

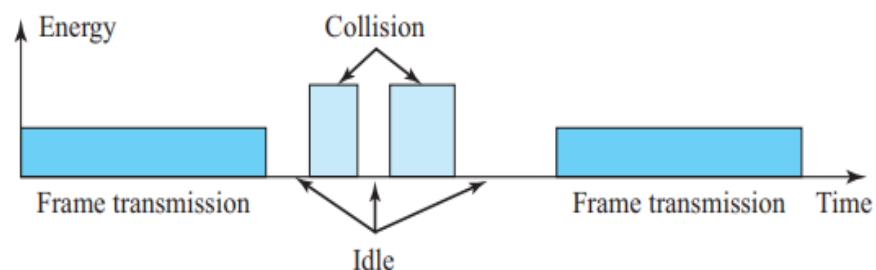


Figure 8: Energy level during transmission, idleness, or collision

Throughput

CSMA/CD has a higher throughput than pure or slotted ALOHA. The maximum throughput depends on the persistence method and the value of p in the p -persistent approach and occurs at a distinct value of G . When $G = 1$, the 1-persistent method's maximum throughput is around 50%. When G is between 3 and 8, the nonpersistent method's maximum throughput can increase to 90%.

Traditional Ethernet

One of the LAN protocols that used CSMA/CD is the traditional Ethernet with the data rate of 10 Mbps. the traditional Ethernet was a broadcast LAN that used the 1-persistence method to control access to the common media. Later versions of Ethernet try to move from CSMA/CD access methods.

8.5 CSMA/CA

For wireless networks, carrier sense multiple access with collision avoidance (CSMA/CA) was developed. Figure 9 illustrates the three CSMA/CA techniques for preventing collisions: interframe space, contention window, and acknowledgments.

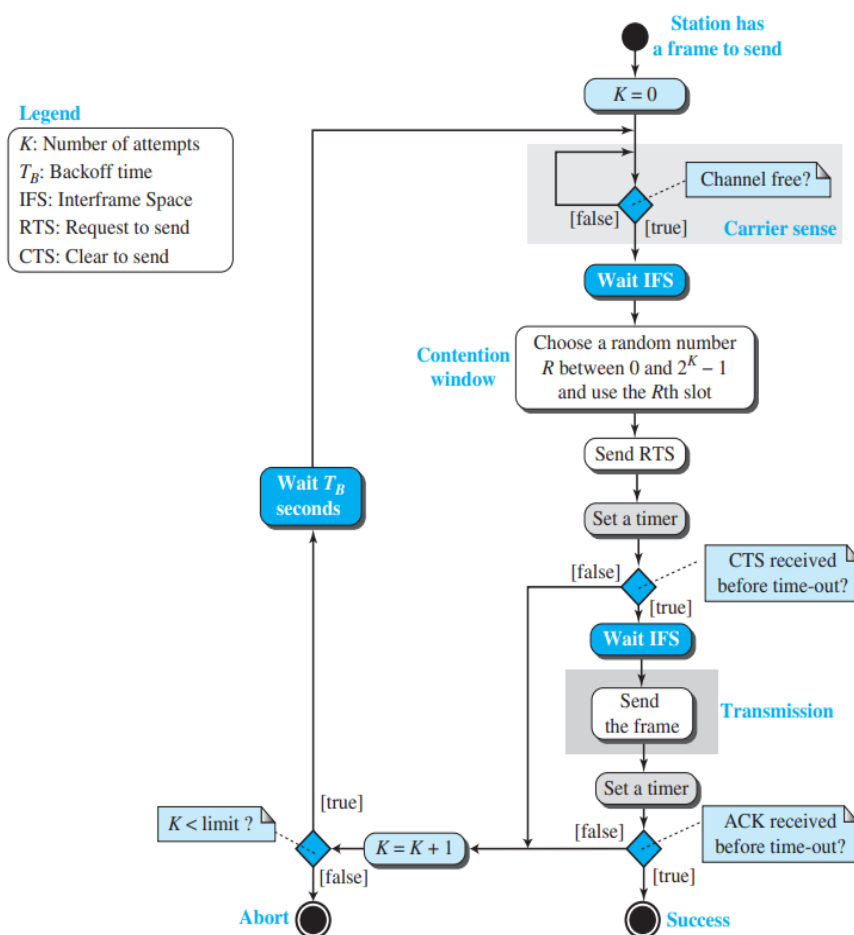


Figure 9: Flow diagram of CSMA/CA

- **Interframe space:** First, transmission is delayed even if the channel is judged to be idle in order to prevent accidents. The station does not immediately send when it finds an unoccupied channel. The interframe space, or IFS, is the time that it waits. Even if the channel could be detected as being empty when it actually has a distant station that has already begun transmitting. The signal from the distant station has not yet reached this station. The IFS time enables this station to receive the sent signal from the distant station's front. The station can send if the channel is still open after waiting an IFS period, but it must first wait for a time equal to the contention window. Stations or frame kinds can also be given higher priority using the IFS option. A station with a shorter IFS, for instance, has a higher priority.
- **Contention window:** A period of time with slots is known as the contention window. A station selects a random number of slots as its wait time when it is prepared to send. The binary exponential backoff approach modifies the number of slots in the window. Accordingly, it starts out at one slot and doubles every time the station is unable to find an idle channel after the IFS duration. The only difference between this and the p-persistent method is that a random outcome determines how many slots the waiting station will occupy. The station must sense the channel after each time slot, which is an intriguing aspect of the contention window. The process, however, is not restarted if the station determines that the channel is busy; instead, it simply pauses the timer and starts it again when the channel is determined to be empty. The station with the longest wait time is given priority as a result. Refer figure 10.

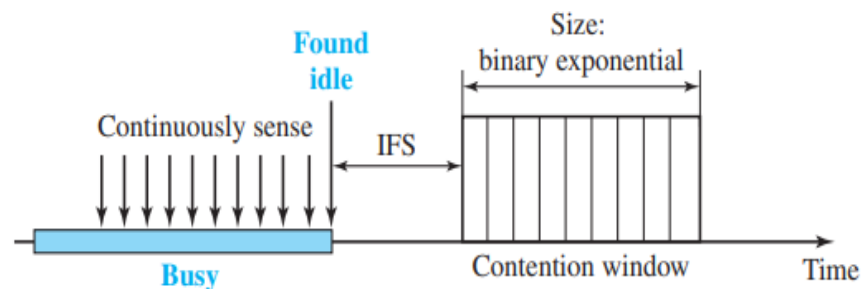


Figure 10: Contention window

- **Acknowledgement:** Even with all these safeguards, a collision that destroys data is still possible. Additionally, the transmission of the data could result in data corruption. The timer and positive acknowledgment can help confirm that the receiver has received the frame.

8.6 CONTROLLED ACCESS

To determine which station has the authority to send in controlled access, the stations confer with one another. A station cannot send without the consent of other stations. We'll go over three controlled-access techniques.

8.6.1 Reservation

A station using the reservation method must make a reservation before providing data. Intervals are used to divide time. The data frames sent in each interval are preceded by a reservation frame.

There are precisely N reservation micro slots in the reservation frame if there are N stations in the system. A station is responsible for each tiny slot. A station reserves its own tiny slot for when it wants to send a data frame. After the reservation frame, the stations that have made reservations can send their data frames.

Five stations and a reservation frame for five mini slots are shown in Figure 11. Only stations 1, 3, and 4 have reservations for the first interval. Only station 1 has made a reservation for the second intermission.

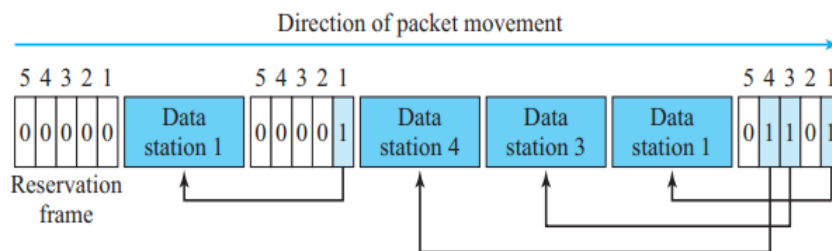


Figure 11: Reservation access method

8.6.2 Polling

When one device is designated as the primary station and the other devices are secondary stations, polling can be used. Even if a secondary device is the final destination, all data transfers must occur through the primary device. The secondary devices execute commands from the primary device, which controls the link. Which device is permitted to use the channel at any one time will depend on the primary device. As a result, a session always begins with the primary device (see Figure 12). To avoid collisions, this approach makes use of the poll and select functions. The disadvantage is that the system fails if the primary station does.

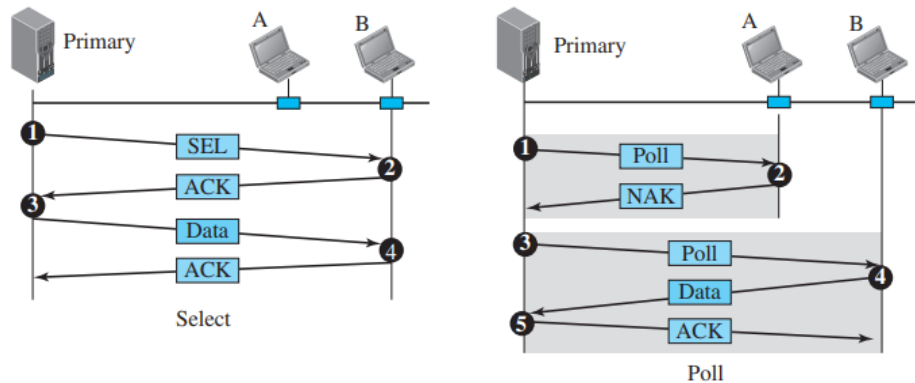


Figure 12: Select and poll functions in polling-access method

- Select:** Every time the primary device has something to send, the select function is used. Keep in mind that the main determines the link. The main is aware that the link is active if it is neither sending nor receiving data. The main device sends anything it has to send. It is unaware of the target device's readiness for reception, though. As a result, the primary must inform the secondary of an impending transmission and wait for the secondary to acknowledge that it is ready. The address of the proposed secondary is contained in one field of the select (SEL) frame that the primary creates and transmits prior to delivering data.
- Poll:** The primary device asks the subsidiary devices for communications using the poll function. When the primary is ready to receive data, each device must be individually queried (called a poll) to see if it has anything to send. When the first secondary is contacted, it either responds with a NAK frame or data (in the form of a data frame) depending on whether it has anything to relay. In the event of a negative response (a NAK frame), the primary then polls each subsequent secondary in the same way until it locates one with information to deliver. When the response is affirmative, the primary reads the response (a data frame) and sends back an acknowledgment (an ACK frame), confirming receipt.

8.6.3 Token Passing

The stations in a network are arranged in a logical ring when using the token-passing approach. In other words, there is a predecessor and a successor for each station. The station that is logically before the station in the ring is the predecessor, and the station that is after the station in the ring is the successor. The station that is currently accessing the channel is the current station. The present station has received the privilege to this access from the predecessor. When the current station has no more data to send, the right will be transferred to the successor.

But how does one station grant another station access to a channel? In this procedure, a unique packet known as a token is passed around the ring. The station has the right to access the channel and send data as long as it has the token. A station waits till it receives the token from its predecessor before

sending any data. Holding the token now, it transmits its info. The token is passed to the following logical station in the ring when the station runs out of data to send. Until it gets the token again in the following round, the station is unable to communicate data. When a station in this process receives the token and is out of data to send, it simply passes the data on to the following station.

This access method needs token management. The amount of time stations may possess the token must be restricted. The token needs to be watched to make sure it hasn't vanished or been obliterated. For instance, the token will vanish from the network if a station holding it fails. Assigning priorities to the stations and the different sorts of data being transmitted is another aspect of token management. Finally, in order to ensure that low-priority stations release the token to high-priority stations, token management is required.

Logical ring

Stations in a token-passing network do not necessarily need to be physically arranged in a ring; they might be arranged logically instead. Four distinct physical topologies that can result in a logical ring are depicted in Figure 13.

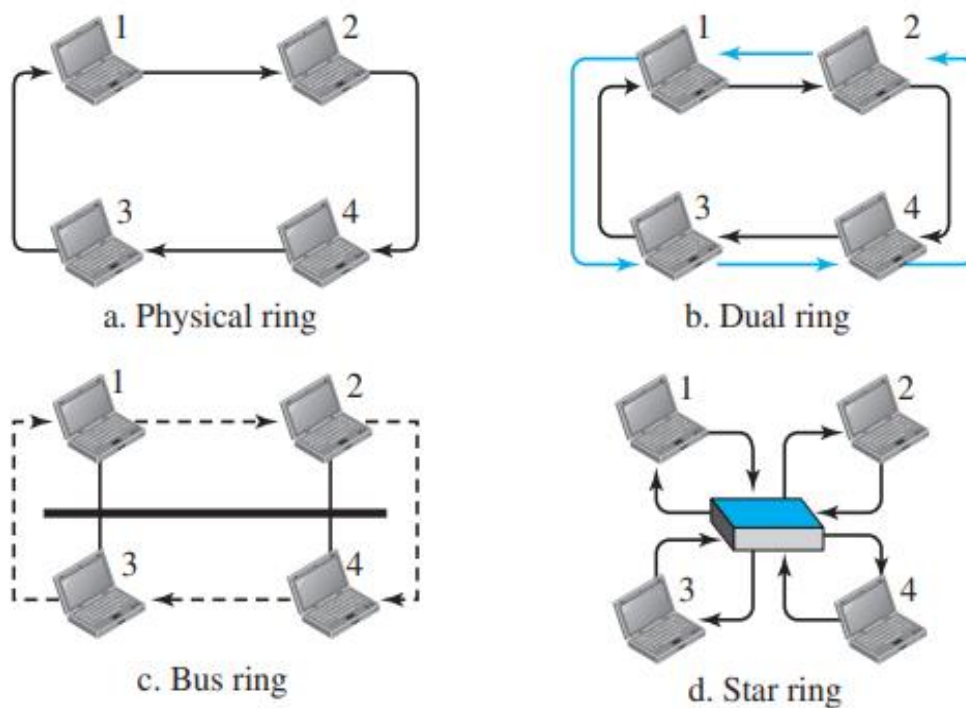


Figure 13: Logical ring and physical topology in token-passing access method

When a station transfers a token to its successor in the physical ring topology, the token cannot be observed by other stations since the successor is the next station in line. As a result, the token need not contain the address of the following successor. The issue with this topology is that the system

as a whole failure if one of the links the medium connecting two neighbouring stations fails.

A second (auxiliary) ring is used in the dual ring topology, and it rotates anticlockwise to the primary ring. Only emergencies (such as a spare tyre for an automobile) should use the second ring. The mechanism automatically merges the two rings to create a temporary ring if one of the main ring's links breaks. The auxiliary ring becomes idle once the broken link is repaired. Keep in mind that each station must have two transmitter ports and two reception ports for this architecture to function. This topology is used by the fast Token Ring networks FDDI (Fiber Distributed Data Interface) and CDDI (Copper Distributed Data Interface).

The stations are linked together by a single cable known as a bus in the bus ring topology, sometimes known as a token bus. However, because each station is aware of the address of its successor, they form a logical ring (and also predecessor for token management purposes). A station puts the address of its successor into the token after releasing it once it has finished delivering all of its data. The token for accessing shared material is only given to the station whose address matches the token's destination address. This topology is utilised by the IEEE-standard Token Bus LAN.

The physical topology in a star ring topology is a star. But there is a hub that serves as the link. The ring is created by the wiring inside the hub, and the stations are linked to this ring by the two wire connections. This design reduces the likelihood of network failure because if a connection fails, the hub will bypass it so that the remaining stations can continue to function. It is also simpler to add and remove stations from the ring. In the IBM Token Ring LAN, this topology is still in use. The physical topology in a star ring topology is a star. But there is a hub that serves as the link. The ring is created by the wiring inside the hub, and the stations are linked to this ring by the two wire connections. This design reduces the likelihood of network failure because if a connection fails, the hub will bypass it so that the remaining stations can continue to function. It is also simpler to add and remove stations from the ring. In the IBM Token Ring LAN, this topology is still in use.

8.7 CHANNELIZATION

A multiple-access technique known as channelization (or channel partition) divides the link's available bandwidth among many stations according to time, frequency, or by the use of a code. Three channelization protocols - FDMA, TDMA, and CDMA are covered in this section.

8.7.1 FDMA

In frequency-division multiple access (FDMA), frequency bands are used to split the available bandwidth. Each station is given a band to use for data transmission. In other words, each band is set aside for a particular station, and it always belongs to that station. A bandpass filter is also used by each station to limit the transmitter frequencies. The assigned bands are separated from one another by tiny guard bands to avoid station interference. Figure 14 depicts the FDMA concept.

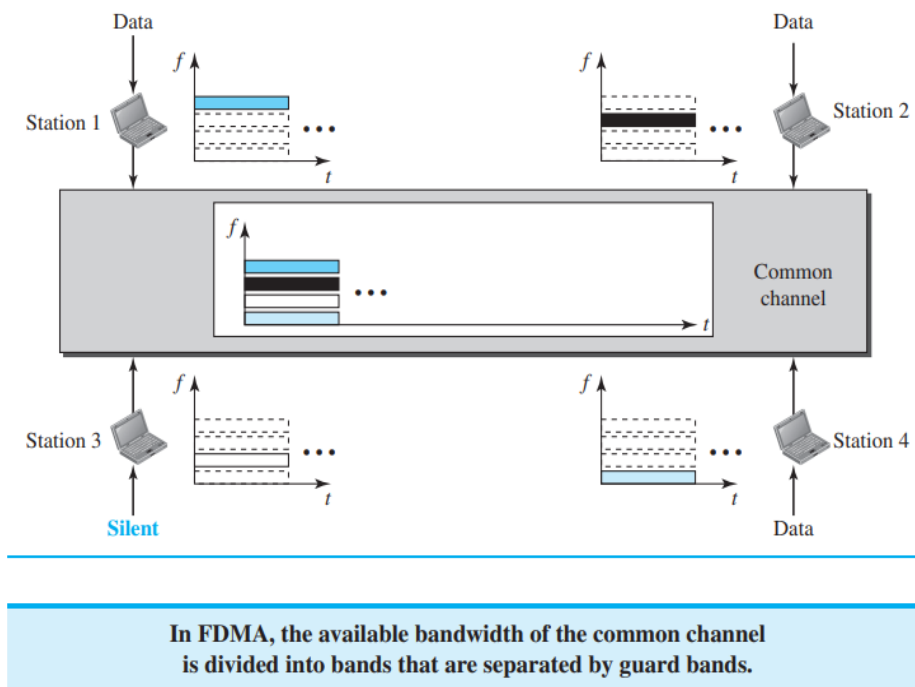


Figure 14: FDMA

A fixed frequency band is specified by FDMA for the entire transmission period. As a result, stream data—a continuous flow of data that may not be packetized—can be used with FDMA without much difficulty.

It is important to note that, despite their basic similarity, FDMA and frequency-division multiplexing (FDM) differ from one another. By combining the loads from low-bandwidth channels and transmitting them over a high-bandwidth channel, FDM is a physical layer approach. Low-pass filters are used to merge the channels. The signals are combined, modulated, and turned into a bandpass signal by the multiplexer. The multiplexer shifts each channel's bandwidth.

On the other hand, the data-link layer's access technique is FDMA. Each station's datalink layer instructs the physical layer to create a bandpass signal using the data that has been provided to it. The signal needs to be produced inside the designated band. At the physical layer, there is no physical multiplexer. Each station's broadcasts are automatically bandpass filtered. When they are sent to the common channel, they are mixed.

8.7.2 TDMA

The stations with time-division multiple access share the channel's capacity. A time slot is assigned to each station during which it may submit data. In the designated time slot, each station transmits its data. Figure 15 depicts the concept underlying TDMA.

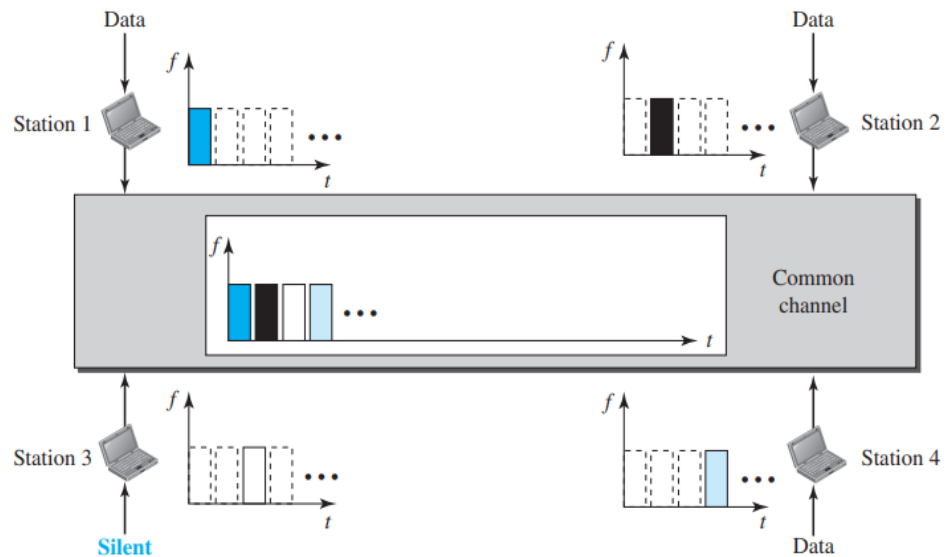


Figure 15: TDMA

The primary issue with TDMA is synchronisation between the various stations. Each station must be aware of the start and placement of their respective slots. If the stations are dispersed over a vast area, this could be challenging due to system propagation delays. We can introduce guard times to make up for the delays. A few synchronisation bits, also known as preamble bits, are often present at the start of each slot to facilitate synchronisation.

We must also stress that there are distinctions between TDMA and time-division multiplexing (TDM), despite the conceptual similarity between the two. Using a faster channel, TDM is a physical layer technology that mixes data from slower channels and sends it. The procedure interleaves data units from each channel using a hardware multiplexer.

On the other hand, the data-link layer's access technique is TDMA. Each station's data-link layer instructs its physical layer to use the designated time slot. At the physical layer, there is no physical multiplexer.

8.7.3 CDMA

Several decades ago, code-division multiple access (CDMA) was developed. It can now be implemented thanks to recent developments in electronic technology. In contrast to FDMA, just one channel uses the entire

link's bandwidth with CDMA. There is no timesharing, unlike TDMA, therefore all stations can deliver data concurrently.

Analogy

Let's start by using an analogy. Simply put, CDMA stands for communication via various codes. For instance, if no one else understands English in a large room full of people, two persons can speak in private in English. If they are the only ones who can communicate in Chinese, the next two persons can speak it, and so on. In other words, communication between numerous couples, though in different languages, can simply be facilitated by the common channel, in this example the room (codes).

Idea

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 16.

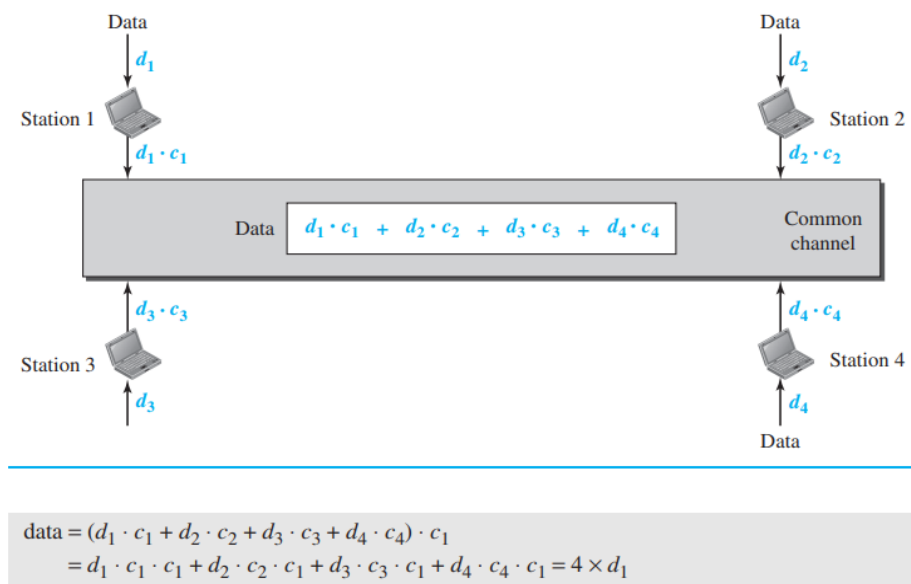


Figure 16: Simple idea of communication with code

Station 1 multiplies (a special kind of multiplication, as we will see) its data by its code to get $d_1 \cdot c_1$. Station 2 multiplies its data by its code to get $d_2 \cdot c_2$, and so on. The data that go on the channel are the sum of all these terms, as shown in the box. Any station that wants to receive data from one of the other three multiplies the data on the channel by the code of the sender. For

example, suppose stations 1 and 2 are talking to each other. Station 2 wants to hear what station 1 is saying. It multiplies the data on the channel by c_1 , the code of station 1. Because $(c_1 \cdot c_1)$ is 4, but $(c_2 \cdot c_1)$, $(c_3 \cdot c_1)$, and $(c_4 \cdot c_1)$ are all 0s, station 2 divides the result by 4 to get the data from station 1.

Chips

Based on coding theory, CDMA operates. As seen in Figure 17, each station is given a code, which is a series of integers known as chips.



Figure 17: Chip sequence

The codes apply to the first instance. The sequences were carefully chosen; they were not chosen at random. These sequences, known as orthogonal sequences, have the following characteristics:

1. Each sequence is made of N elements, where N is the number of stations.
2. If we multiply a sequence by a number, every element in the sequence is multiplied by that element. This is called multiplication of a sequence by a scalar.
3. If we multiply two equal sequences, element by element, and add the results, we get N , where N is the number of elements in each sequence. This is called the inner product of two equal sequences.
4. If we multiply two different sequences, element by element, and add the results, we get 0. This is called the inner product of two different sequences.
5. Adding two sequences means adding the corresponding elements. The result is another sequence.

Data representation

We encode using the following guidelines: A station encodes a 0 bit as -1 and a 1 bit as +1 depending on the number of bits it needs to communicate. A station sends no signal when it is not in use, which is read as a 0. They're displayed in Figure 18.

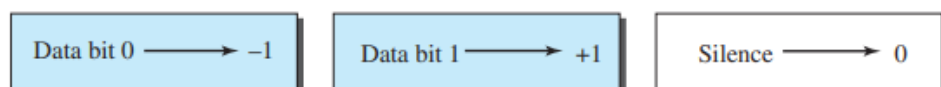


Figure 18: Data representation in CDMA

We provide a straightforward illustration of how four stations can share the link during a 1-bit interval. Repeating the process for more intervals is simple. Stations 1 and 2 are thought to be broadcasting a 0 bit, while channel 4 is thought to be sending a 1. Three is silent. The data are converted to 1, 1, 0, and +1 at the sender site. The orthogonal sequence, which is distinct for each station, is multiplied by the corresponding number at each station. A new sequence is produced and delivered to the channel as a result. We make the assumption for convenience that all stations send the generated sequences simultaneously. The four sequences we previously described are added together to create the sequence you see on the channel. Figure 19 illustrates this situation.

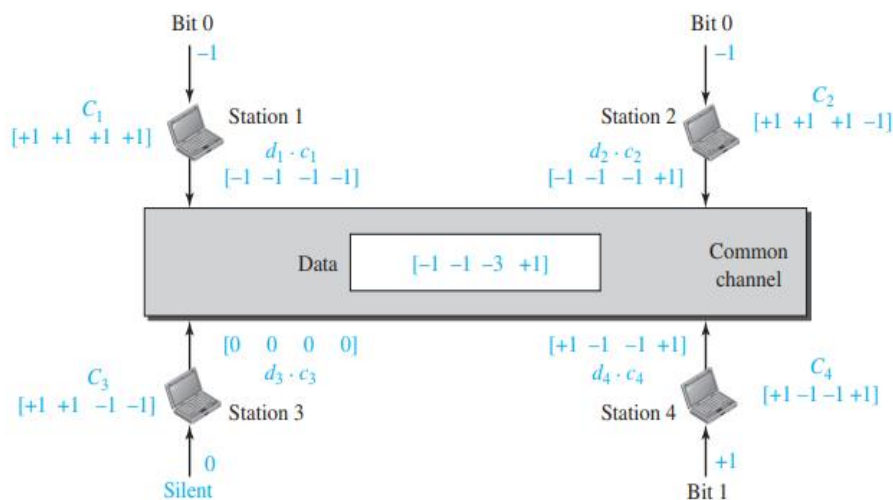


Figure 19: Sharing channel in CDMA

Signal level

If we display the digital signal generated by each station and the data recovered at the destination, the process will be easier to comprehend (see Figure 20). The signal on the common channel is shown in the figure together with the matching signals for each station (NRZ-L is used for simplicity).

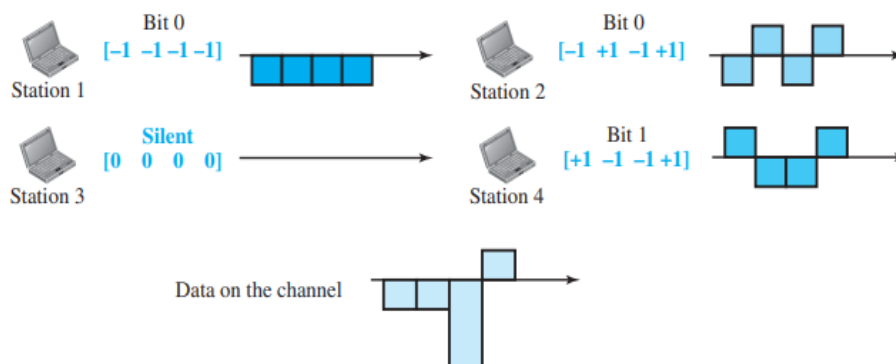


Figure 20: Digital signal created by four stations in CDMA

Figure 21 demonstrates how station 3 can use station 2's code to detect data delivered by that station. To create a new signal, the signal denoting station 2 chip code is multiplied (inner product operation) by the total amount of data on the channel. Following that, the station integrates and adds the signal's undersampling to obtain the number 4, which is then split by 4 and decoded as bit 0.

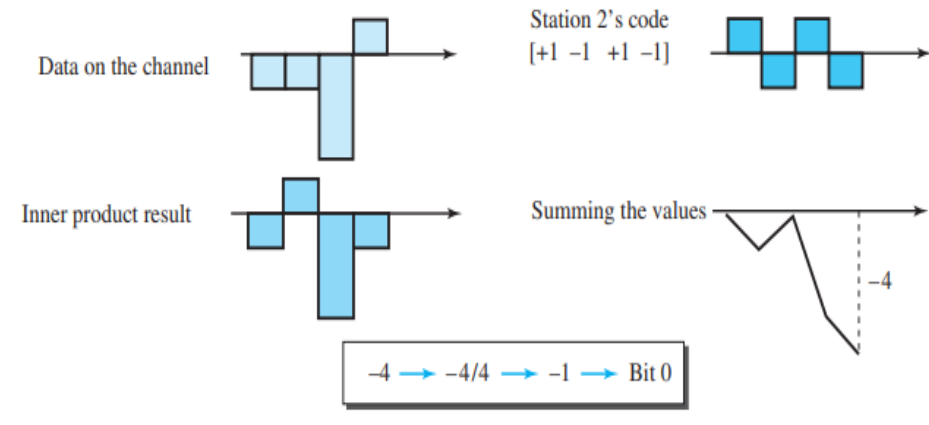


Figure 21: Decoding of the composite signal for one in CDMA

Sequence generation

To generate chip sequences, we use a Walsh table, which is a two-dimensional table with an equal number of rows and columns, as shown in Figure 22.

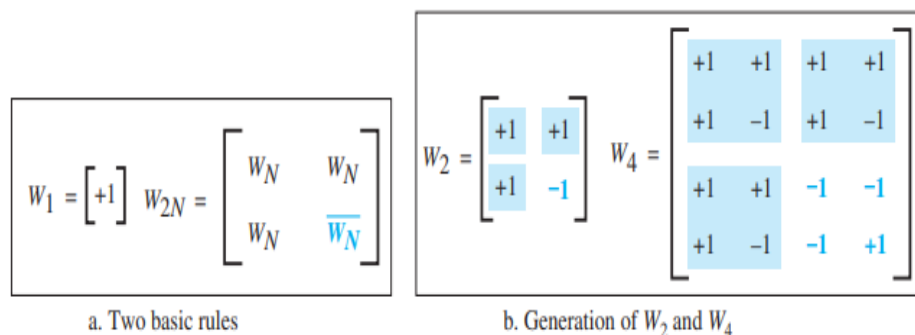


Figure 22: General rule and examples of creating Walsh tables

In the Walsh table, each row is a sequence of chips. W_1 for a one-chip sequence has one row and one column. We can choose -1 or $+1$ for the chip for this trivial table (we chose $+1$). According to Walsh, if we know the table for N sequences W_N , we can create the table for $2N$ sequences W_{2N} , as shown in Figure 22. The W_N with the overbar $\overline{W_N}$ stands for the complement of W_N , where each $+1$ is changed to -1 and vice versa. Figure 12.29 also shows how we can create W_2 and W_4 from W_1 . After we select W_1 , W_2 can be made from four W_1 s, with the last one the complement of W_1 . After W_2 is generated, W_4 can be made of four W_2 s, with the last one the complement of W_2 . Of course, W_8 is composed of four W_4 s, and so on.

Note that after W_N is made, each station is assigned a chip corresponding to a row.

Something we need to emphasize is that the number of sequences, N , needs to be a power of 2. In other words, we need to have $N = 2^m$.

SUMMARY

To manage access to a shared link, numerous formal protocols have been developed. They are divided into three categories: channelization protocols, restricted access protocols, and random-access protocols.

No station has superiority over another station in random access or contention methods, and no station is given control over another. Multiple access (MA) to the shared medium is permitted by ALOHA. This configuration raises the possibility of collisions. The CSMA method was created to reduce the possibility of collision and, as a result, improve performance. If a station perceives the medium first before attempting to use it, the likelihood of collision can be decreased. Each station must first listen to the media when using carrier sense multiple access (CSMA).

The CSMA technique is enhanced by carrier sense multiple access with collision detection (CSMA/CD) to handle collisions. In this technique, a station sends a frame and then checks the medium to verify if the transmission was successful. In that case, the station is complete. The frame is transmitted again if there is a collision, though. Carrier sensing multiple access with collision avoidance (CSMA/CA) was developed to prevent collisions on wireless networks. The interframe space, the contention window, and acknowledgments are three techniques for preventing collisions.

To determine which station has the authority to send in controlled access, the stations confer with one another. A station cannot send without the consent of other stations. We covered three widely used controlled-access techniques: token passing, polling, and reservations. A station must make a reservation before sending data using the reservation access technique. Intervals are used to divide time. The data frames sent in each interval are preceded by a reservation frame. In the polling approach, even when the secondary device is the final destination, all data exchanges must take place through the primary device. The secondary devices execute commands from the primary device, which controls the link. The stations in a network are arranged in a logical ring when using the token-passing approach. Tokens, a unique packet, are passed around the ring.

A multiple-access technique called channelization divides the link's available bandwidth among several stations according to time, frequency, or through the use of a code. FDMA, TDMA, and CDMA were the three channelization protocols that were covered. Frequency bands are used in

FDMA to split the available bandwidth. Each station is given a band to use for data transmission. In other words, each band is set aside for a particular station, and it always belongs to that station. The stations in TDMA split the channel's time-shared bandwidth. A time slot is assigned to each station during which it may submit data. In the designated time slot, each station transmits its data. In CDMA, the stations use different codes to achieve multiple access. CDMA is based on coding theory and uses sequences of numbers called chips. The sequences are generated using orthogonal codes such as the Walsh tables.

LIST OF REFERENCES

- 1] Data Communications and Networking, Behrouz A. Forouzan, Fifth Edition, TMH, 2013.
- 2] Computer Network, Andrew S. Tanenbaum, David J. Wetherall, Fifth Edition, Pearson Education, 2011.
- 3] Computer Network, Bhushan Trivedi, Oxford University Press.
- 4] Data and Computer Communication, William Stallings, PHI.

UNIT END EXERCISES

- 1] What do you mean by random access?
- 2] What is CSMA? Explain in detailed.
- 3] Explain the concept of CSMA/CD.
- 4] Illustrate the meaning of CSMA/CA with example.
- 5] Write a note on control access.
- 6] Write a note on Reservation.
- 7] Explain the concept of pooling.
- 8] Write a note on token passing.
- 9] Explain channelization.
- 10] What do you mean by FDMA. Illustrate the concept with the help of a diagram.
- 11] Write a detailed note on TDMA.
- 12] With the help of a diagram explain the concept of CDMA.

CONNECTING DEVICES AND VIRTUAL LANs

Unit Structure:

- 9.0 Objectives
- 9.1 Introduction
- 9.2 Connecting devices
 - 9.2.1 Hubs
 - 9.2.2 Link-Layer Switches
 - 9.2.3 Routers
- 9.3 Virtual LANs
 - 9.3.1 Membership
 - 9.3.2 Configuration
 - 9.3.3 Communication between Switches
 - 9.3.4 Advantages

Summary

List of References

Unit End Exercises

9.0 OBJECTIVES

- To understand the fundamentals of virtual LAN's and various connecting devices
- To acquaint with the concepts of hubs, routers and virtual LAN's along with the membership, configuration and communication between the switches

9.1 INTRODUCTION

LANs or hosts typically function together. They are linked to the Internet or to one another. We utilise connecting devices to link hosts or local area networks. The Internet model's connecting devices can function at various Internet layer depths. Following a discussion of a few connecting devices, we demonstrate how to use them to build virtual local area networks (VLANs).

The chapter is divided into two sections:

- The connected devices are covered in the first section. The features of hubs are first described. After that, the section talks about link-layer

switches (also known as switches) and demonstrates how connecting LANs with broadcast domains can result in loops.

- Virtual LANs, or VLANs, are discussed in the second section. The section begins by demonstrating how a VLAN's membership can be specified. The configuration of the VLAN is then covered in the section. The process of switching between VLANs is then demonstrated. The section concludes with a discussion of a VLAN's benefits.

9.2 CONNECTING DEVICES

Networks and hosts typically interact with one another. In order to create networks or the internet, we join networks together or hosts together using connecting devices. The Internet model's connecting devices can function at various Internet layer depths. We examine link-layer switches, routers, and hubs as three different categories of connecting devices. Today's hubs function in the Internet model's first layer. In the top two layers, link-layer switches function. In the first three layers, routers work. Refer figure 1.

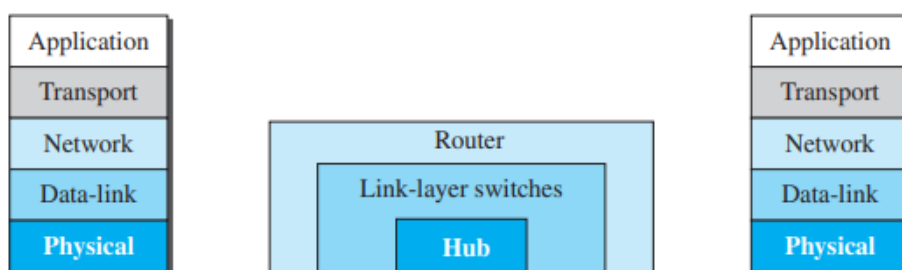


Figure 1: Three categories of connecting devices

9.2.1 Hubs

A hub is a piece of equipment that only works at the physical layer. Within a network, signals that contain information can only travel so far before attenuation compromises the integrity of the data. When a signal is received, a repeater regenerates and retimes the original bit pattern before it is too weak or corrupted. The signal is then sent by the repeater. When Ethernet LANs used bus topologies in the past, a repeater was used to connect two LAN segments in order to get around the coaxial cable's length limitation. But nowadays, Ethernet LANs employ star topology.

A repeater is a multiport device, frequently referred to as a hub, that may act as both a repeater and a connecting point in a star architecture. Figure 2 demonstrates how the hub transmits packets from all outgoing ports but the one from where the signal was received when a packet from station A to station B arrives at the hub. This is done to remove any potential corrupting noise from the signal representing the frame. The frame is aired, to put it another way. The frame is received by every station in the LAN, but only station B saves it. The other stations ignore it. The function of a repeater or hub in a switched LAN is depicted in Figure 2.

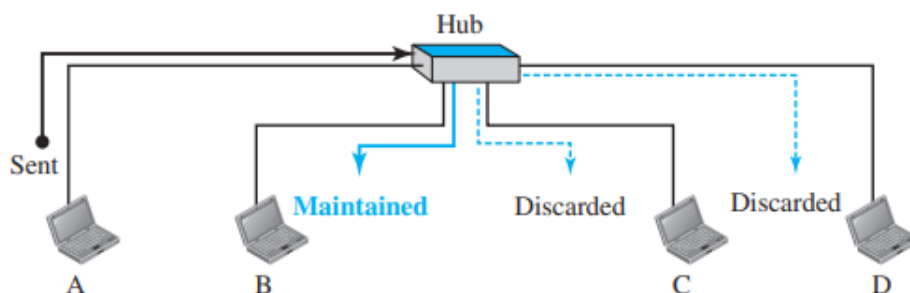


Figure 2: A hub

The figure clearly demonstrates that a hub lacks a filtering capacity and the intelligence to determine which port the frame should be sent out of.

A physical-layer device is a hub or repeater. They do not examine the link-layer address of the receiving frame and do not have a link-layer address of their own. They just regenerate the damaged bits and release them through each port.

9.2.2 Link-Layer Switches

A switch that functions in both the physical and data-link layers is referred to as a link-layer switch. It regenerates the signal it receives because it is a physical-layer device. The link-layer switch may examine the MAC addresses (source and destination) in the frame because it is a link-layer device.

Filtering

One can wonder what functions a link-layer switch and a hub perform differently. A link-layer switch has the capacity to filter. It has the ability to examine a frame's destination address and choose the appropriate outgoing port from which to send the frame.

Consider an example. In Figure 3, we have a LAN with four stations that are connected to a link-layer switch. If a frame destined for station 71:2B:13:45:61:42 arrives at port 1, the link-layer switch consults its table to find the departing port. According to its table, frames for 71:2B:13:45:61:42 should be sent out only through port 2; therefore, there is no need for forwarding the frame through other ports.

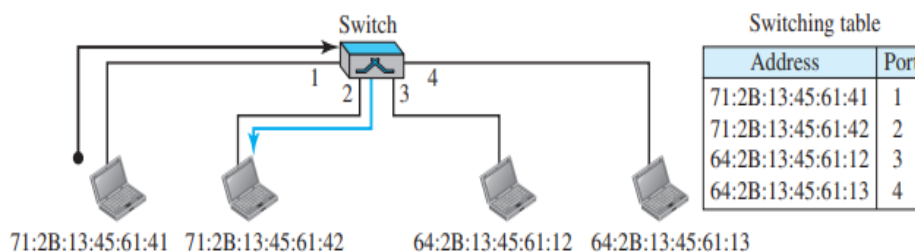


Figure 3: Link-layer switch

Transparent switches

Transparent switches are those in which the stations have no idea that the switch even exists. Reconfiguring the stations is not essential if a switch is added or removed from the system. The IEEE 802.1d specification states that a system using transparent switches must satisfy the following three requirements:

- From one station to another, frames must be sent.
- The forwarding table is created automatically by studying network frame motions.
- Loops in the system need to be avoided

Forwarding

A transparent switch must correctly forward the frames, as discussed in the previous section.

Learning

Switching tables on the first switches were static. During switch configuration, the system administrator would manually enter each table entry. Even though the procedure was straightforward, it was impractical. A station required to be manually added or removed from the database. If a station's MAC address changed, which is a common occurrence, the same held true. A new MAC address, for instance, results from installing a new network card.

A dynamic table that automatically translates addresses to ports (interfaces) is a superior alternative to the static database. We require a switch that gradually picks up on the movements of the frames in order to create a dynamic table. In order to accomplish this, the switch examines both the source and destination addresses of each frame that moves through it. The source address is used to add entries to the table and to update it, while the destination address is utilised to make forwarding decisions (table lookup). Let's use Figure 4 to explain this procedure.

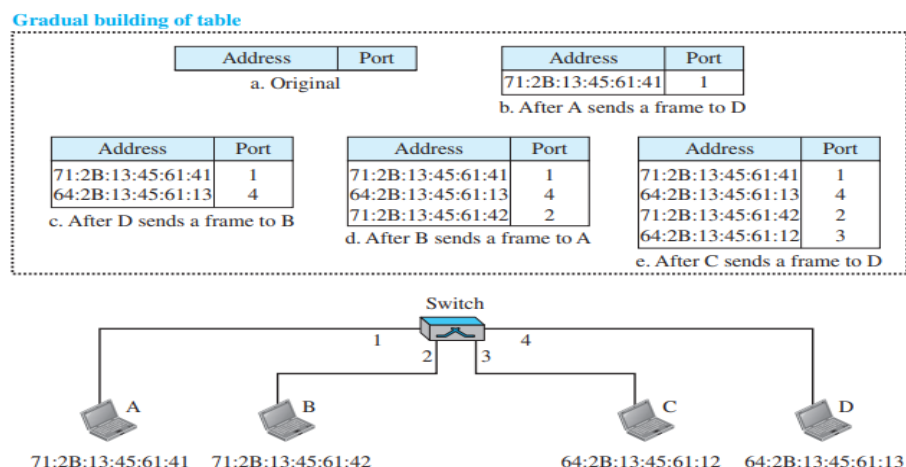


Figure 4: Learning switch

1. The switch does not include an entry for either station D or A when station A sends a frame to station D. The frame floods the network as it leaves all three ports open. However, the switch discovers that station A has to be linked to port 1 by examining the source address. This indicates that moving forward, frames to destination A must be sent via port 1. This row is added to the switch's table. Now there is a row in the table.
2. Because the switch does not have an entry for station B, it floods the network once more when station D transmits a frame to station B. It does, however, add one extra entry to the station D-related table.
3. The process of learning continues until every port's data is present in the table. However, keep in mind that learning may take a while. For instance, a station will never have an entry in the table if it never sends out a frame, which is a rare occurrence.

Loop problem

As long as there are no redundant switches present in the system, transparent switches function as intended. However, in order to increase system reliability, systems administrators prefer to have redundant switches (more than one switch between a pair of LANs). When a switch malfunctions, another one takes over until the malfunctioning switch is fixed or replaced. Redundancy has the potential to lead to loops, which is very undesirable. Only when two or more broadcasting LANs (such as those employing hubs) are connected by more than one switch may loops be formed.

A very basic example of a loop that can be made in a system with two LANs connected by two switches is shown in Figure 5.

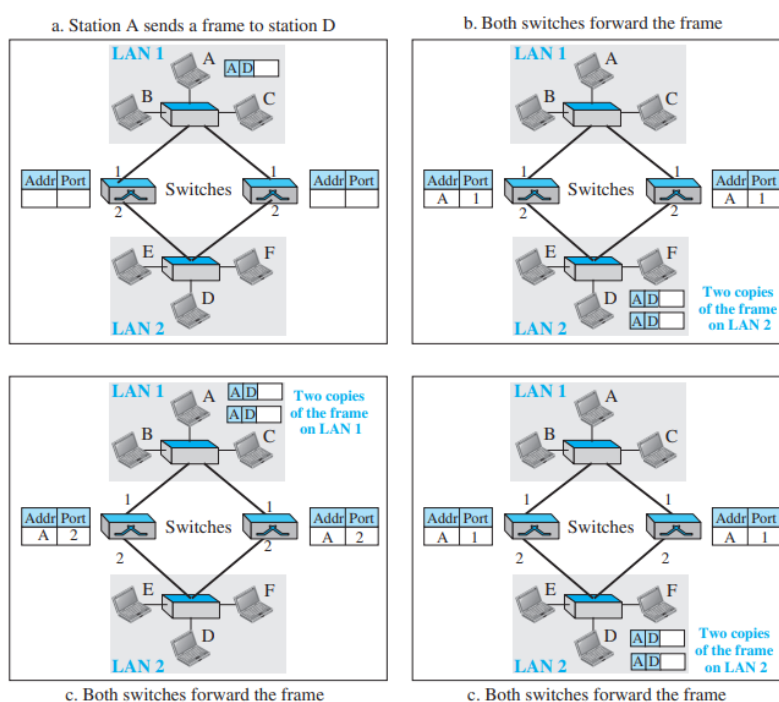


Figure 5: Loop problem in a learning switch

1. Station A transmits station D a frame. Both switches' tables are empty. On the basis of the source address A, both forward the frame and update their tables.
2. The frame is now duplicated twice on LAN 2. The copy that the left switch sent out is received by the right switch, which forwards the frame without knowing the destination address D. The left switch receives the copy that the right switch sent out and sends it because it has no information about D. The fact that switches use an access technique like CSMA/CD as two nodes on a broadcast network sharing the media explains why each frame is treated independently. Although both switches' databases have been updated, destination D still lacks information.
3. The frame is now duplicated twice on LAN 1. Repetition of Step 2 results in sending both copies to LAN2.
4. The procedure keeps going and going. Keep in mind that switches also serve as repeaters and frame regenerators. Therefore, fresh copies of the frames are generated at each cycle.

Spanning tree algorithm

The IEEE specification mandates that switches employ the spanning tree approach to establish a loop-less topology in order to address the looping issue. A graph without a loop is referred to as a spanning tree in graph theory. This entails setting up a topology in a switched LAN where each LAN can only be accessible from any other LAN via a single link (no loop). Due to the physical connections between cables and switches, we are unable to alter the system's physical topology, but we can design a logical topology that sits on top of the real one.

A system with four LANs and five switches is depicted in Figure 6. We have demonstrated the physical system and its graph theory representation. We have displayed both LANs and switches as nodes, despite the fact that some textbooks only show the switches as the connecting arcs. The connecting arcs demonstrate the link between a switch and a LAN. Assigning a cost (metric) to each arc is necessary in order to determine the spanning tree. The systems administrator is in charge of interpreting the cost. The fewest number of hops was decided upon. The hop count, however, is typically 1 from a switch to the LAN and 0 in the opposite direction.

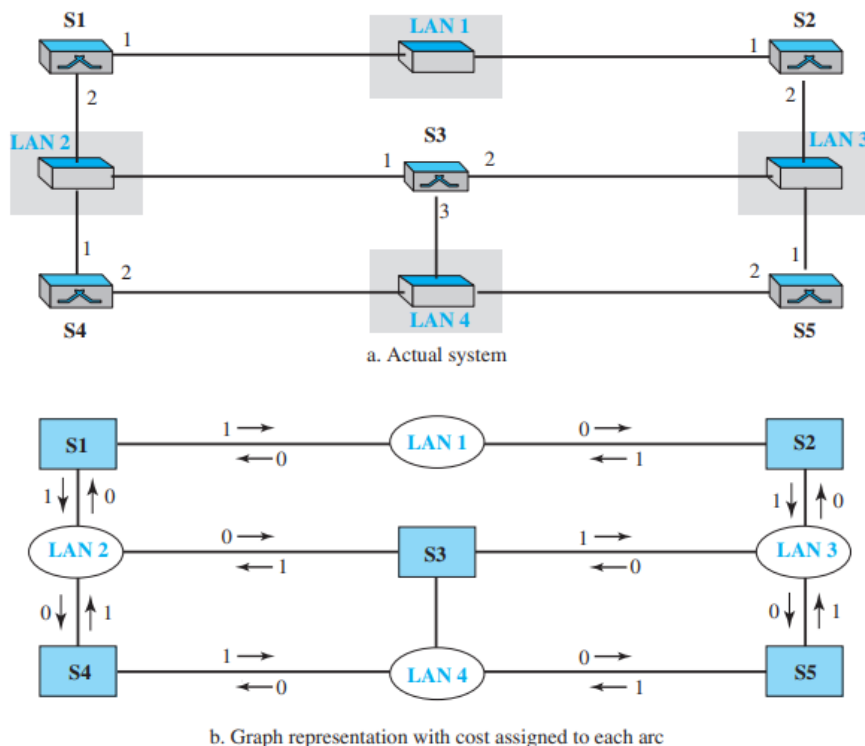


Figure 6: A system of connected LANs and its graph representation

There are three steps to finding the spanning tree:

1. Each switch has an integrated ID (normally the serial number, which is unique). To let other switches know which switch has the smallest ID, each switch broadcasts its ID. As the root switch, the switch with the smallest ID is chosen (root of the tree). Assumed to have the smallest ID is switch S1. As a result, it gets chosen as the root switch.
2. The method looks for the route from the root switch to every other switch or LAN that has the shortest cost. The entire cost from the root switch to the destination can be looked at in order to determine the shortest route. The quickest routes are shown in Figure 7. The Dijkstra algorithm was employed by us.

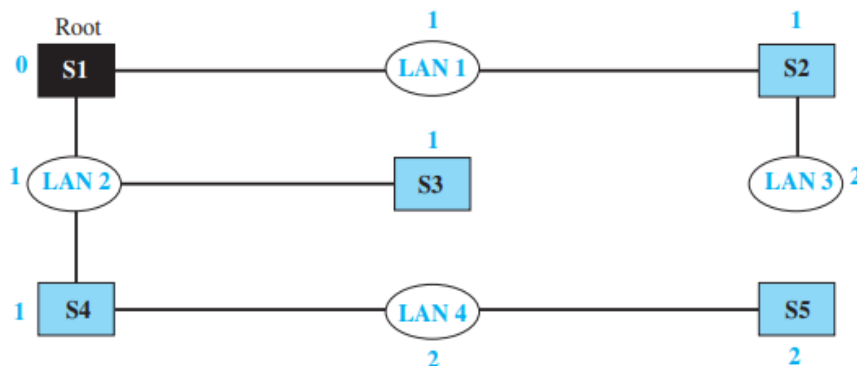


Figure 7: Finding the shortest paths and the spanning tree in a system of switches

3. The shortest tree is produced by combining the shortest pathways, as seen in Figure 7.
4. We identify the forwarding ports, or ports that are a member of the spanning tree, that forward a frame that the switch receives. We also identify the ports that block the frames the switch receives and are not a part of the spanning tree. The logical systems of LANs are depicted in Figure 8 as forwarding ports (solid lines) and blocking ports (broken lines).

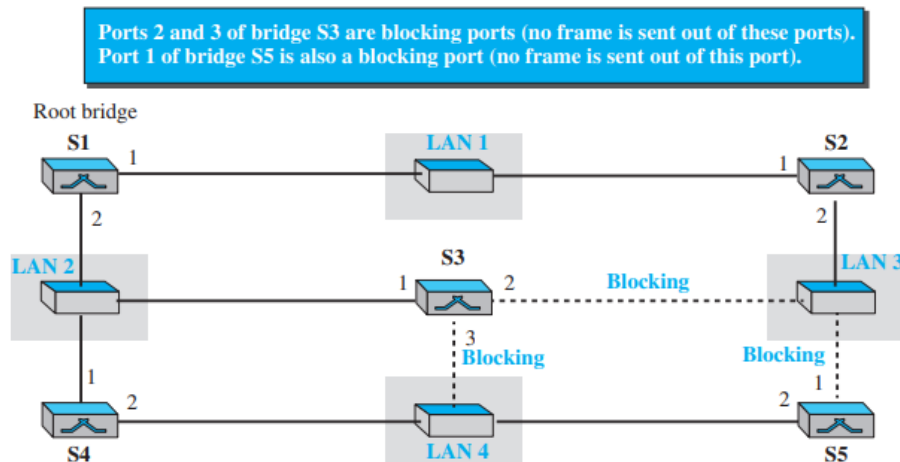


Figure 8: Forwarding and blocking ports after using spanning tree algorithm

In the spanning tree structure, it should be noted that there is only one path from any LAN to any other LAN. This indicates that there is only one route between any two LANs. Loops are not produced. That there is only one route from LAN 1 to LAN 2, LAN 3, or LAN 4 can be demonstrated by you. The same is true for the connections between LAN 2 and LANs 1, 3, and 4. The same applies to LANs 3 and 4.

The spanning tree algorithm has been explained as if manual entries are necessary. That is untrue. A software package that performs this function dynamically is installed in each switch.

Advantages of switches

Compared to a hub, a link-layer switch has a number of benefits. Here, we only touch about two of them.

1] Elimination of Collisions

The collision is eliminated by a link-layer swap. By doing this, the network host's average bandwidth will be increased. Each host can broadcast at any time in a switched LAN, eliminating the requirement for carrier sensing and collision detection.

2] Using Cables to Connect Diverse Devices

Devices that employ various protocols at the physical layer (data rates) and various transmission media can be connected via a link-layer switch. A switch can receive a frame from a device that uses twisted-pair cable and delivers data at 10 Mbps and transfer the frame to another device that uses fiber-optic cable and can receive data at 100 Mbps as long as the frame format at the data-link layer does not change.

9.2.3 Routers

In this section, we discuss routers in comparison to hubs and two-layer switches. A router functions at the physical, data-link, and network layers; it is a three-layer device. It regenerates the signal it receives because it is a physical-layer device. The router examines the packet's source and destination physical addresses as a link-layer device. A router examines the network-layer addresses because it is a network-layer device.

Networks can be linked through a router. To put it another way, a router is an internetworking tool that joins separate networks to create an internetwork. This definition states that an internet or an internetwork is created when two networks are connected via a router.

A router differs significantly from a repeater or switch in three ways.

1. Each of a router's interfaces has a physical and logical (IP) address.
2. A router only responds to packets whose destination address at the link layer coincides with the interface address at which the packet arrives.
3. When a router forwards a packet, it modifies the source and destination link-layer addresses.

Consider the example and refer the figure 9 below. Consider a company with two distinct buildings, each of which has a Gigabit Ethernet LAN installed. Every LAN in the company is equipped with switches. Using 10 Gigabit Ethernet technology, which speeds up the connection to the Ethernet and the connection to the organisation server, the two LANs can be joined to create a larger LAN. The entire system can then be connected to the Internet using a router.

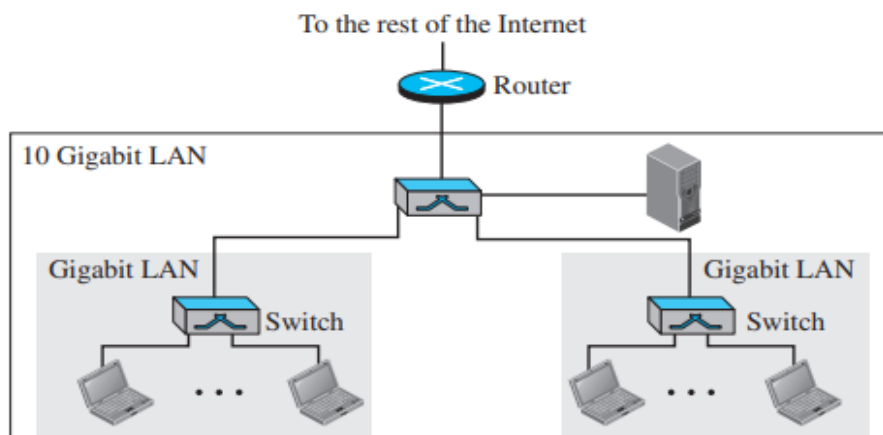


Figure 9: Routing example

A router will change the MAC addresses it receives because the MAC addresses have only local jurisdictions.

9.3 VIRTUAL LANS

If a station physically connects to a LAN, that LAN is regarded as its member. Geographic location is a requirement for membership. What would happen if we required a virtual link between two computers that were connected to two different physical LANs? A virtual local area network (VLAN), as opposed to a physical local area network, is one that is configured by software.

To further explain this definition, let's give an illustration. Figure 10 depicts a switched LAN in an engineering firm with three LANs connected by a switch, each with nine stations.

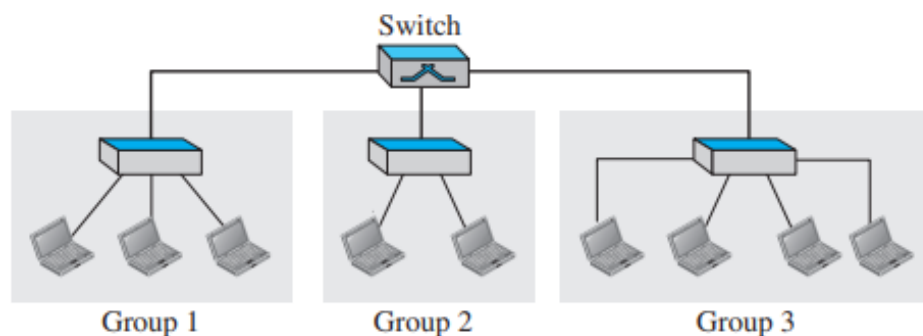


Figure 10: A switch connecting three LANs

The first group of engineers consists of the first three engineers, the second group consists of the next two engineers, and the third group consists of the last four engineers. This configuration of the LAN makes it possible.

To speed up the third group's project, what would happen if the administration required to transfer two engineers from the first group to the third group? It would be necessary to modify the LAN setup. The network specialist has to wire. If the two engineers return to their old team in a week, the issue will recur. Changes in the work group translate into actual changes in the network configuration in a switched LAN.

The same switched LAN is partitioned into VLANs in Figure 11. A LAN is divided into logical, as opposed to physical, parts using VLAN technology. A LAN can be split up into a number of logical LANs, or VLANs. Every VLAN serves as a work group for the company. There is no requirement to alter the physical setup when a person switches between groups. VLAN group membership is governed by software rather than hardware. It is possible to logically migrate any station to another VLAN. Broadcast messages delivered to a certain VLAN can be received by all members of that VLAN.

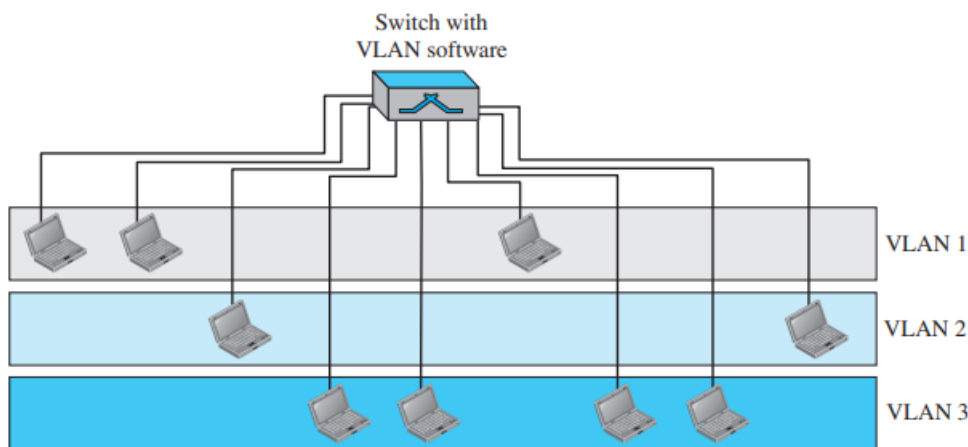


Figure 11: A switch using VLAN software

This indicates that if a station switches from VLAN 1 to VLAN 2, it continues to receive broadcast messages delivered to VLAN 2, but no longer does so.

It is clear that employing VLANs will make it simple to remedy the issue in our prior example. Software makes switching engineers across groups easier than modifying the real network's configuration.

Even the grouping of stations connected to several switches into a VLAN is made possible by VLAN technology. A backbone local area network with two switches and three VLANs is depicted in Figure 12. Each VLAN has stations from switches A and B.

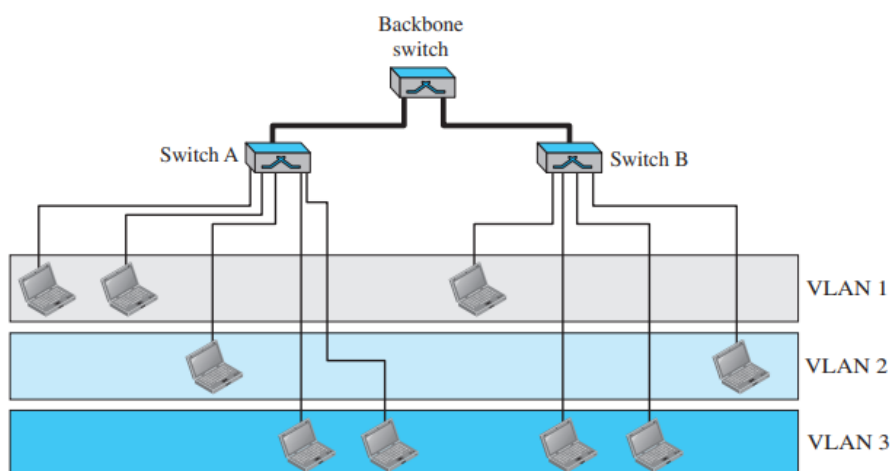


Figure 12: Two switches in a backbone using VLAN software

This arrangement works well for a business with two distinct buildings. Each building may have a backbone-connected switching LAN. Despite being connected to different physical LANs, people from the first and second buildings might be in the same work group.

We can see from these three examples that a VLAN specifies broadcast domains. Stations from one or more physical LANs are grouped together

into broadcast domains by VLANs. A VLAN allows stations to interact with one another as if they were physically connected to a segment.

9.3.1 Membership

What qualifies as a VLAN station grouping characteristic? Various characteristics are employed by vendors, including interface numbers, port numbers, MAC addresses, IP addresses, IP multicast addresses, or a combination of two or more of these.

Interface Numbers

Switch interface numbers are a membership criterion used by some VLAN suppliers. The administrator, for instance, can specify that devices connected to ports 1, 2, 3, and 7 are part of VLAN 1, devices connected to ports 4, 10, and 12 are part of VLAN 2, and so forth.

MAC Addresses

The 48-bit MAC address is a membership trait used by some VLAN suppliers. The administrator could, for instance, specify that devices with the MAC addresses E2:13:42:A1:23:34 and F2:A1:23:BC:D3:41 are part of VLAN 1.

IP address

Some VLAN manufacturers employ the 32-bit IP address as a membership characteristic (see Chapter 18). For instance, the administrator may specify that VLAN 1 is the home network for stations with the IP addresses 181.34.23.67, 181.34.23.72, 181.34.23.98, and 181.34.23.112.

IP addresses for multicast

The multicast IP address is a membership trait used by some VLAN suppliers. At the data-link layer, multicasting at the IP layer is now translated.

Combination

Recently, all of these traits have been possible to be integrated thanks to software offered by some suppliers. When installing the software, the administrator has the option of selecting one or more qualities. The parameters can also be modified by reconfiguring the software.

9.3.2 Configuration

In what way are the stations divided up into various VLANs? There are three ways to set up stations: manually, semi-automatically, and automatically.

Manual configuration

In a manual configuration, the network administrator manually places the stations in various VLANs during setup using the VLAN software.

Additionally, manual migration from one VLAN to another is done later. It should be noted that this structure is intellectual in nature rather than physical. Using the VLAN programme, the administrator types the port numbers, IP addresses, or other attributes manually.

Automatic configuration

In an automatic configuration, the administrator-defined criteria are used to determine whether to automatically connect or disconnect the stations from a VLAN. For instance, the administrator may specify that a group's membership requirements include the project number. A user transitions to a new VLAN whenever they switch projects.

Semi-automatic configuration

Between a manual configuration and an automatic configuration is a semiautomatic configuration. The initialising is often done manually, and migrations are carried out automatically.

9.3.3 Communication between Switches

In a multi-switched backbone, each switch must know not only which station belongs to which VLAN, but also the membership of stations linked to other switches. Switch A, for instance, must be aware of the membership status of stations linked to Switch B, and Switch B, similarly, must be aware of Switch A. For this, three techniques have been developed: frame tagging, temporal division multiplexing, and table maintenance.

Table maintenance

In this way, the switch makes an entry in a database and logs station membership whenever a station delivers a broadcast frame to a member of its group. Periodically, the switches update one another's tables by sending them.

Frame tagging

This approach adds a second header to the MAC frame to specify the destination VLAN as a frame moves between switches. The receiving switches identify the VLANs that will be receiving the broadcast message using the frame tag.

Time division multiplexing

This approach divides the trunk (connection) between switches into time-shared channels. For instance, if a backbone has five total VLANs, each trunk is split into five channels. Channel 1 is used for traffic headed for VLAN 1, Channel 2 is used for traffic headed for VLAN 2, and so on. By examining the channel from which the frame originated, the receiving switch determines the destination VLAN.

IEEE standard

The syntax for frame tagging is specified in a standard called 802.1Q that was adopted in 1996 by the IEEE 802.1 subcommittee. The standard also permits the use of multivendor equipment in VLANs and specifies the format to be used in multi switched backbones. The development of IEEE 802.1Q has paved the way for additional VLAN-related standardisation in other areas. The majority of vendors have already embraced the norm.

9.3.4 Advantages

Using VLANs has a number of benefits:

1] Time and Cost Saving

Station migration costs between groups can be decreased via VLANs. Physical reconfiguration is time-consuming and expensive. It is simpler and faster to move a station using software than it is to physically move it to another segment or even another switch.

2] Establishing Virtual Work Groups

To establish virtual work groups, use VLANs. On a college campus, for instance, teachers working on the same topic can communicate with one another via broadcast messages without having to be in the same department. If the multicasting functionality of IP was previously employed, this may minimise traffic.

3] Safety

VLANs add an additional layer of protection. Users within the same group can send broadcast messages with the certainty that no users inside other groups will receive them.

SUMMARY

A repeater is a connecting device that functions in the Internet model's physical layer. A repeater connects LAN segments, regenerates signals, but it cannot filter data. In the physical and data-link levels of the Internet model, a link-layer switch is a connecting device. A transparent switch has the ability to forward, filter, and generate its forwarding table automatically. The spanning tree approach can be used by a switch to build a loop-less topology.

Instead of actual wiring, software is used to construct a virtual local area network (VLAN). Port numbers, MAC addresses, IP addresses, IP multicast addresses, or any combination of these features may be used to determine membership in a VLAN. VLANs can lessen network traffic, are time and money-effective, and add an additional layer of protection.

LIST OF REFERENCES

- 1] Data Communications and Networking, Behrouz A. Forouzan, Fifth Edition, TMH, 2013.
- 2] Computer Network, Andrew S. Tanenbaum, David J. Wetherall, Fifth Edition, Pearson Education, 2011.
- 3] Computer Network, Bhushan Trivedi, Oxford University Press.
- 4] Data and Computer Communication, William Stallings, PHI.

UNIT END EXERCISES

- 1] What do you mean by connecting devices?
- 2] Explain the concept of hubs.
- 3] Write a note on link layer switches.
- 4] Describe routers in detail.
- 5] With the help of diagram illustrate the concept of virtual LANs.
- 6] What do you mean by the concept of membership.
- 7] State and explain the term configuration associated with virtual LANs.
- 8] Write a note on communication between switches.
- 9] State the advantages of virtual LANs.

INTRODUCTION TO NETWORK LAYER

Unit Structure:

- 10.0 Objectives
- 10.1 Introduction
- 10.2 Network layer services
 - 10.2.1 Packetizing
 - 10.2.2 Routing and Forwarding
 - 10.2.3 Other Services
- 10.3 IPv4 addresses
 - 10.3.1 Address Space
 - 10.3.2 Classful Addressing
- 10.4 Unicast Routing
 - 10.4.1 General Idea
 - 10.4.2 Least-Cost Routing
- 10.5 Routing Algorithms
 - 10.5.1 Distance-Vector Routing
 - 10.5.2 Link-State Routing
 - 10.5.3 Path-Vector Routing

Summary

List of References

Unit End Exercises

10.0 OBJECTIVES

10.1 INTRODUCTION

The TCP/IP protocol suite's network layer is in charge of delivering datagrams from host to host. Both the data-link layer and the transport layer offer services to this layer. We introduce the fundamental ideas and problems of the network layer in this chapter. The addressing method utilised in the network layer is also covered in this chapter. The following chapters will explore further network-layer difficulties after this one.

Before discussing the network layer in the Internet today, let's briefly discuss the network-layer services that, in general, are expected from a network-layer protocol. Figure 1 depicts Alice and Bob's interaction at the network layer.

The figure demonstrates how the Internet is made up of numerous networks (or connections) joined via connecting mechanisms. To put it another way, LANs and WANs combine to form the Internet, which is an internetwork. We must consider the connecting hardware (routers or switches) that link LANs and WANs in order to comprehend the network layer's (or internetwork layer's) function.

The source host, destination host, and each router along the path are all affected by the network layer, as shown in the image (R2, R4, R5, and R7). The network layer at the source host (Alice) receives a packet from the transport layer, wraps it in a datagram, and sends it to the data-link layer. The datagram is decapsulated, extracted, and sent to the appropriate transport layer at the target host (Bob). The routers only require three of the five TCP/IP layers while routing packets, even if the source and destination hosts are engaged in all five. For control purposes, they may need the transport and application levels nevertheless. Because it receives a packet from one network and sends it to another, a router in the path is typically depicted with two data-link levels and two physical layers.

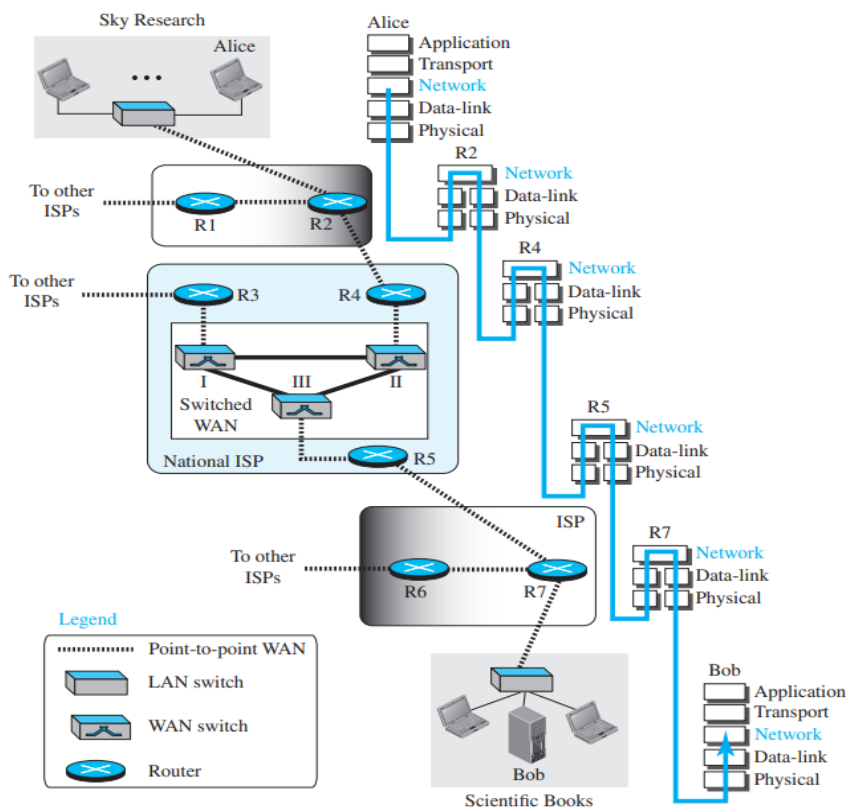


Figure 1: Communication at the network layer

10.2.1 Packetizing

The network layer's primary responsibility is unquestionably packetization, which involves enclosing the payload (data received from the top layer) in a network-layer packet at the source and removing the payload from the packet at the destination. In other words, one responsibility of the network layer is to transmit a payload, unchanged or unmodified, from the source to the destination. In place of a carrier like the post office, which is in charge of delivering packages from a sender to a recipient without altering or utilising the contents, the network layer performs this function.

The source host receives the payload from an upper-layer protocol, adds a header with the source and destination addresses and a few other pieces of data necessary for the network-layer protocol (which will be covered later), and then sends the packet to the data-link layer. Unless the payload is too huge for delivery and needs to be fragmented, the source is not permitted to change its contents.

The destination host receives the network-layer packet at the data-link layer, decapsulates it, and sends the payload to the appropriate upper-layer protocol. The network layer is responsible for waiting until all pieces arrive, reassembling them, and sending them to the upper-layer protocol if the packet is fragmented at the source or at routers along the path.

The packets the routers in the path received are not permitted to be decapsulated unless they require fragmentation. Additionally, source and destination addresses cannot be changed by the routers. To forward the message to the following network on the path, they merely inspect the addresses. The header must be copied to all fragments with certain alterations if a packet is fragmented.

10.2.2 Routing and Forwarding

Routing and forwarding, which are related functions of the network layer, are additional responsibilities that are equally significant as the first.

Routing

The packet must be routed from its source to its destination by the network layer. A physical network is made up of the routers that connect the various networks (LANs and WANs). This indicates that there are several ways to get from one place to another. The optimum route among these potential ones must be found by the network layer. The optimum path needs to be determined by the network layer using some specific tactics. To do this on the Internet today, routing protocols are employed to help the routers coordinate their local knowledge and create consistent tables that will be used when a packet arrives. Before any communication takes place, the routing protocols need be executed.

Forwarding

The action taken by each router when a packet arrives at one of its interfaces can be described as forwarding if routing is the application of strategies and

the execution of specific routing protocols to generate the decision-making tables for each router. The decision-making table that a router often utilises to carry out this function is referred to as either the routing table or the forwarding table. A router must forward a packet it receives from one of its connected networks to another connected network (in unicast routing) or to several connected networks (in multicast routing). The router determines the appropriate output interface number in the forwarding table by using a piece of information from the packet header, which may be the destination address or a label. The concept of a router's forwarding mechanism is depicted in Figure 2.

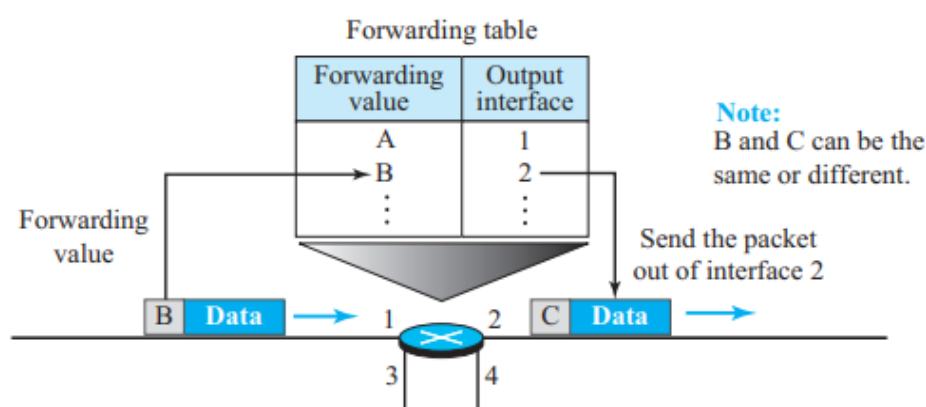


Figure 2: Forwarding process

10.2.3 Other Services

Let's quickly go over additional services that might be expected from the network layer.

1] Error control

Although error control can be done in the network layer, the Internet's designers chose to disregard this problem because the network layer carries data. This choice was made in part because error checking at the network layer is ineffective because packets there may be fragmented at each router.

To prevent any corruption in the header but not the entire datagram, the network layer's designers have however included a checksum field to the datagram. The header of the datagram might not be altered or corrupted thanks to its checksum.

The Internet utilises an auxiliary protocol called ICMP that provides some sort of error control if the datagram is discarded or has some unknown information in the header, despite the fact that the network layer does not directly provide error control.

2] Flow control

A source's ability to convey a certain amount of data without overloading the recipient is controlled by flow control. The receiver will get overloaded with data if the upper layer at the source computer

produces data more quickly than the upper layer at the destination computer can process it. The receiver must communicate with the sender that it is experiencing data overload in order to control the flow of data.

However, flow control is not directly provided by the Internet's network layer. Without regard to the receiver's readiness, the sender sends the datagrams when they are ready.

There are a few reasons why flow control was not included in the network layer's design. First off, the network layer at the receiver's task is so straightforward because there is no error control at this layer that it rarely gets overloaded. Second, the top layers that rely on the network layer's services can implement buffers so they can receive data from the network layer when it is ready and do not have to use it right away. Third, adding another level of flow control makes the network layer more complex and the entire system less efficient because flow control is offered for the majority of upper-layer protocols that utilise network layer services.

3] Congestion control

Congestion control in a network-layer protocol is another problem. When there are too many datagrams in one area of the Internet, there is congestion at the network layer. When source computers send more datagrams than the network or routers can handle, congestion may result. Some routers may drop some of the datagrams in this scenario. However, as additional datagrams are dropped, the situation could worsen since the sender might send copies of the dropped packets as a result of the error control mechanism at the upper levels. Sometimes a scenario may reach a point where the system collapses and no datagrams are transmitted if the congestion persists.

4] Quality of service

The quality of service (QoS) of the connection has become increasingly crucial as new uses for the Internet, such as multimedia communication in particular, real-time audio and video communication have been made possible. By offering higher quality services to serve these applications, the Internet has flourished. These provisions are usually used on the higher layer, nevertheless, to protect the network layer.

5] Security

Security is a concern pertaining to network layer communication as well. When the Internet was first developed, security was not a problem because it was only utilised by a limited group of university students for research purposes; other people had no access to the Internet. There is no security provision in the network layer's design. Security, however, is a major issue today. Another virtual level that converts the connectionless service into a connection-oriented service is required in order to guarantee security for a connectionless network layer. This IPSec-based virtual layer

10.3 IPV4 ADDRESSES

The term "Internet address" or "IP address" refers to the identifier used in the IP layer of the TCP/IP protocol suite to identify each device's connection to the Internet. An IPv4 address is a 32-bit number that specifically and universally identifies how a host or router connects to the Internet. The IP address, not the host or router, is what identifies a connection because it could change if a device is relocated to a different network.

Since each IPv4 address designates a single connection to the Internet, they are distinctive. A device has two IPv4 addresses if it connects to the Internet over two different networks. Any host that wishes to connect to the Internet must accept the IPv4 addressing system, making IPv4 addresses ubiquitous.

10.3.1 Address Space

A protocol like IPv4 that defines addresses has an address space. An address space is the total number of addresses used by the protocol. If a protocol uses b bits to define an address, the address space is 2^b because each bit can have two different values (0 or 1). IPv4 uses 32-bit addresses, which means that the address space is 2^{32} or 4,294,967,296 (more than four billion). If there were no restrictions, more than 4 billion devices could be connected to the Internet.

Notation

An IPv4 address is typically displayed using one of three notations: binary (base 2), dotted-decimal (base 256), or hexadecimal (base 16). An IPv4 address is shown as 32 bits in binary format. Each octet is typically separated by one or more spaces to make the address easier to read (8 bits). A byte is a common term used to describe each octet. The IPv4 address is typically expressed in decimal form with a decimal point (dot) separating the bytes in order to make it shorter and simpler to read. Dotted-decimal notation is the name given to this style of writing. It should be noted that each value in the dotted-decimal notation ranges from 0 to 255 because each byte (octet) only has 8 bits. Hexadecimal notation for IPv4 addresses is occasionally seen. Hexadecimal digits correspond to four bits each. Thus, a 32-bit address consists of 8 hexadecimal digits. The programming of networks frequently employs this notation. An IP address is depicted in each of the three notations in Figure 3.

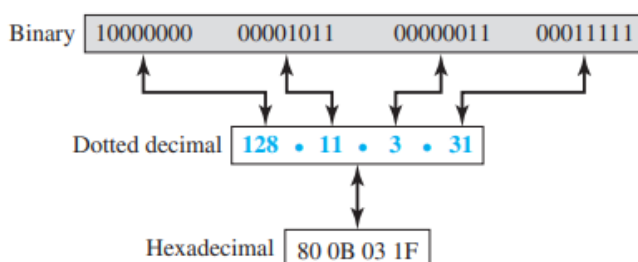


Figure 3: Three different notations in IPv4 addressing

Hierarchy in addressing

The addressing system is hierarchical in any communication network that incorporates delivery, such as a telephone network or a postal network. The country, state, city, street, house number, and name of the mail recipient are all included in the postal address (also known as the mailing address) in a postal network. A phone number is broken down into the local exchange, connection, country code, and area code.

Although it is just separated into two parts, a 32-bit IPv4 address is also hierarchical. The prefix, or first part, of the address designates the network, and the suffix, or second part, designates the node (connection of a device to the Internet). A 32-bit IPv4 address's prefix and suffix are shown in Figure 4. The suffix length is (32-n) bits, while the prefix length is n bits.

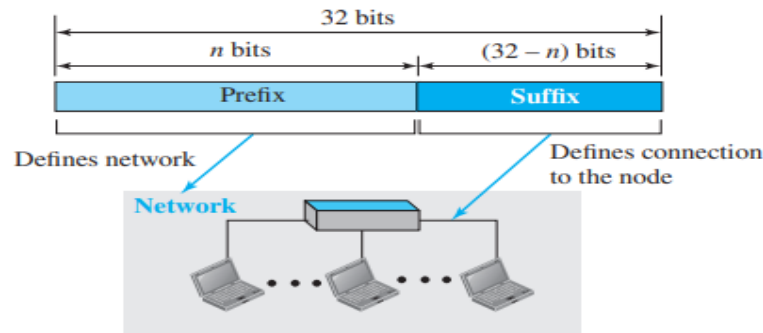


Figure 4: Hierarchy in addressing

Prefixes can have variable or fixed lengths. The IPv4 network identification was initially intended to be a fixed-length prefix. Classful addressing is the term used to describe this outmoded system. The brand-new method, known as classless addressing, makes use of a variable-length network prefix.

10.3.2 Classful Addressing

An IPv4 address originally had a fixed-length prefix, but three fixed-length prefixes ($n = 8$, $n = 16$, and $n = 24$) were created in order to support both small and big networks. According to Figure 5, the entire address space was split into five classes (class A, B, C, D, and E). Classful addressing is the term used to describe this system.

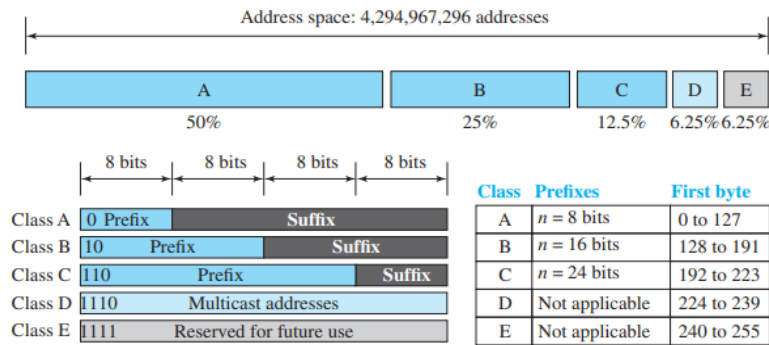


Figure 5: Occupation of the address space in classful addressing

The network length in class A is 8 bits, but since the class is defined by the first bit, which is 0, we can only use seven bits as the network identifier. This indicates that only $2^7 = 128$ networks can have a class A address globally.

The network length in class B is 16 bits, but since the class is defined by the first two bits, which are $(10)_2$, we can only use 14 bits as the network identifier. As a result, only $2^{14} = 16,384$ networks in the entire world are capable of using a class B address.

All addresses that start with $(110)_2$ belong to class C. In class C, the network length is 24 bits, but since three bits define the class, we can have only 21 bits as the network identifier. This means there are $2^{21} = 2,097,152$ networks in the world that can have a class C address.

Prefix and suffix are not separated into classes D and E. Addresses for multicast are used with it. All binary addresses with the prefix 1111 fall under class E. Class E, like Class D, does not have a prefix or a suffix and is used as a reserve.

Address depletion

Address depletion is the cause of classful addressing's obsolescence. The Internet encountered the issue of the addresses being rapidly used up because they were not dispersed effectively. As a result, there were no longer any addresses accessible for businesses and individuals who needed to connect to the Internet. Consider class A in order to comprehend the issue. There are only 128 organisations in the world that can be given this class, but each of those 128 organisations must have a single network with 16,777,216 nodes that is visible to the rest of the world (computers in this single network). Since there may only be a small number of firms this size, the majority of the addresses in this class were useless (unused). Class B addresses were created for medium-sized businesses, however many of these addresses were never used. Class C addresses suffer from a very different design issue. Since each network only allows for 256 addresses, most businesses did not feel safe employing a block in this address class. The class as a whole was wasted because Class E addresses were seldom ever used.

Subnetting and supernetting

Two approaches subnetting and supernetting were suggested and, to a certain extent, used to overcome depletion. A class A or class B block is partitioned into multiple subnets during subnetting. The prefix length of each subnet is longer than that of the original network. A network in class A, for instance, may be partitioned into four subnets, each of which would have the prefix $n_{\text{sub}} = 10$. At the same time, subnetting enables the division of addresses among multiple companies if all of the addresses in a network are not in use. This plan failed because the majority of large businesses objected to splitting the block and allocating some of the unused addresses to smaller organisations.

While supernetting combines multiple class C blocks into one larger block to appeal to enterprises that require more than the 256 addresses offered by a class C block, subnetting was developed to divide a large block into smaller ones. This concept also failed because it made packet routing more challenging.

Advantage

Although classful addressing had a number of drawbacks and was eventually rendered obsolete, it did have one benefit: given an address, we can quickly determine the address' class and, since the prefix length for each class is fixed, we can also quickly determine the prefix length. In other words, classful addressing has an intrinsic prefix length that can be extracted from an address without additional information.

10.4 UNICAST ROUTING

Delivering a datagram from its source to its destination or destinations is the network layer's aim in an internet. We have unicast routing if a datagram is only intended for one destination (one-to-one delivery). Multicast routing is used when a datagram is going to multiple locations (one-to-many delivery). Only hierarchical routing - routing in a series of steps utilising several routing algorithms can be used for unicast routing on the Internet, which has a high number of routers and hosts.

10.4.1 General Idea

With the aid of forwarding tables, a packet is routed from its source to its destination, hop by hop, in unicast routing. Because it sends its packet to the default router in its local network, the source host doesn't require a forwarding table. Because it receives the packet from its local network's default router, the destination host does not require a forwarding table either. Therefore, only the routers that link the various internet networks require forwarding tables. With the aforementioned justification, sending a packet from its source to its destination entails sending it from a source router (the host's default router) to a destination router (the router connected to the destination network). The question is which other routers a packet should traverse in addition to the source and destination routers. In other words, there are various routes a packet might take to get from its source to its destination; the choice of which route it should follow must be made.

An internet as a graph

An internet can be treated as a graph to determine the optimum path. In computer science, a graph is a collection of nodes connected by edges (lines). Each router can be viewed as a node in a graph that represents the internet, and every network connecting two routers can be viewed as an edge. In reality, the internet is modelled as a weighted graph with costs attached to each edge. The nodes and edges of a weighted graph used to represent a geographic area can be cities and the weights in this case are the distances between the cities. The cost of an edge in routing, however, is interpreted differently in various routing protocols, which we examine in a

later section. For the time being, we'll suppose that each edge has a price. The cost is infinite if there is no edge between the nodes. Figure 6 demonstrates how a graph can be used to model the internet.

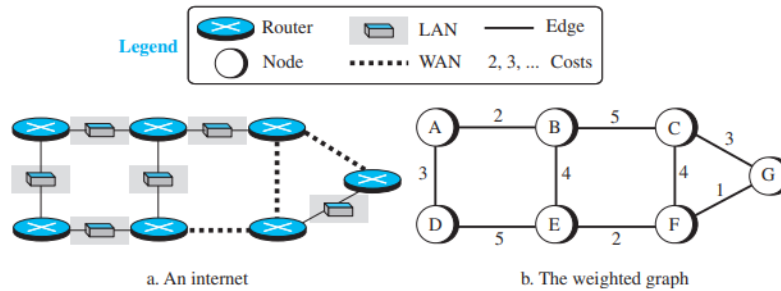


Figure 6: An internet and its graphical representation

10.4.2 Least-Cost Routing

Finding the least expensive route between the source and destination routers is one technique to interpret the optimal route when the internet is modelled as a weighted graph. To put it another way, the source router selects a path to the destination router so that the overall cost of the route is the lowest of all feasible paths. The cheapest route between A and E in Figure 6 is A-B-E, which costs 6. In order to route a packet using this criteria, each router must determine the least-cost route between itself and every other router.

Least-cost trees

There are $(N - 1)$ least-cost pathways connecting any two routers in an internet with N routers. This suggests that for the entire internet, we require $N(N - 1)$ least-cost paths. 90 least-cost paths are required in an internet with only 10 routers. Combining all of these paths into a least-cost tree will enable you to visualise them more clearly. A least-cost tree is one that spans the entire graph (visits every other node), has the source router as its root, and has the shortest path possible between all of its nodes. In this method, each node can only have one shortest-path tree, whereas the entire internet can have N least-cost trees. Later in this section, we demonstrate how to generate a least-cost tree for each node; for now, Figure 7 displays the seven least-cost trees for the internet in Figure 6.

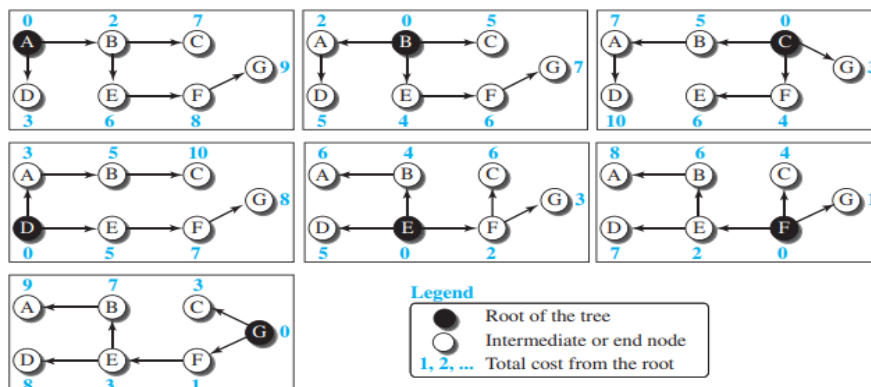


Figure 7: Least-cost trees for nodes in the internet of Figure 6

The least-cost trees for a weighted graph can have several properties if they are created using consistent criteria.

1. The least-cost route from X to Y in X's tree is the inverse of the least-cost route from Y to X in Y's tree; the cost in both directions is the same. For example, in Figure 7, the route from A to F in A's tree is (A → B → E → F), but the route from F to A in F's tree is (F → E → B → A), which is the inverse of the first route. The cost is 8 in each case.
2. Instead of travelling from X to Z using X's tree, we can travel from X to Y using X's tree and continue from Y to Z using Y's tree. For example, in Figure 7, we can go from A to G in A's tree using the route (A → B → E → F → G). We can also go from A to E in A's tree (A → B → E) and then continue in E's tree using the route (E → F → G). The combination of the two routes in the second case is the same route as in the first case. The cost in the first case is 9; the cost in the second case is also 9 (6 + 3).

10.5 ROUTING ALGORITHMS

Now that we've covered the fundamental theory underlying least-cost trees and the forwarding tables that can be created using them, it's time to focus on the routing algorithms. There have been numerous routing algorithms developed in the past. The interpretation of least cost and the process used to build the least-cost tree for each node differ between different methods. Here, we go over the typical algorithms.

10.5.1 Distance-Vector Routing

A router is required by a distance-vector routing (DVR) protocol to periodically notify its neighbours of topological changes. The previous ARPANET routing method was known as (or known as Bellman-Ford algorithm).

Bellman Ford Basis - Each router keeps track of its distance from ALL potential destination nodes in a Distance Vector table. The neighbours' distance vectors are used to provide information for computing distances based on a particular metric.

Information kept by DV router -

- Each router has an ID
- Associated with each link connected to a router, there is a link cost (static or dynamic).
- Intermediate hops

Distance Vector Table Initialization -

- Distance to itself = 0
- Distance to ALL other routers = infinity number.

Distance Vector Algorithm –

1. A router transmits its distance vector to each of its neighbors in a routing packet.
2. Each router receives and saves the most recently received distance vector from each of its neighbors.
3. A router recalculates its distance vector when:
 - It receives a distance vector from a neighbor containing different information than before.
 - It discovers that a link to a neighbor has gone down.

The DV calculation is based on minimizing the cost to each destination

```

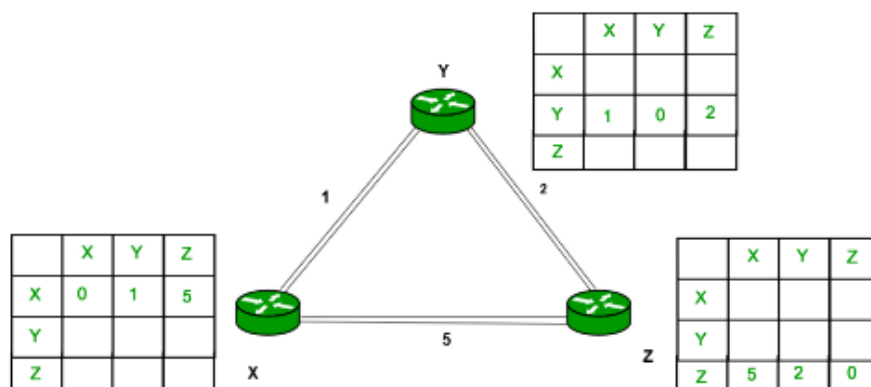
Dx(y) = Estimate of least cost from x to y
C(x,v) = Node x knows cost to each neighbor v
Dx = [Dx(y): y ∈ N] = Node x maintains distance vector
Node x also maintains its neighbors' distance vectors
- For each neighbor v, x maintains Dv = [Dv(y): y ∈ N]
  
```

From time-to-time, each node sends its own distance vector estimate to neighbors.

When a node x receives new DV estimate from any neighbor v, it saves v's distance vector and it updates its own DV using B-F equation:

$$Dx(y) = \min \{C(x,v) + Dv(y), Dx(y)\} \text{ for each node } y \in N$$

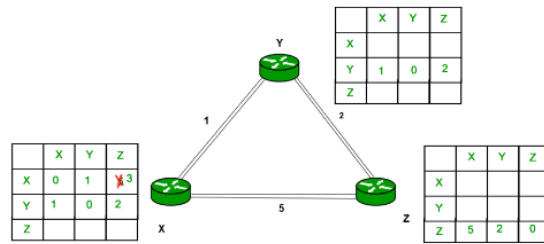
Example – Consider 3-routers X, Y and Z as shown in figure. Each router have their routing table. Every routing table will contain distance to the destination nodes.



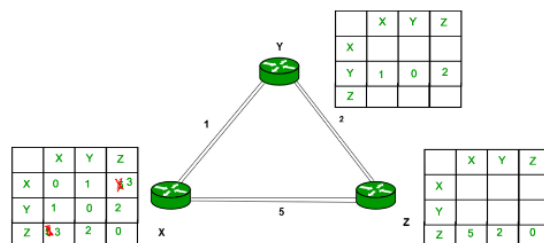
Consider router X, X will share its routing table to neighbors and neighbors will share their routing table to it. X and distance from node X to destination will be calculated using the Bellman-Ford equation.

$$D_x(y) = \min \{ C(x,v) + D_v(y) \} \text{ for each node } y \in N$$

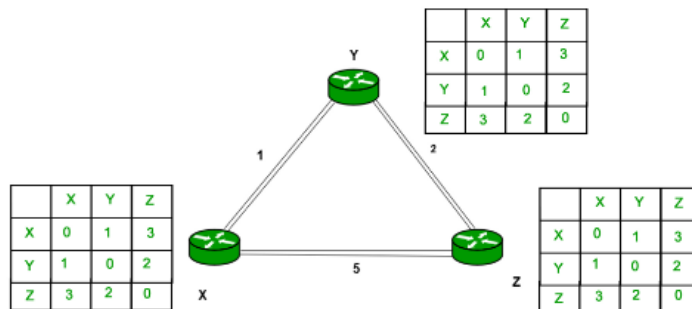
As we can see that distance will be less going from X to Z when Y is intermediate node(hop) so it will be update in routing table X.



Similarly for Z also –



Finally, the routing table for all –



Advantages of Distance Vector routing –

- It is simpler to configure and maintain than link state routing.

Disadvantages of Distance Vector routing –

- It is slower to converge than link state.
- It is at risk from the count-to-infinity problem.
- It creates more traffic than link state since a hop count change must be propagated to all routers and processed on each router. Hop count updates take place on a periodic basis, even if there are no changes

in the network topology, so bandwidth-wasting broadcasts still occur.

- For larger networks, distance vector routing results in larger routing tables than link state since each router must know about all other routers. This can also lead to congestion on WAN links.

10.5.2 Link-State Routing

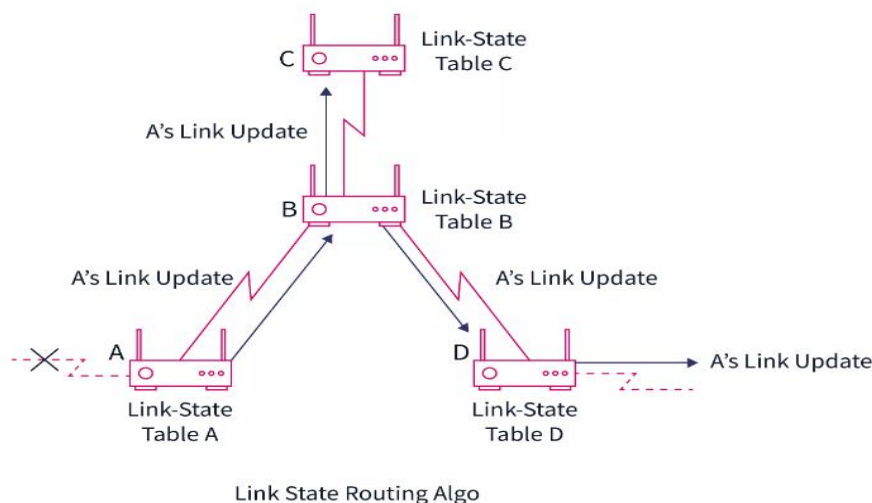
Every router uses the internal Link State Routing Algorithm protocol to communicate with and learn about other routers on the network. Every router computes its routing table using the distributed link state routing method.

A router can create its routing table using knowledge of the network topology. Now, a router uses a shortest path computation algorithm, such as Dijkstra's algorithm, together with the topological knowledge to create the routing table. Each router builds a routing table that is shared with the other routers in the network, enabling faster and more dependable data transmission.

With the other routers in the inter-network, a router does not send its whole routing table. It solely transmits data about its neighbours. This information, together with details on all of the routers it is directly linked to and the cost of the connection, is broadcast by a router.

Flooding now refers to the process of transmitting data about a router's neighbours. Except for its neighbours, a router transmits data to all other routers in the inter-network. Each router that gets the data duplicates it and transmits copies to all of its neighbours. The information is shared throughout the network's interconnected routers in this fashion.

Only when there is a change in the information does this information exchange take place. Consequently, the link state routing algorithm works well. For a simple explanation of the router and the updates made by the link state routing method, see the figure below.



Important Points Related to the Link State Routing Algorithm:

- 1] The link state routing technique only exchanges data when the connection changes.
- 2] It needs a lot of RAM because it keeps a routing database.
- 3] The quickest path must be calculated, which adds overhead to the CPU.
- 4] Each router's data must be distributed throughout the network.

Link state routing protocols

The optimum route from the source to the destination is provided by a routing protocol, which is a routing algorithm.

A router notifies its surrounding routers of its IP address, MAC address, and signature using the link state routing protocol. The surrounding routers now generate a record by fusing the IP address and the MAC address using the data (i.e. IP address, MAC address, and signature). This data enables the router to send the data packet along the best route. Additionally, it informs a router of all the potential paths.

Let's talk about the several protocols that employ the link state routing standard.

- Open Shortest Path First, or OSPF, is a routing protocol that employs the link state routing algorithm to communicate with other inter-network routers and exchange data (such as cost of the route and information about surrounding routers).
- Mobile ad hoc networks and wireless ad hoc networks use the OLSR, or Optimized Link State Routing Protocol, which is an optimised link state routing protocol. To determine the nearby routers that are connected as well as the cost of the connection, the OLSR sends a hello message. It also makes use of the Topology Control messages in addition to the greeting message.

Phases of link state routing

Now let's talk about the link state routing algorithm's two phases. The link state routing method has two steps, which are:

- 1] **Reliable Flooding:** As previously mentioned, a router uses the flooding approach to distribute its information. The gathering and transmission of neighbourhood information takes place in this initial stage. The initial state and the final state are the two parts of the first phase, which is reliable flooding.
 - **Initial State:** Each router learns the cost of connection to its neighbours in the initial stage of dependable flooding.
 - **Final State:** In the dependable flooding's final state, each router is aware of the details of the complete router network (graph).
- 2] **Route Calculation:** The second phase, or route calculation, involves each router using a shortest path computing algorithm, such as

Dijkstra's algorithm, to determine the least expensive, or most ideal, paths to each router.

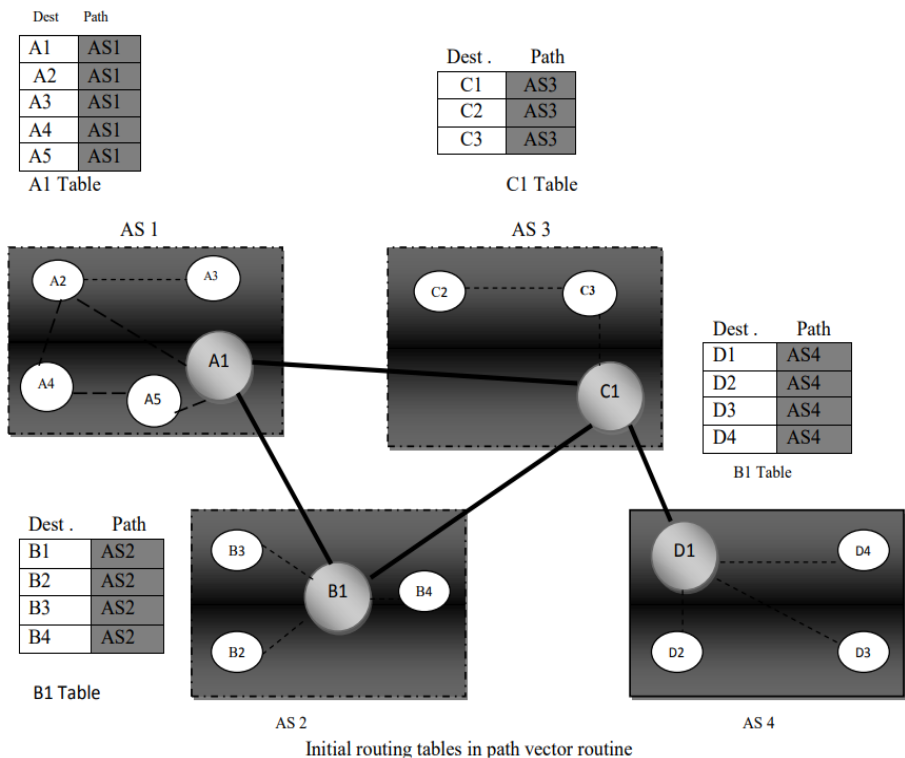
Features of link state routing algorithm

Now let's talk about the different aspects of the link state routing method.

- 1] Link state packet: The link state packet is a brief data packet that contains details on the routing data.
- 2] Link state database: The link state database is a collection of data on the neighbours of a specific router and the cost of a connection.
- 3] Shortest path communication algorithm: The Dijkstra Algorithm is used to determine the shortest path or the best path to travel in order to reach a specific router, as was previously stated.
- 4] Routing table: The list of all the pathways and interfaces is contained in a routing table.

11.5.3 Path-Vector Routing

Path Vector Routing is a routing technique used in the network layer's unicast routing protocol. It is helpful for interdomain routing. The idea behind distance vector routing and path vector routing are similar. It is predicated that each autonomous system contains a speaker node, which acts on behalf of the entire autonomous system. In order to reach the speaker nodes in the neighbouring ASs, the speaker node in an AS builds a routing table. A speaker node promotes the path in its autonomous system or other autonomous systems, not the nodes' metrics.



In a system with four ASs, it serves as the starting table for each speaker node. Here, Node A1 is the speaker node for AS1, B1 is the speaker node

for AS2, C1 is the speaker node for AS3, and D1 is the speaker node for AS4. Node A1 sets an initial table showing A1 to A5 and these are situated in AS1, and these may be reached through it.

In an autonomous system, a speaker sits at a table with its closest neighbours. In this case, Node A1 shares it with Nodes B1 and C1, Node C1 with Nodes A1, B1, and D1, Node B1 with Nodes A1 and C1, and Node D1 with Node C1.

When router A1 receives a packet for node A3, it recognises that the path is in AS1, but when it receives a packet for node D1, it recognises that the packet should go from AS1, to AS2, and then to AS3. On the other hand, if node D1 in AS4 receives a packet for node A2, it recognises that the path should go through AS4, AS3, and AS1.

Functions

1] LOOP PREVENTION

Path vector routing allows for the avoidance of loop formation. When a router receives a message, it determines whether its autonomous system is on the list of paths leading to the destination or, if looping is present, ignores the message.

2] ROUTING POLICIES

When a router receives a message, it has the option to check the path; if one of the autonomous systems specified in the path violates its policy, it can ignore the message's path and destination, failing to update its routing table with this path or failing to relay the message to nearby routers.

3] OPTIMUM PATH

A route that will lead to the greatest result for the entity in charge of the autonomous system.

SUMMARY

The transport layer in the Internet receives services from the network layer and offers services to it. The network layer's primary functions include packetizing and directing packets from their source to their destination. Other services like flow, error, or congestion control are not seriously addressed by the network layer of the Internet.

The network layer's primary responsibility is to perform packet switching. The datagram approach and the virtual-circuit approach are the two methods used in packet switching. In a connection-oriented network, the second is employed; the first is used in a connectionless network. The network layer now employs the first strategy, but the second strategy is more likely to become the norm.

Addressing is one of the major problems at the network layer. We studied IPv4 addressing in this chapter (the current version). We discussed the IPv4 address space and the two methods of allocating addresses: classful and classless addressing. Despite being obsolete, the first still aids in understanding the second. The entire address space is partitioned into five fixed-size classes when classful addressing is used. Using DHCP and NAT protocols, several issues with the present version's address scarcity can be momentarily fixed.

With the aid of forwarding tables, a packet is routed from its source to its destination, hop by hop, in unicast routing. A packet can take a number of different paths from its source to its destination, but the question is which should be the best. The cost and policy imposed on the journey will determine how best to interpret the phrase.

To determine which route is the best among them, several routing algorithms and the accompanying protocols have been developed; just three have proved successful. Each node in distance-vector routing first constructs its own least-cost tree using the incomplete knowledge it has about its immediate neighbours. To make the trees more and more full and to represent the entire internet, the incomplete trees are traded between close neighbours. In other words, a router continuously shares with all of its neighbours what it knows about the entire internet while using distance-vector routing. The Routing Information Protocol is the protocol that carries out distance-vector routing (RIP).

Link-state routing is another routing method that has been applied to the Internet. The property of a link (an edge) that indicates a network in the internet is referred to in this manner as the link-state. The state of the link in this method is determined by the cost attached to an edge. All routers participate in this mechanism, which saturates the internet with data about link states. A link-state database can be constructed once every router is aware of every state. The link-state database can be used to create the least-cost tree for each router and the accompanying forwarding table. Open Shortest Path First (OSPF) is a link-state routing protocol (OSPF).

The least-cost objective is the foundation for both link-state and distance-vector routing. There are circumstances, nevertheless, in which this objective is not the top priority. Algorithms called path-vector routing were created for this. A packet can always be blocked from visiting a particular router by adding policies to the forwarding table. In path-vector routing, the best path is the one that complies with the imposed policy and is the best route from the source. The Border Gateway Protocol is the protocol that uses path-vector routing (BGP)

LIST OF REFERENCES

- 1] Data Communications and Networking, Behrouz A. Forouzan, Fifth Edition, TMH, 2013.
- 2] Computer Network, Andrew S. Tanenbaum, David J. Wetherall, Fifth Edition, Pearson Education, 2011.
- 3] Computer Network, Bhushan Trivedi, Oxford University Press.
- 4] Data and Computer Communication, William Stallings, PHI.

UNIT END EXERCISES

- 1] Illustrate different services associated with network layer.
- 2] Explain the term packetizing.
- 3] Write a note on routing and forwarding.
- 4] Write a note on IPv4 addresses.
- 5] Explain classful addressing.
- 6] Write a note on unicast routing.
- 7] What do you mean by least-cost routing.
- 8] Explain distance-vector routing.
- 9] Write a note on link-state routing.
- 10] Explain Path vector routing.

❄❄❄❄❄❄❄❄

UNICAST ROUTING

Unit Structure:

- 11.0 Objectives
- 11.1 Introduction
 - 11.1.1 Routing algorithms
 - 11.1.2 Unicast routing protocols
- 11.2 Let us Sum Up
- 11.3 List of References
- 11.4 Bibliography
- 11.5 Unit End Exercises

11.0 OBJECTIVES

After going through this unit, you will be able to:

- To deliver a datagram from its source to its destination or destinations.
- Understand one-to-one delivery
- Understand one-to-many delivery.
- Understand hierarchical routing.
- Understand Source router
- Understand Destination router

11.1 INTRODUCTION

Unicast routing in the Internet, with a large number of routers and a huge number of hosts. It can be done only by using hierarchical routing. Hierarchical routing is a routing in several steps using different routing algorithm. We will try to understand the concept of Internet which is defined as an internetwork made of networks connected by routers.

Also need to understand routing concepts and algorithms. Once the routing concepts and algorithms are understood, we can apply them to the Internet using hierarchical routing.

11.1.1 Routing algorithms

- **Distance-Vector Routing:**
- In distance-vector routing, the first thing each node creates is its own least-cost tree with the rudimentary information it has about its immediate neighbors.
- The incomplete trees are exchanged between immediate neighbors to make the trees more and more complete and to represent the whole internet.

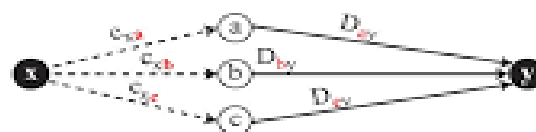
- We can say that in distance-vector routing, a router continuously tells all of its neighbors what it knows about the whole internet (although the knowledge can be incomplete).
- Before we show how incomplete least-cost trees can be combined to make complete ones, we need to discuss two important topics:
 - a) The Bellman-Ford equation
 - b) Distance vectors

a) Bellman-Ford Equation:

- The heart of distance-vector routing is the famous Bellman-Ford equation.
- This equation is used to find the least cost (shortest distance) between a source node, x, and a destination node, y, through some intermediary nodes (a, b, c, . . .) when the costs between the source and the intermediary nodes and the least costs between the intermediary nodes and the destination are given.
- The following shows the general case in which D_{ij} is the shortest distance and c_{ij} is the cost between nodes i and j.
- $D_{xy} = \min\{(C_{xa} + D_{ay}), (C_{xb} + D_{by}), (C_{xc} + D_{cy}), \dots\}$
- In distance-vector routing, normally we want to update an existing least cost with a least cost through an intermediary node, such as z, if the latter is shorter.
- In this case, the equation becomes simpler, as shown below:

$$D_{xy} = \min\{D_{xy}, (C_{xz} + D_{zy})\}$$

Figure 11.1 shows the idea graphically for both cases



a. General case with three intermediate nodes



b. Updating a path with a new route

- We can say that the Bellman-Ford equation enables us to build a new least-cost path from previously established least-cost paths.
- In Figure 11.1, we can think of $(a \rightarrow y)$, $(b \rightarrow y)$, and $(c \rightarrow y)$ as previously established least-cost paths and $(x \rightarrow y)$ as the new least-cost path.
- We can even think of this equation as the builder of a new least-cost tree from previously established least-cost trees if we use the equation repeatedly.
- In other words, the use of this equation in distance-vector routing is a witness that this method also uses least-cost trees, but this use may be in the background.

Distance Vectors:

- The concept of a distance vector is the rationale for the name distance-vector routing.
- A least-cost tree is a combination of least-cost paths from the root of the tree to all destinations.
- These paths are graphically glued together to form the tree.
- Distance-vector routing unglues these paths and creates a distance vector, a one-dimensional array to represent the tree.
- Figure 11.1 shows the tree for node A in the internet in Figure 11.2 and the corresponding distance vector.
- Note that the name of the distance vector defines the root, the indexes define the destinations, and the value of each cell defines the least cost from the root to the destination.
- A distance vector does not give the path to the destinations as the least-cost tree does; it gives only the least costs to the destinations.
- Later we show how we can change a distance vector to a forwarding table, but we first need to find all distance vectors for an internet.
- We know that a distance vector can represent least-cost paths in a least-cost tree, but the question is how each node in an internet originally creates the corresponding vector.
- Each node in an internet, when it is booted, creates a very rudimentary distance vector with the minimum information the node can obtain from its neighborhood.
- The node sends some greeting messages out of its interfaces and discovers the identity of the immediate neighbors and the distance between itself and each neighbor.
- It then makes a simple distance vector by inserting the discovered distances in the corresponding cells and leaves the value of other cells as infinity.

- Do these distance vectors represent least-cost paths? They do, considering the limited information a node has.
- When we know only one distance between two nodes, it is the least cost.
- Figure 11.3 shows all distance vectors for our internet.
- However, we need to mention that these vectors are made asynchronously, when the corresponding node has been booted; the existence of all of them in a figure does not mean synchronous creation of them.

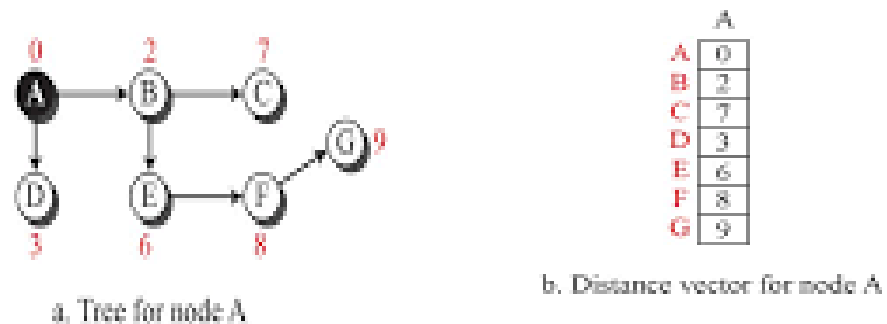


Figure 11.2: The distance vector corresponding to a tree

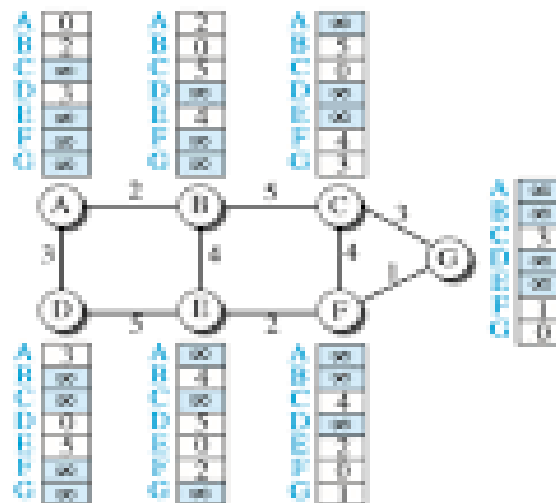


Figure 11.3: The first distance vector for an internet

- These rudimentary vectors cannot help the internet to effectively forward a packet.
- For example, node A thinks that it is not connected to node G because the corresponding cell shows the least cost of infinity.

- To improve these vectors, the nodes in the internet need to help each other by exchanging information.
- After each node has created its vector, it sends a copy of the vector to all its immediate neighbors.
- After a node receives a distance vector from a neighbor, it updates its distance vector using the Bellman-Ford equation (second case).
- However, we need to understand that we need to update, not only one least cost, but N of them in which N is the number of the nodes in the internet.
- If we are using a program, we can do this using a loop; if we are showing the concept on paper, we can show the whole vector instead of the N separate equations.
- We show the whole vector instead of seven equations for each update in Figure 11.4.
- The figure shows two asynchronous events, happening one after another with some time in between.

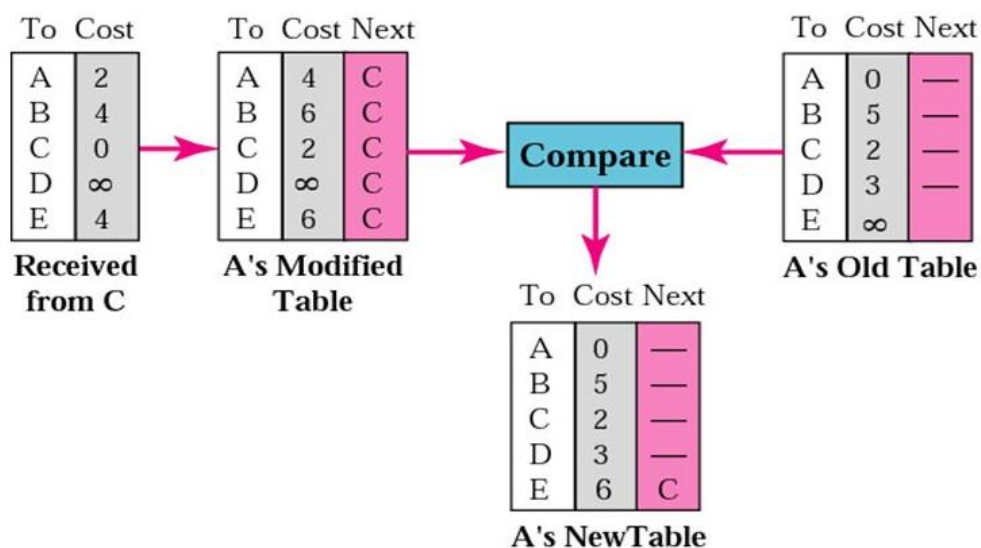


Figure 11.4: Updating distance vectors

- **Count to Infinity:**
- A problem with distance-vector routing is that any decrease in cost propagates quickly, but any increase in cost will propagate slowly.
- For a routing protocol to work properly, if a link is broken, every other router should be aware of it immediately, but in distance-vector routing, this takes some time.

- The problem is referred to as count to infinity.
- It sometimes takes several updates before the cost for a broken link is recorded as infinity by all routers.
- **Two-Node Loop:**
 - One example of count to infinity is the two-node loop problem.
 - To understand the problem, let us look at the scenario depicted in Figure 11.5.
 - The figure shows a system with three nodes.
 - We have shown only the portions of the forwarding table needed for our discussion.
 - At the beginning, both nodes A and B know how to reach node X.

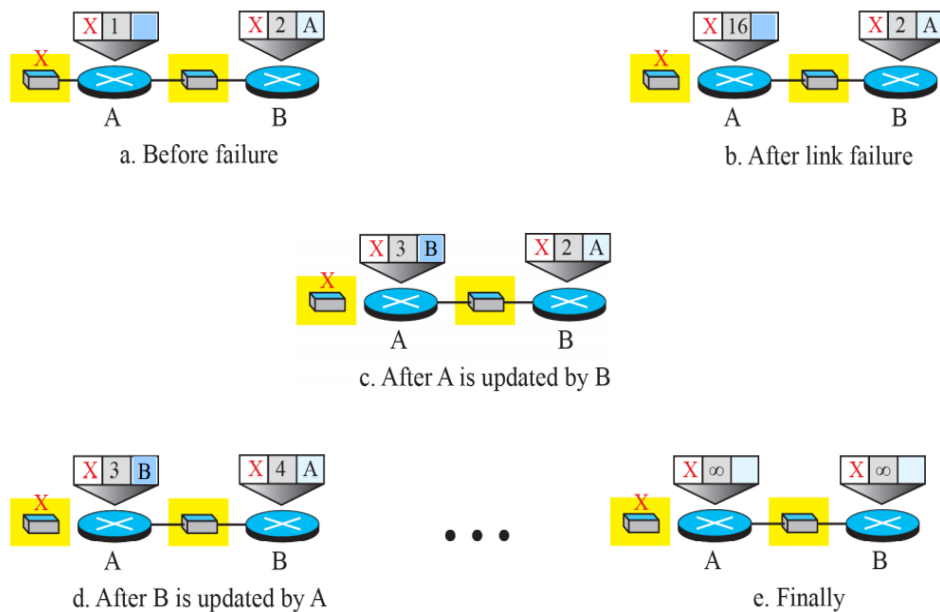


Figure 11.5: Two-node instability

- But suddenly, the link between A and X fails.
- Node A changes its table.
- If A can send its table to B immediately, everything is fine.
- However, the system becomes unstable if B sends its forwarding table to A before receiving A's forwarding table.
- Node A receives the update and, assuming that B has found a way to reach X, immediately updates its forwarding table.

- Now A sends its new update to B.
- Now B thinks that something has been changed around A and updates its forwarding table.
- The cost of reaching X increases gradually until it reaches infinity.
- At this moment, both A and B know that X cannot be reached.
- However, during this time the system is not stable.
- Node A thinks that the route to X is via B; node B thinks that the route to X is via A.
- If A receives a packet destined for X, the packet goes to B and then comes back to A.
- Similarly, if B receives a packet destined for X, it goes to A and comes back to B.
- Packets bounce between A and B, creating a two-node loop problem.
- A few solutions have been proposed for instability of this kind.

Split Horizon

- One solution to instability is called split horizon.
- In this strategy, instead of flooding the table through each interface, each node sends only part of its table through each interface.
- If, according to its table, node B thinks that the optimum route to reach X is via A, it does not need to advertise this piece of information to A; the information has come from A (A already knows).
- Taking information from node A, modifying it, and sending it back to node A is what creates the confusion.
- In our scenario, node B eliminates the last line of its forwarding table before it sends it to A.
- In this case, node A keeps the value of infinity as the distance to X.
- Later, when node A sends its forwarding table to B, node B also corrects its forwarding table.
- The system becomes stable after the first update: both node A and node B know that X is not reachable

Poison Reverse

- Using the split-horizon strategy has one drawback.
- Normally, the corresponding protocol uses a timer, and if there is no news about a route, the node deletes the route from its table.
- When node B in the previous scenario eliminates the route to X from its advertisement to A, node A cannot guess whether this is due to the split-horizon strategy (the source of information was A) or because B has not received any news about X recently.

- In the poison reverse strategy B can still advertise the value for X, but if the source of information is A, it can replace the distance with infinity as a warning: “Do not use this value; what I know about this route comes from you.”

Three-Node Instability

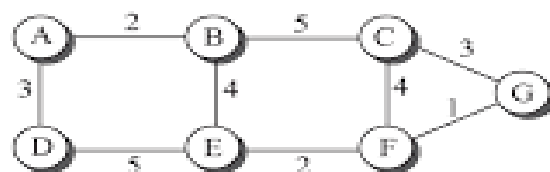
- The two-node instability can be avoided using split horizon combined with poison reverse.
- However, if the instability is between three nodes, stability cannot be guaranteed.

Link-State Routing

- A routing algorithm that directly follows our discussion for creating least-cost trees and forwarding tables is link-state (LS) routing.
- This method uses the term link-state to define the characteristic of a link (an edge) that represents a network in the internet.
- In this algorithm the cost associated with an edge defines the state of the link.
- Links with lower costs are preferred to links with higher costs; if the cost of a link is infinity, it means that the link does not exist or has been broken.

Link-State Database (LSDB)

- To create a least-cost tree with this method, each node needs to have a complete map of the network, which means it needs to know the state of each link.
- The collection of states for all links is called the link-state database (LSDB).
- There is only one LSDB for the whole internet; each node needs to have a duplicate of it to be able to create the least-cost tree.
- Figure 11.6 shows an example of an LSDB for the graph.
- The LSDB can be represented as a two-dimensional array(matrix) in which the value of each cell defines the cost of the corresponding link.



a. The weighted graph

	A	B	C	D	E	F	G
A	0	2	∞	3	∞	∞	∞
B	2	0	5	∞	4	∞	∞
C	∞	5	0	∞	∞	4	3
D	3	∞	∞	0	5	∞	∞
E	∞	4	∞	5	0	2	∞
F	∞	∞	4	∞	2	0	1
G	∞	∞	3	∞	∞	1	0

b. Link state database

- Now the question is how each node can create this LSDB that contains information about the whole internet.
- This can be done by a process called flooding.
- Each node can send some greeting messages to all its immediate neighbors (those nodes to which it is connected directly) to collect two pieces of information for each neighboring node: the identity of the node and the cost of the link.
- The combination of these two pieces of information is called the LS packet (LSP); the LSP is sent out of each interface, as shown in Figure 11.6.
- When a node receives an LSP from one of its interfaces, it compares the LSP with the copy it may already have.
- If the newly arrived LSP is older than the one it has (found by checking the sequence number), it discards the LSP.
- If it is newer or the first one received, the node discards the old LSP (if there is one) and keeps the received one.
- It then sends a copy of it out of each interface except the one from which the packet arrived.
- This guarantees that flooding stops somewhere in the network (where a node has only one interface).
- We need to convince ourselves that, after receiving all new LSPs, each node creates the comprehensive LSDB as shown in Figure 11.7.
- This LSDB is the same for each node and shows the whole map of the internet.
- In other words, a node can make the whole map if it needs to, using this LSDB.

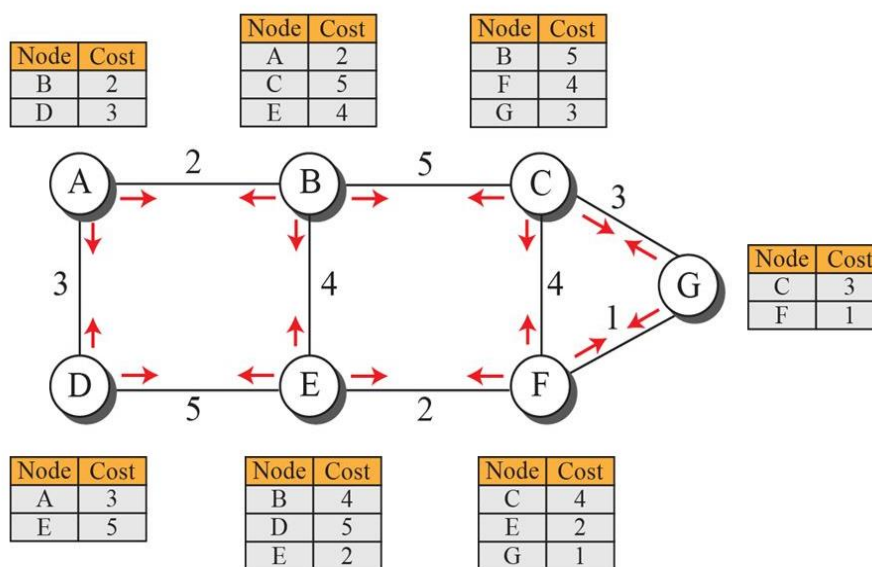


Figure 11.7: LSPs created and sent out by each node to build LSDB

- We can compare the link-state routing algorithm with the distance-vector routing algorithm.
- In the distance-vector routing algorithm, each router tells its neighbors what it knows about the whole internet; in the link-state routing algorithm, each router tells the whole internet what it knows about its neighbors.

Path-Vector Routing

- Both link-state and distance-vector routing are based on the least-cost goal. However, there are instances where this goal is not the priority.
- For example, assume that there are some routers in the internet that a sender wants to prevent its packets from going through.
- For example, a router may belong to an organization that does not provide enough security or it may belong to a commercial rival of the sender which might inspect the packets for obtaining information.
- Least-cost routing does not prevent a packet from passing through an area when that area is in the least-cost path.
- In other words, the least-cost goal, applied by LS or DV routing, does not allow a sender to apply specific policies to the route a packet may take.
- Aside from safety and security, there are occasions, as discussed in the next section, in which the goal of routing is merely reachability: to allow the packet to reach its destination more efficiently without assigning costs to the route.
- To respond to these demands, a third routing algorithm, called path-vector (PV) routing has been devised.
- Path-vector routing does not have the drawbacks of LS or DV routing as described above because it is not based on least-cost routing.
- The best route is determined by the source using the policy it imposes on the route.
- In other words, the source can control the path.
- Although path-vector routing is not actually used in an internet, and is mostly designed to route a packet between ISPs, we discuss the principle of this method in this section as though applied to an internet.
- In the next section, we show how it is used in the Internet.

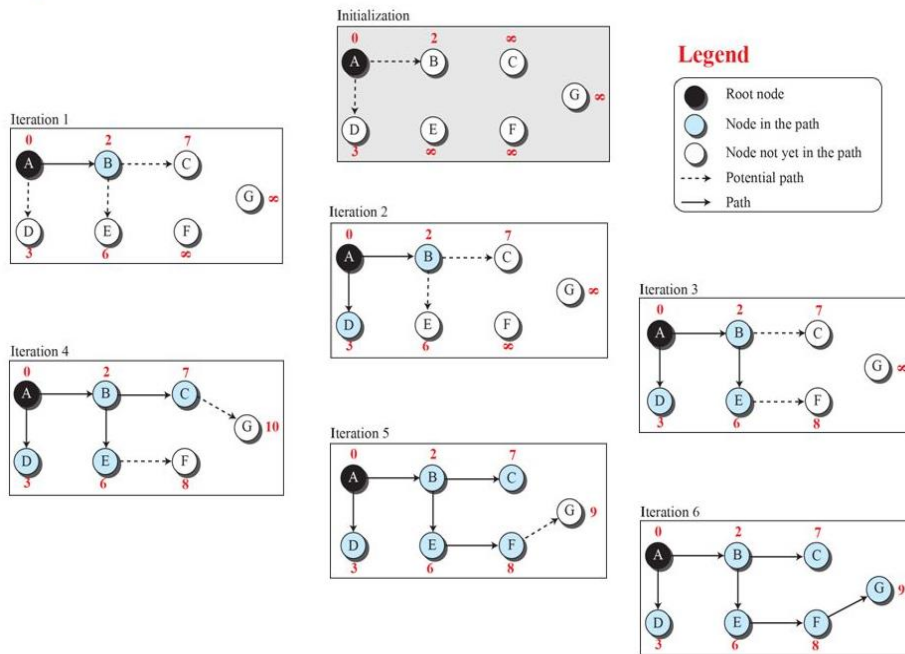


Figure 11.8: Least-cost tree

Spanning Trees

- ☐ In path-vector routing, the path from a source to all destinations is also determined by the best spanning tree.
- ☐ The best spanning tree, however, is not the least-cost tree; it is the tree determined by the source when it imposes its own policy.
- ☐ If there is more than one route to a destination, the source can choose the route that meets its policy best.
- ☐ A source may apply several policies at the same time.
- ☐ One of the common policies uses the minimum number of nodes to be visited (something similar to least-cost).
- ☐ Another common policy is to avoid some nodes as the middle node in a route.
- ☐ Figure 11.9 shows a small internet with only five nodes.
- ☐ Each source has created its own spanning tree that meets its policy.
- ☐ The policy imposed by all sources is to use the minimum number of nodes to reach a destination.
- ☐ The spanning tree selected by A and E is such that the communication does not pass through D as a middle node.
- ☐ Similarly, the spanning tree selected by B is such that the communication does not pass through C as a middle node.

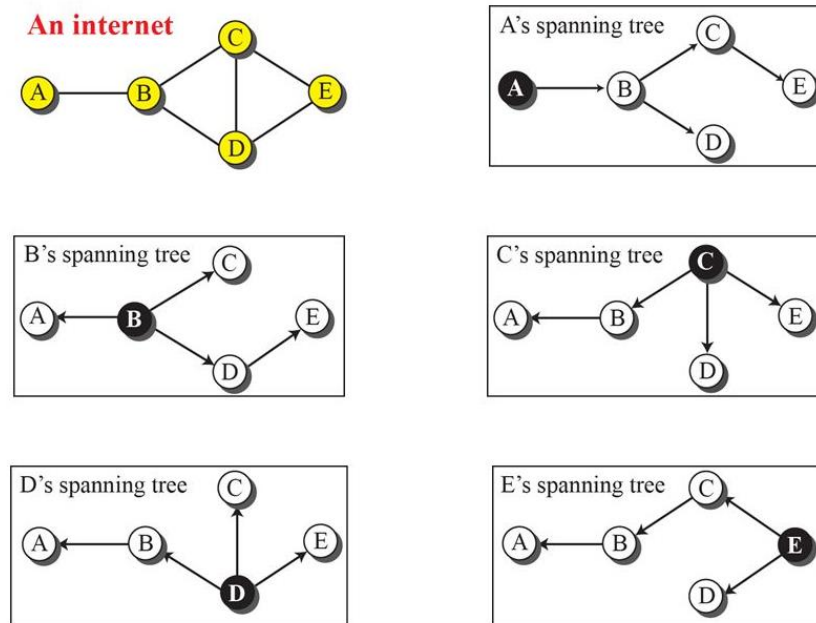


Figure 11.9: Spanning trees in path-vector routing

Creation of Spanning Trees

- Path-vector routing, like distance-vector routing, is an asynchronous and distributed routing algorithm.
- The spanning trees are made, gradually and asynchronously, by each node.
- When a node is booted, it creates a path vector based on the information it can obtain about its immediate neighbor.
- A node sends greeting messages to its immediate neighbors to collect these pieces of information.
- Figure 11.10 shows all of these path vectors for our internet in Figure 11.9.
- Note, however, that we do not mean that all of these tables are created simultaneously; they are created when each node is booted.
- The figure also shows how these path vectors are sent to immediate neighbors after they have been created (arrows).
- Each node, after the creation of the initial path vector, sends it to all its immediate neighbors.
- Each node, when it receives a path vector from a neighbor, updates its path vector using an equation similar to the Bellman-Ford, but applying its own policy instead of looking for the least cost.
- We can define this equation as
- $\text{Path}(x, y) = \text{best} \{ \text{Path}(x, y), [(x + \text{Path}(v, y))] \}$ for all v 's in the internet
- In this equation, the operator (+) means to add x to the beginning of the path.
- We also need to be cautious to avoid adding a node to an empty path because an empty path means one that does not exist.

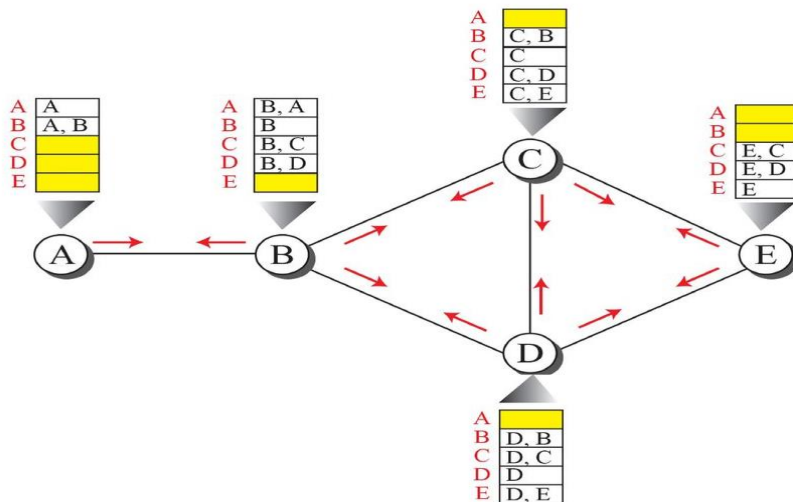
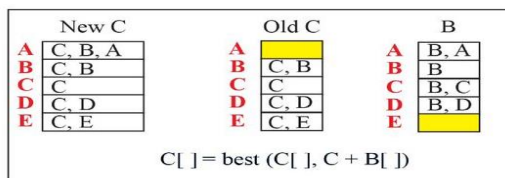


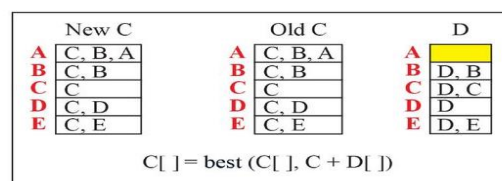
Figure 11.10: Path vectors made at booting time

- The policy is defined by selecting the best of multiple paths.
- Path-vector routing also imposes one more condition on this equation: If Path (v, y) includes x, that path is discarded to avoid a loop in the path.
- In other words, x does not want to visit itself when it selects a path to y.
- Figure 11.11 shows the path vector of node C after two events.
- In the first event, node C receives a copy of B's vector, which improves its vector: now it knows how to reach node A.
- In the second event, node C receives a copy of D's vector, which does not change its vector.
- As a matter of fact, the vector for node C after the first event is stabilized and serves as its forwarding table.



Event 1: C receives a copy of B's vector

Note:
 X []: vector X
 Y: node Y



Event 2: C receives a copy of D's vector

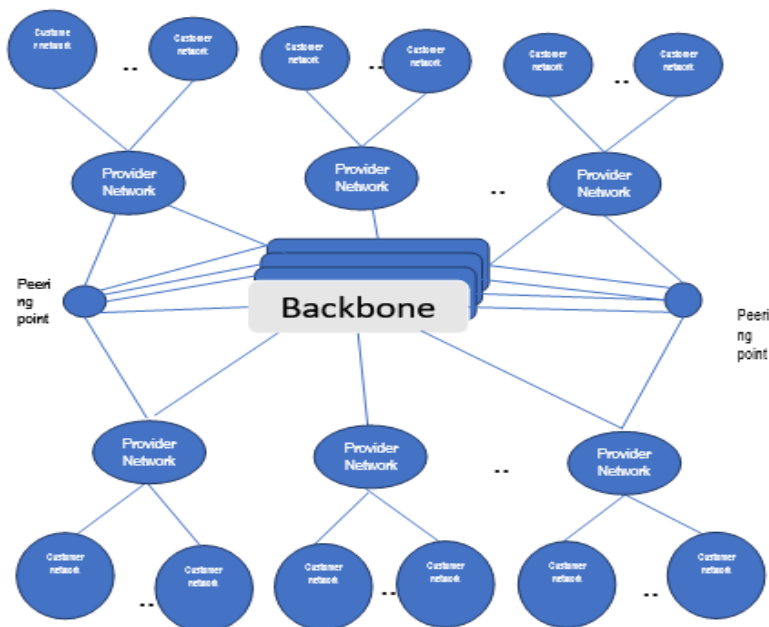
Figure 11.11: Updating path vectors

11.1.2 Unicast routing protocols

A protocol needs to define its domain of operation, the messages exchanged, communication between routers, and interaction with protocols in other domains.

Three common protocols used in the Internet: Routing Information Protocol (RIP), based on the distance-vector algorithm, Open Shortest Path First (OSPF), based on the link-state algorithm, and Border Gateway Protocol (BGP), based on the path-vector algorithm.

- Internet Structure
- The Internet has changed from a tree-like structure, with a single backbone, to a multi-backbone structure run by different private corporations today.
- Although it is difficult to give a general view of the Internet today, we can say that the Internet has a structure similar to what is shown in Figure 11.12



- There are several backbones run by private communication companies that provide global connectivity.
- These backbones are connected by some peering points that allow connectivity between backbones.
- At a lower level, there are some provider networks that use the backbones for global connectivity but provide services to Internet customers.

- Finally, there are some customer networks that use the services provided by the provider networks.
- Any of these three entities (backbone, provider network, or customer network) can be called an Internet Service Provider or ISP.
- They provide services, but at different levels.

Hierarchical Routing

- The Internet today is made of a huge number of networks and routers that connect them.
- It is obvious that routing in the Internet cannot be done using a single protocol for two reasons: a scalability problem and an administrative issue.
- Scalability problem means that the size of the forwarding tables becomes huge, searching for a destination in a forwarding table becomes time-consuming, and updating creates a huge amount of traffic.
- The administrative issue is related to the Internet structure described in Figure 11.12.
- As the figure shows, each ISP is run by an administrative authority.
- The administrator needs to have control in its system.
- The organization must be able to use as many subnets and routers as it needs, may desire that the routers be from a particular manufacturer, may wish to run a specific routing algorithm to meet the needs of the organization, and may want to impose some policy on the traffic passing through its ISP.
- Hierarchical routing means considering each ISP as an autonomous system (AS).
- Each AS can run a routing protocol that meets its needs, but the global Internet runs a global protocol to glue all ASs together.
- The routing protocol run in each AS is referred to as intra-AS routing protocol, intradomain routing protocol, or interior gateway protocol (IGP); the global routing protocol is referred to as inter-AS routing protocol, interdomain routing protocol, or exterior gateway protocol (EGP).
- □ We can have several intradomain routing protocols, and each AS is free to choose one, but it should be clear that we should have only one interdomain protocol that handles routing between these entities.

- Presently, the two common intradomain routing protocols are RIP and OSPF; the only interdomain routing protocol is BGP. The situation may change when we move to IPv6.

Autonomous Systems

- Each ISP is an autonomous system when it comes to managing networks and routers under its control.
- Although we may have small, medium-size, and large ASs, each AS is given an autonomous number (ASN) by the ICANN.
- Each ASN is a 16-bit unsigned integer that uniquely defines an AS.
- The autonomous systems, however, are not categorized according to their size; they are categorized according to the way they are connected to another ASs.
- We have stub ASs, multihomed ASs, and transient ASs.
- The type affects the operation of the interdomain routing protocol in relation to that AS
 - o Stub AS. A stub AS has only one connection to another AS. The data traffic can be either initiated or terminated in a stub AS; the data cannot pass through it. A good example of a stub AS is the customer network, which is either the source or the sink of data.
 - o Multihomed AS. A multihomed AS can have more than one connection to other ASs, but it does not allow data traffic to pass through it. A good example of such an AS is some of the customer ASs that may use the services of more than one provider network, but their policy does not allow data to be passed through them.
 - o Transient AS. A transient AS is connected to more than one other AS and also allows the traffic to pass through. The provider networks and the backbone are good examples of transient ASs.

Routing Information Protocol (RIP)

- The Routing Information Protocol (RIP) is one of the most widely used intradomain routing protocols based on the distance-vector routing algorithm we described earlier.
- RIP was started as part of the Xerox Network System (XNS), but it was the Berkeley Software Distribution (BSD) version of UNIX that helped make the use of RIP widespread.

Hop Count

- A router in this protocol basically implements the distance-vector routing algorithm shown in Table 11.1.
- However, the algorithm has been modified as described below.

- First, since a router in an AS needs to know how to forward a packet to different networks (subnets) in an AS, RIP routers advertise the cost of reaching different networks instead of reaching other nodes in a theoretical graph.
- In other words, the cost is defined between a router and the network in which the destination host is located.
- Second, to make the implementation of the cost simpler (independent from performance factors of the routers and links, such as delay, bandwidth, and so on), the cost is defined as the number of hops, which means the number of networks (subnets) a packet needs to travel through from the source router to the final destination host.
- Note that the network in which the source host is connected is not counted in this calculation because the source host does not use a forwarding table; the packet is delivered to the default router.
- Figure 11.13 shows the concept of hop count advertised by three routers from a source host to a destination host.
- In RIP, the maximum cost of a path can be 15, which means 16 is considered as infinity (no connection).

For this reason, RIP can be used only in autonomous systems in which the diameter of the AS is not more than 15 hops.

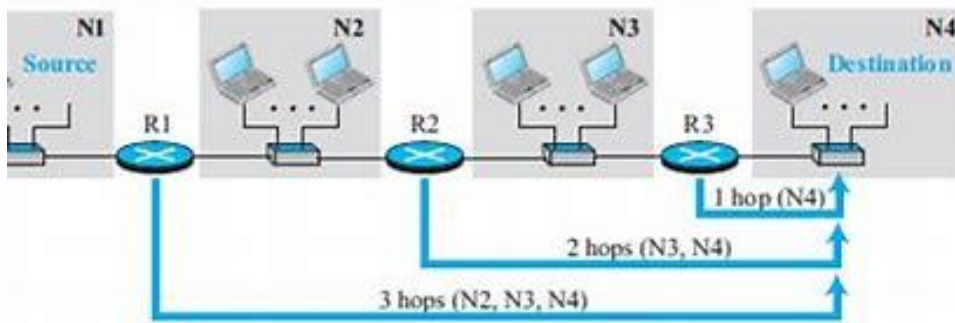


Figure 11.13 hop count

Forwarding Tables

Although the distance-vector algorithm we discussed in the previous section is concerned with exchanging distance vectors between neighboring nodes, the routers in an autonomous system need to keep forwarding tables to forward packets to their destination networks. A forwarding table in RIP is a three-column table in which the first column is the address of the destination network, the second column is the address of the next router to which the packet should be forwarded, and the third column is the cost (the number of hops) to reach the destination network. Figure 11.14 shows the three forwarding tables for the routers in Figure 11.13. Note that the first and the third columns together convey the same information as does a

distance vector, but the cost shows the number of hops to the destination networks.

Forwarding table for R1			Forwarding table for R2			Forwarding table for R3		
Destination network	Next router	Cost in hops	Destination network	Next router	Cost in hops	Destination network	Next router	Cost in hops
N1	—	1	N1	R1	2	N1	R2	3
N2	—	1	N2	—	1	N2	R2	2
N3	R2	2	N3	—	1	N3	—	1
N4	R2	3	N4	R3	2	N4	—	1

Fig 11.14 forwarding tables for the routers

Although a forwarding table in RIP defines only the next router in the second column, it gives the information about the whole least-cost tree based on the second property of these trees, discussed in the previous section.

For example, R1 defines that the next router for the path to N4 is R2; R2 defines that the next router to N4 is R3; R3 defines that there is no next router for this path.

The tree is then $R1 \rightarrow R2 \rightarrow R3 \rightarrow N4$.

A question often asked about the forwarding table is what the use of the third column is.

The third column is not needed for forwarding the packet, but it is needed for updating the forwarding table when there is a change in the route, as we will see shortly.

RIP Implementation

- RIP is implemented as a process that uses the service of UDP on the well-known port number 520.
- In BSD, RIP is a daemon process (a process running in the background), named routed (abbreviation for route daemon and pronounced route-dee).
- This means that, although RIP is a routing protocol to help IP route its datagrams through the AS, the RIP messages are encapsulated inside UDP user datagrams, which in turn are encapsulated inside IP datagrams.
- In other words, RIP runs at the application layer, but creates forwarding tables for IP at the network layer.
- RIP has gone through two versions: RIP-1 and RIP-2.
- The second version is backward compatible with the first section; it allows the use of more information in the RIP messages that were set to 0 in the first version.
- We discuss only RIP-2 in this section.

- Two RIP processes, a client and a server, like any other processes, need to exchange messages.
- RIP-2 defines the format of the message, as shown in Figure 11.15.
- Part of the message, which we call entry, can be repeated as needed in a message.
- Each entry carries the information related to one line in the forwarding table of the router that sends the message.

Com	Ver	Reserved
Family		Tag
Network Address		
Subnet Mask		
Next-hop address		
Distance		

Figure 11.15 format of the message

Explanation of fields: Com: Command, request (1), response (2) Ver: Version, current version is 2 Family: Family of protocol, for TCP/IP value is 2 Tag: Information about autonomous system Network address: Destination address Subnet mask: Prefix length Next-hop address: Address length Distance: Number of hops to the destination

RIP has two types of messages: request and response.

- A request message is sent by a router that has just come up or by a router that has some time-out entries.
- A request message can ask about specific entries or all entries.
- A response (or update) message can be either solicited or unsolicited.
- A solicited response message is sent only in answer to a request message.
- It contains information about the destination specified in the corresponding request message.
- An unsolicited response message, on the other hand, is sent periodically, every 30 seconds or when there is a change in the forwarding table.

RIP Algorithm

- RIP implements the same algorithm as the distance-vector routing algorithm we discussed in the previous section. However, some changes need to be made to the algorithm to enable a router to update its forwarding table:
- Instead of sending only distance vectors, a router needs to send the whole contents of its forwarding table in a response message.
- The receiver adds one hop to each cost and changes the next router field to the address of the sending router. We call each route in the modified forwarding table the received route and each route in the old forwarding table the old route.
- The received router selects the old routes as the new ones except in the following three cases:
 1. If the received route does not exist in the old forwarding table, it should be added to the route.
 2. If the cost of the received route is lower than the cost of the old one, the received route should be selected as the new one.
 3. If the cost of the received route is higher than the cost of the old one, but the value of the next router is the same in both routes, the received route should be selected as the new one.
 - This is the case where the route was actually advertised by the same router in the past, but now the situation has been changed.
 - For example, suppose a neighbor has previously advertised a route to a destination with cost 3, but now there is no path between this neighbor and that destination.
 - The neighbor advertises this destination with cost value infinity (16 in RIP).
 - The receiving router must not ignore this value even though its old route has a lower cost to the same destination.

The new forwarding table needs to be sorted according to the destination route (mostly using the longest prefix first).

Timers in RIP

- RIP uses three timers to support its operation.
- The periodic timer controls the advertising of regular update messages.
- Each router has one periodic timer that is randomly set to a number between 25 and 35 seconds (to prevent all routers sending their messages at the same time and creating excess traffic).

- The timer counts down; when zero is reached, the update message is sent, and the timer is randomly set once again.
- The expiration timer governs the validity of a route.
- When a router receives update information for a route, the expiration timer is set to 180 seconds for that particular route.
- Every time a new update for the route is received, the timer is reset.
- If there is a problem on an internet and no update is received within the allotted 180 seconds, the route is considered expired and the hop count of the route is set to 16, which means the destination is unreachable.
- Every route has its own expiration timer.
- The garbage collection timer is used to purge a route from the forwarding table.
- When the information about a route becomes invalid, the router does not immediately purge that route from its table.
- Instead, it continues to advertise the route with a metric value of 16.
- At the same time, a garbage collection timer is set to 120 seconds for that route.
- When the count reaches zero, the route is purged from the table.
- This timer allows neighbors to become aware of the invalidity of a route prior to purging.

11.2 LET US SUM UP

- We have seen that unicast routing is a part of network layer.
- Also, we have understood Bellman-Ford equation enables us to build a new least-cost path from previously established least-cost paths
- We have seen that the distance-vector (DV) routing uses to find the best route.
- We have understood the concept of a distance vector is the rationale for the name distance-vector routing.
- Also, we have understood the Routing information protocol.

11.3 LIST OF REFERENCES

<https://ieeexplore.ieee.org/document/4976812?arnumber=4976812>

<https://ieeexplore.ieee.org/document/8551330>

11.4 BIBLIOGRAPHY

1. Data Communications and Networking, Behrouz A. Forouzan, Fifth Edition, TMH, 2018.
2. Computer Network, Bhushan Trivedi, Oxford University Press, 2016

11.5 UNIT END EXERCISES

1. Explain unicast routing.
2. Explain Bellman-Ford equation with example.
3. Explain distance-vector (DV) routing.
4. List the steps of creation of spanning tree.
5. Explain path vector routing.
6. Explain Link state database.
7. What is RIP?

NEXT GENERATION IP

Unit Structure:

12.0 Objectives

12.1 Introduction

12.1.1 IPv6 addressing

12.1.2 IPv6 protocol

12.1.3 ICMPv6 protocol

12.1.4 Transition from IPv4 to IPv6

12.2 Let us Sum Up

12.3 List of References

12.4 Bibliography

12.5 Unit End Exercises

12.0 OBJECTIVES

After going through this unit, you will be able to:

- Understand the addressing mechanism in the new generation of the Internet.
- Understand the new packet format
- Understand how the new protocol replaces several auxiliary protocols in version 4.
- Understand strategies that need to be followed for this smooth transition.
- Understand IPv6 ADDRESSING
- Understand Address Types

12.1 INTRODUCTION

The address depletion of IPv4 and other shortcomings of this protocol prompted a new version of IP in the early 1990s. The new version, which is called Internet Protocol version 6 (IPv6) or IP new generation (IPng) was a proposal to augment the address space of IPv4 and at the same time redesign the format of the IP packet and revise some auxiliary protocols such as ICMP.

It is interesting to know that IPv5 was a proposal, based on the OSI model, that never materialized. The following lists the main changes in the IPv6 protocol: larger address space, better header format, new options, allowance for extension, support for resource allocation, and support for more security. The implementation of these changes made it necessary to create a new version of the ICMP protocol, ICMPv6.

12.1.1 IPv6 addressing

- The main reason for migration from IPv4 to IPv6 is the small size of the address space in IPv4.
- The new addressing responds to some problems in the IPv4 addressing mechanism.
- An IPv6 address is 128 bits or 16 bytes (octets) long, four times the address length in IPv4.
- **Representation**
 - A computer normally stores the address in binary, but it is clear that 128 bits cannot easily be handled by humans.
 - Several notations have been proposed to represent IPv6 addresses when they are handled by humans.
 - The following shows two of these notations: binary and colon hexadecimal.
 - Binary (128 bits) 1111111011110110 ... 1111111100000000
 - Colon Hexadecimal
FEF6:BA98:7654:3210:ADEF:BBFF:2922:FF00
 - Binary notation is used when the addresses are stored in a computer.
 - The colon hexadecimal notation (or colon hex for short) divides the address into eight sections, each made of four hexadecimal digits separated by colons.
 - Abbreviation
 - Although an IPv6 address, even in hexadecimal format, is very long, many of the digits are zeros. In this case, we can abbreviate the address.
 - The leading zeros of a section can be omitted. Using this form of abbreviation, 0074 can be written as 74, 000F as F, and 0000 as 0. Note that 3210 cannot be abbreviated.
 - Further abbreviation, often called zero compression, can be applied to colon hex notation if there are consecutive sections consisting of zeros only.
 - We can remove all the zeros and replace them with a double semicolon.
 - FDEC:0:0:0:0:BBFF:0:FFFF ➡ FDEC::BBFF:0:FFFF

- Note that this type of abbreviation is allowed only once per address.
- If there is more than one run of zero sections, only one of them can be
- Compressed

Mixed Notation

- Sometimes we see a mixed representation of an IPv6 address: colon hex and dotted decimal notation.
- This is appropriate during the transition period in which an IPv4 address is embedded in an IPv6 address (as the rightmost 32 bits).
- We can use the colon hex notation for the leftmost six sections and four-byte dotted-decimal notation instead of the rightmost two sections.
- However, this happens when all or most of the leftmost sections of the IPv6 address are 0s.
- For example, the address (::130.24.24.18) is a legitimate address in IPv6, in which the zero compression shows that all 96 leftmost bits of the address are zeros.

CIDR Notation

- IPv6 allows slash or CIDR notation.
- For example, the following shows how we can define a prefix of 60 bits using CIDR.
- We will later show how an IPv6 address is divided into a prefix and a suffix. FDEC::BBFF:0:FFFF/60

Address Space

- The address space of IPv6 contains 2^{128} addresses.
- This address space is 2^{96} times the IPv4 address—definitely no address depletion—as shown, the size of the space is
340, 282, 366, 920, 938, 463, 374, 607, 431, 768, 211, 456
- To give some idea about the number of addresses, we assume that only 1/64 (almost 2 percent) of the addresses in the space can be assigned to the people on planet Earth and the rest are reserved for special purposes.
- We also assume that the number of people on the earth is soon to be 2^{34} (more than 16 billion).
- Each person can have 2^{88} addresses to use.
- Address depletion in this version is impossible.

Three Address Types

In IPv6, a destination address can belong to one of three categories: unicast, anycast, and multicast.

Unicast Address:

- A unicast address defines a single interface (computer or router).
- The packet sent to a unicast address will be routed to the intended recipient.

Anycast Address:

- An anycast address defines a group of computers that all share a single address.
- A packet with an anycast address is delivered to only one member of the group, the most reachable one.
- An anycast communication is used, for example, when there are several servers that can respond to an inquiry.
- The request is sent to the one that is most reachable.
- The hardware and software generate only one copy of the request; the copy reaches only one of the servers.
- IPv6 does not designate a block for any casting; the addresses are assigned from the unicast block.

Multicast Address:

- A multicast address also defines a group of computers.
- However, there is a difference between any casting and multicasting.
- In any casting, only one copy of the packet is sent to one of the members of the group; in multicasting each member of the group receives a copy.
- As we will see shortly, IPv6 has designated a block for multicasting from which the same address is assigned to the members of the group.
- It is interesting that IPv6 does not define broadcasting, even in a limited version.
- IPv6 considers broadcasting as a special case of multicasting.

Address Space Allocation

- Like the address space of IPv4, the address space of IPv6 is divided into several blocks of varying size and each block is allocated for a special purpose.

- Most of the blocks are still unassigned and have been set aside for future use.
- Table 12.1 shows only the assigned blocks.

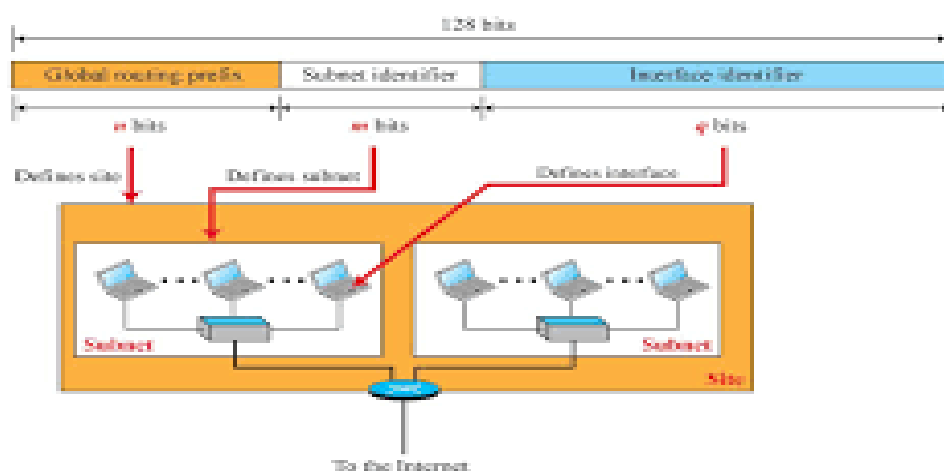
Block prefix	CIDR	Block assignment	Fraction
0000 0000	0000::/8	Special addresses	1/256
001	2000::/3	Global unicast	1/8
1111 110	FC00::/7	Unique local unicast	1/128
1111 1110 10	FE80::/10	Link local addresses	1/1024
1111 1111	FF00::/8	Multicast addresses	1/256

Table 12.1 : Prefixes for assigned IPv6 addresses

- In this table, the last column shows the fraction each block occupies in the whole address space.

Global Unicast Addresses

- The block in the address space that is used for unicast (one-to-one) communication between two hosts in the Internet is called the global unicast address block.
- CIDR for the block is 2000::/3, which means that the three leftmost bits are the same for all addresses in this block (001).
- The size of this block is 2^{125} bits, which is more than enough for Internet expansion for many years to come.
- An address in this block is divided into three parts: global routing prefix (n bits), subnet identifier (m bits), and interface identifier (q bits), as shown in Figure 12.1.
- The figure also shows the recommended length for each part.



- The global routing prefix is used to route the packet through the Internet to the organization site, such as the ISP that owns the block.

- Since the first three bits in this part are fixed (001), the rest of the 45 bits can be defined for up to 2^{45} sites (a private organization or an ISP).
- The global routers in the Internet route a packet to its destination site based on the value of n .
- The next m bits (16 bits based on recommendation) define a subnet in an organization.
- This means that an organization can have up to $2^{16} = 65,536$ subnets, which is more than enough.
- The last q bits (64 bits based on recommendation) define the interface identifier.
- The interface identifier is similar to host id in IPv4 addressing, although the term interface identifier is a better choice because, as we discussed earlier, the host identifier actually defines the interface, not the host.
- If the host is moved from one interface to another, its IP address needs to be changed.
- In IPv4 addressing, there is not a specific relation between the host id (at the IP level) and link-layer address (at the data-link layer) because the link-layer address is normally much longer than the host id.
- The IPv6 addressing allows this relationship.
- A link-layer address whose length is less than 64 bits can be embedded as the whole or part of the interface identifier, eliminating the mapping process.
- Two common link layer addressing schemes can be considered for this purpose: the 64-bit extended unique identifier (EUI-64) defined by IEEE and the 48-bit link-layer address defined by Ethernet.

Autoconfiguration

- One of the interesting features of IPv6 addressing is the autoconfiguration of hosts.
- As we discussed in IPv4, the host and routers are originally configured manually by the network manager.
- However, the Dynamic Host Configuration Protocol, DHCP, can be used to allocate an IPv4 address to a host that joins the network.
- In IPv6, DHCP protocol can still be used to allocate an IPv6 address to a host, but a host can also configure itself.
- When a host in IPv6 joins a network, it can configure itself using the following process:
 1. The host first creates a link local address for itself. This is done by taking the 10-bit link local prefix (1111 1110 10), adding 54 zeros, and adding the 64-bit interface identifier, which any host knows how

to generate from its interface card. The result is a 128-bit link local address.

2. The host then tests to see if this link local address is unique and not used by other hosts. Since the 64-bit interface identifier is supposed to be unique, the link local address generated is unique with a high probability. However, to be sure, the host sends a neighbour solicitation message and waits for a neighbour advertisement message. If any host in the subnet is using this link local address, the process fails and the host cannot autoconfigure itself; it needs to use other means such as DHCP for this purpose.
3. If the uniqueness of the link local address is passed, the host stores this address as its link local address, but it still needs a global unicast address. The host then sends a router solicitation message to a local router. If there is a router running on the network, the host receives a router advertisement message that includes the global unicast prefix and the subnet prefix that the host needs to add to its interface identifier to generate its global unicast address.

Renumbering

- To allow sites to change the service provider, renumbering of the address prefix (n) was built into IPv6 addressing.
- Each site is given a prefix by the service provider to which it is connected. If the site changes the provider, the address prefix needs to be changed.
- A router to which the site is connected can advertise a new prefix and let the site use the old prefix for a short time before disabling it.
- In other words, during the transition period, a site has two prefixes.
- The main problem in using the renumbering mechanism is the support of the DNS, which needs to propagate the new addressing associated with a domain name.
- A new protocol for DNS, called Next Generation DNS, is under study to provide support for this mechanism.

12.1.2 THE IPv6 PROTOCOL

- The change of the IPv6 address size requires the change in the IPv4 packet format.
- The designer of IPv6 decided to implement remedies for other shortcomings now that a change is inevitable.
- The following shows other changes implemented in the protocol in addition to changing address size and format.

Better header format.

- IPv6 uses a new header format in which options are separated from the base header and inserted, when needed, between the base header and the data.

- This simplifies and speeds up the routing process because most of the options do not need to be checked by routers.

New options.

- IPv6 has new options to allow for additional functionalities.

Allowance for extension.

- IPv6 is designed to allow the extension of the protocol if required by new technologies or applications.

Support for resource allocation.

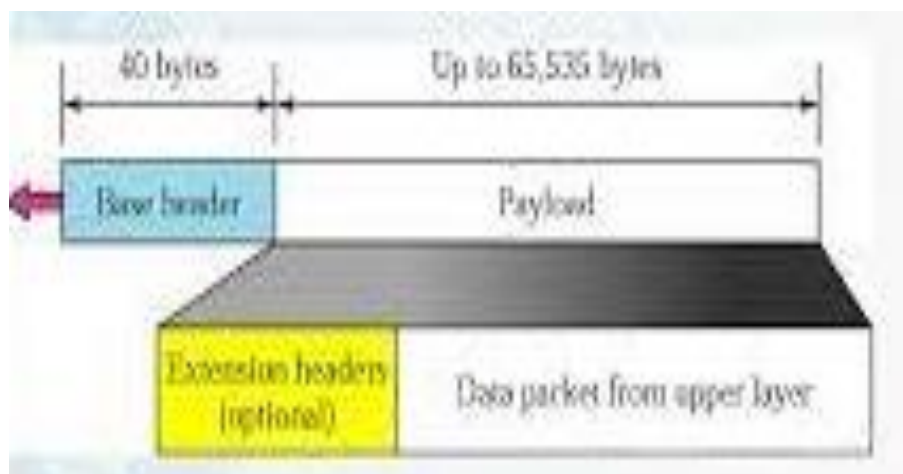
- In IPv6, the type-of-service field has been removed, but two new fields, traffic class and flow label, have been added to enable the source to request special handling of the packet.
- This mechanism can be used to support traffic such as real-time audio and video.

Support for more security.

- The encryption and authentication options in IPv6 provide confidentiality and integrity of the packet.

Packet Format

- The IPv6 packet is shown in Figure 12.2.
- Each packet is composed of a base header followed by the payload.
- The base header occupies 40 bytes, whereas payload can be up to 65,535 bytes of information.
- The description of fields follows.



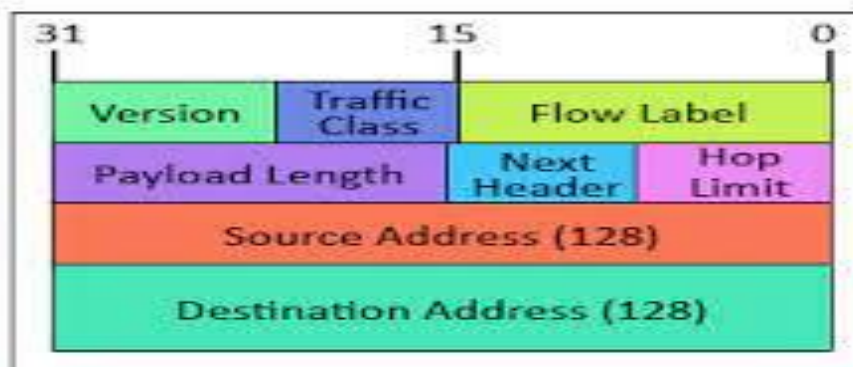


Figure 12.2

Version.

The 4-bit version field defines the version number of the IP. For IPv6, the value is 6.

Traffic class.

- The 8-bit traffic class field is used to distinguish different payloads with different delivery requirements.
- It replaces the type-of-service field in IPv4.

Flow label.

- The flow label is a 20-bit field that is designed to provide special handling for a particular flow of data.

Payload length.

- The 2-byte payload length field defines the length of the IP datagram excluding the header.
- Note that IPv4 defines two fields related to the length: header length and total length.
- In IPv6, the length of the base header is fixed (40 bytes); only the length of the payload needs to be defined.

Next header.

- The next header is an 8-bit field defining the type of the first extension header or the type of the data that follows the base header in the datagram.
- This field is similar to the protocol field in IPv4, but we talk more about it when we discuss the payload.

Hop limit.

- The 8-bit hop limit field serves the same purpose as the TTL field in IPv4.

Source and destination addresses.

- The source address field is a 16-byte (128-bit) Internet address that identifies the original source of the datagram.
- The destination address field is a 16-byte (128-bit) Internet address that identifies the destination of the datagram.

Payload

Compared to IPv4, the payload field in IPv6 has a different format and meaning,

Concept of Flow and Priority in IPv6

- The IP protocol was originally designed as a connectionless protocol.
- However, the tendency is to use the IP protocol as a connection-oriented protocol.
- The MPLS technology described earlier allows us to encapsulate an IPv4 packet in an MPLS header using a label field.
- In version 6, the flow label has been directly added to the format of the IPv6 datagram to allow us to use IPv6 as a connection-oriented protocol.
- To a router, a flow is a sequence of packets that share the same characteristics, such as traveling the same path, using the same resources, having the same kind of security, and so on.
- A router that supports the handling of flow labels has a flow label table.
- The table has an entry for each active flow label; each entry defines the services required by the corresponding flow label.
- When the router receives a packet, it consults its flow label table to find the corresponding entry for the flow label value defined in the packet.
- It then provides the packet with the services mentioned in the entry.
- However, note that the flow label itself does not provide the information for the entries of the flow label table; the information is provided by other means, such as the hop-by-hop options or other protocols.
- In its simplest form, a flow label can be used to speed up the processing of a packet by a router.
- When a router receives a packet, instead of consulting the forwarding table and going through a routing algorithm to define the address of the next hop, it can easily look in a flow label table for the next hop.

- In its more sophisticated form, a flow label can be used to support the transmission of real-time audio and video.
- Real-time audio or video, particularly in digital form, requires resources such as high bandwidth, large buffers, long processing time, and so on.
- A process can make a reservation for these resources beforehand to guarantee that real-time data will not be delayed due to a lack of resources.
- The use of real-time data and the reservation of these resources require other protocols such as Real-Time Transport Protocol (RTP) and Resource Reservation Protocol (RSVP) in addition to IPv6.

Fragmentation and Reassembly

- There are still fragmentation and reassembly of datagrams in the IPv6 protocol, but there is a major difference in this respect.
- IPv6 datagrams can be fragmented only by the source, not by the routers; the reassembly takes place at the destination.
- The fragmentation of packets at routers is not allowed to speed up the processing of packets in the router.
- The fragmentation of a packet in a router needs a lot of processing.
- The packet needs to be fragmented; all fields related to the fragmentation need to be recalculated.
- • In IPv6, the source can check the size of the packet and make the decision to fragment the packet or not.
- • When a router receives the packet, it can check the size of the packet and drop it if the size is larger than allowed by the MTU of the network ahead.
- • The router then sends a packet-too-big ICMPv6 error message to inform the source.

Extension Header

- An IPv6 packet is made of a base header and some extension headers.
- The length of the base header is fixed at 40 bytes.
- However, to give more functionality to the IP datagram, the base header can be followed by up to six extension headers.
- Many of these headers are options in IPv4.
- Six types of extension headers have been defined.

- These are hop-by-hop option, source routing, fragmentation, authentication, encrypted security payload, and destination option as shown in Figure 12.3 extension header types as follows

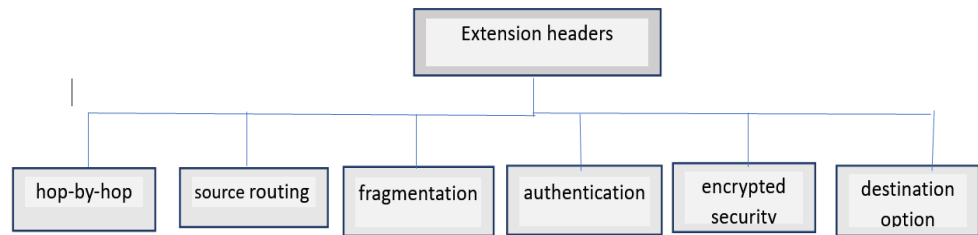


Figure 12.3 extension header types

Hop-by-Hop Option

The hop-by-hop option is used when the source needs to pass information to all routers visited by the datagram. For example, perhaps routers must be informed about certain management, debugging, or control functions. Or, if the length of the datagram is more than the usual 65,535 bytes, routers must have this information.

So far, only three hop by-hop options have been defined: Pad1, PadN, and jumbo payload.

Pad1.

- This option is 1 byte long and is designed for alignment purposes.
- Some options need to start at a specific bit of the 32-bit word.
- If an option falls short of this requirement by exactly one byte, Pad1 is added.

PadN.

- PadN is similar in concept to Pad1.
- The difference is that PadN is used when 2 or more bytes are needed for alignment.

Jumbo payload.

- Recall that the length of the payload in the IP datagram can be a maximum of 65,535 bytes.
- However, if for any reason a longer payload is required, we can use the jumbo payload option to define this longer length.

Destination Option

- The destination option is used when the source needs to pass information to the destination only.
- Intermediate routers are not permitted access to this information.

- The format of the destination option is the same as the hop-by-hop option.
- So far, only the Pad1 and PadN options have been defined

Source Routing

- The source routing extension header combines the concepts of the strict source route and the loose source route options of IPv4.

Fragmentation

- The concept of fragmentation in IPv6 is the same as that in IPv4.
- However, the place where fragmentation occurs differs.
- In IPv4, the source or a router is required to fragment if the size of the datagram is larger than the MTU of the network over which the datagram travels.
- In IPv6, only the original source can fragment.
- A source must use a Path MTU Discovery technique to find the smallest MTU supported by any network on the path.
- The source then fragments using this knowledge.
- If the source does not use a Path MTU Discovery technique, it fragments the datagram to a size of 1280 bytes or smaller.
- This is the minimum size of MTU required for each network connected to the Internet.

Authentication

- The authentication extension header has a dual purpose: it validates the message sender and ensures the integrity of data.
- The former is needed so the receiver can be sure that a message is from the genuine sender and not from an imposter.
- The latter is needed to check that the data is not altered in transition by some hacker.

Encrypted Security Payload

The encrypted security payload (ESP) is an extension that provides confidentiality and guards against eavesdropping.

Comparison of Options between IPv4 and IPv6

- The following shows a quick comparison between the options used in IPv4 and the options used in IPv6 (as extension headers).
- The no-operation and end-of-option options in IPv4 are replaced by Pad1 and Pad N options in IPv6.

- The record route option is not implemented in IPv6 because it was not used.
- The timestamp option is not implemented because it was not used.
- The source route option is called the source route extension header in IPv6.
- The fragmentation fields in the base header section of IPv4 have moved to the fragmentation extension header in IPv6.
- The authentication extension header is new in IPv6.
- The encrypted security payload extension header is new in IPv6.

12.1.2 THE ICMPv6 PROTOCOL

- Another protocol that has been modified in version 6 of the TCP/IP protocol suite is ICMP.
- This new version, Internet Control Message Protocol version 6 (ICMPv6), follows the same strategy and purposes of version 4.
- ICMPv6, however, is more complicated than ICMPv4: some protocols that were independent in version 4 are now part of ICMPv6 and some new messages have been added to make it more useful.
- Figure 12.4 compares the network layer of version 4 to that of version 6. The ICMP, ARP, and IGMP protocols in version 4 are combined into one single protocol, ICMPv6.

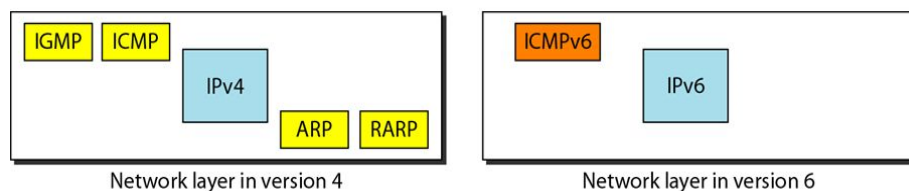
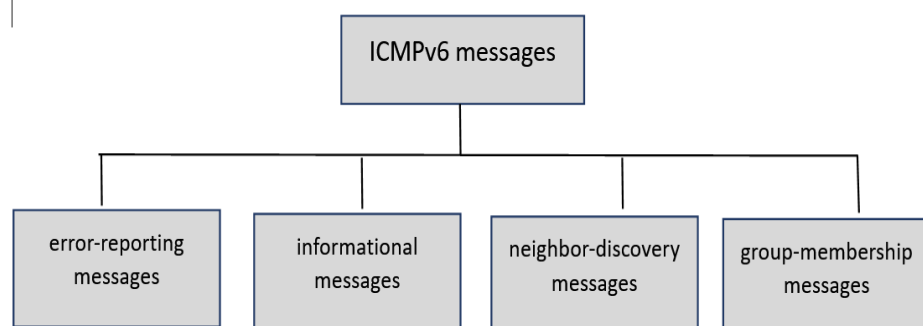


Figure 12.4 Network layer of version 4 to that of version 6

- We can divide the messages in ICMPv6 into four groups: error-reporting messages, informational messages, neighbor-discovery messages, and group-membership messages, as shown in Figure 12.5.



12.1.2 TRANSITION FROM IPv4 TO IPv6

How can we make transition to stop using IPv4 and start using IPv6?

- The first solution that comes to mind is to define a transition day on which every host or router should stop using the old version and start using the new version.
- However, this is not practical; because of the huge number of systems in the Internet, the transition from IPv4 to IPv6 cannot happen suddenly.
- It will take a considerable amount of time before every system in the Internet can move from IPv4 to IPv6.
- The transition must be smooth to prevent any problems between IPv4 and IPv6 systems.
- **Strategies**
- Three strategies have been devised for transition: dual stack, tunneling, and header translation.
- One or all of these three strategies can be implemented during the transition period.

Dual Stack

- It is recommended that all hosts, before migrating completely to version 6, have a dual stack of protocols during the transition.
- In other words, a station must run IPv4 and IPv6 simultaneously until all the Internet uses IPv6.
- See Figure 12.6 for the layout of a dual-stack configuration.

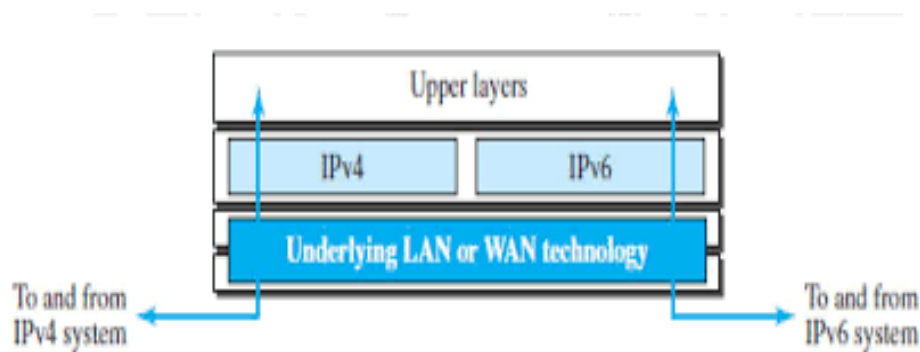


Figure 12.6 Dual-stack

- To determine which version to use when sending a packet to a destination, the source host queries the DNS.
- If the DNS returns an IPv4 address, the source host sends an IPv4 packet.

- If the DNS returns an IPv6 address, the source host sends an IPv6 packet

Tunneling

- Tunneling is a strategy used when two computers using IPv6 want to communicate with each other and the packet must pass through a region that uses IPv4.
- To pass through this region, the packet must have an IPv4 address.
- So the IPv6 packet is encapsulated in an IPv4 packet when it enters the region, and it leaves its capsule when it exits the region.
- It seems as if the IPv6 packet enters a tunnel at one end and emerges at the other end.
- To make it clear that the IPv4 packet is carrying an IPv6 packet as data, the protocol value is set to 41.
- Tunneling is shown in Figure 12.7.

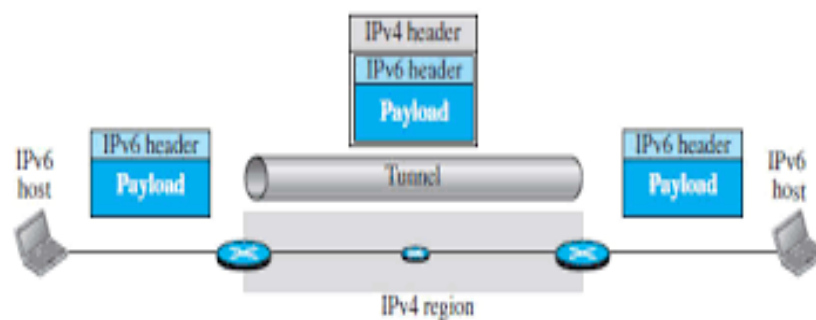


Figure 12.7. Tunneling strategy

Header Translation

- Header translation is necessary when the majority of the Internet has moved to IPv6 but some systems still use IPv4.
- The sender wants to use IPv6, but the receiver does not understand IPv6.
- Tunneling does not work in this situation because the packet must be in the IPv4 format to be understood by the receiver.
- In this case, the header format must be totally changed through header translation.
- The header of the IPv6 packet is converted to an IPv4 header.

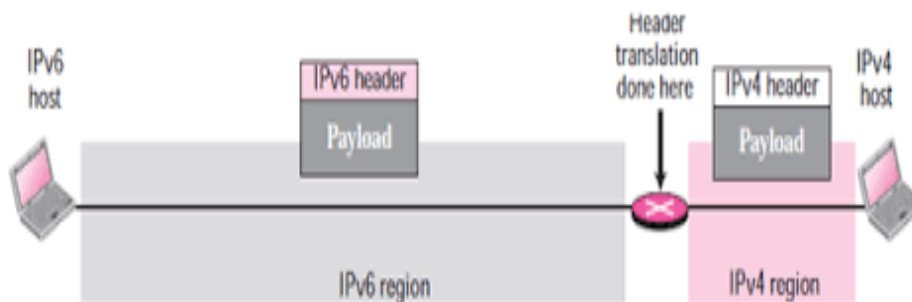


Figure 12.8. Header Translation

Use of IP Addresses

- During the transition a host may need to use two addresses, IPv4 and IPv6.
- When the transition is complete, IPv4 addresses should disappear.
- The DNS servers need to be ready to map a host name to either address type during the transition, but the IPv4 directory will disappear after all hosts in the world have migrated to IPv6.

12.2 LET US SUM UP

- We have seen that modified in version 6 of the TCP/IP protocol suite is ICMP.
- Also, we have understood Internet Control Message Protocol version 6 (ICMPv6)
- We have seen messages in ICMPv6 into four groups: error-reporting messages, informational messages, neighbor-discovery messages, and group-membership messages
- We have understood the concept transition from IPv4 TO IPv6
- Also, we have understood the Strategies in IPV6.

12.3 LIST OF REFERENCES

<https://ieeexplore.ieee.org/document/8912062>

<https://ieeexplore.ieee.org/document/4382215>

12.4 BIBLIOGRAPHY

1. Data Communications and Networking, Behrouz A. Forouzan, Fifth Edition, TMH, 2018.
2. Computer Network, Bhushan Trivedi, Oxford University Press, 2016

12.5 UNIT END EXERCISES

1. Explain the advantages of IPv6 when compared to IPv4.
2. List three protocols in the IPv4 network layer that are combined into a single protocol in IPv6.
3. Explain the benefit of renumbering.
4. What is tunneling strategy?
5. List categories of ICMPv6 messages.
6. List and explain extension header types.

INTRODUCTION TO THE TRANSPORT LAYER

Unit Structure

13.0 Objectives

13.1 Introduction

13.1.1 Transport Layer Protocol

13.1.2 User Datagram Protocol

13.1.3 Transmission Control Protocol

13.1.4 SCTP

13.2 Let us Sum Up

13.3 List of References

13.4 Bibliography

13.5 Unit End Exercises

13.0 OBJECTIVES

- After going through this unit, you will be able to:
- Understand services to the application layer and receives services from the network layer.
- Understand a process-to-process connection
- Understand end-to-end logical vehicle for transferring data
- Understand Internet transport-layer protocols
- Understand multiplexing and demultiplexing
- Understand connectionless and connection-oriented protocols

13.1 INTRODUCTION

The transport layer in the TCP/IP suite is located between the application layer and the network layer. It provides services to the application layer and receives services from the network layer. The transport layer acts as a liaison between a client program and a server program, a process-to-process connection.

The transport layer is the heart of the TCP/IP protocol suite; it is the end-to-end logical vehicle for transferring data from one point to another in the Internet. This chapter is the first chapter devoted to the transport layer.

13.1.1 Transport Layer Protocol

- The transport layer is located between the network layer and the application layer.
- The transport layer is responsible for providing services to the application layer; it receives services from the network layer.
- The services that can be provided by the transport layer.

Process-to-Process Communication

- The first duty of a transport-layer protocol is to provide process-to-process communication.
- A process is an application-layer entity (running program) that uses the services of the transport layer.
- Before we discuss how process-to-process communication can be accomplished, we need to understand the difference between host-to-host communication and process-to-process communication.
- The network layer is responsible for communication at the computer level (host-to-host communication).
- A network-layer protocol can deliver the message only to the destination computer.
- However, this is an incomplete delivery.
- The message still needs to be handed to the correct process.
- This is where a transport-layer protocol takes over.
- A transport-layer protocol is responsible for delivery of the message to the appropriate process.
- Figure 13.1 shows the domains of a network layer and a transport layer.

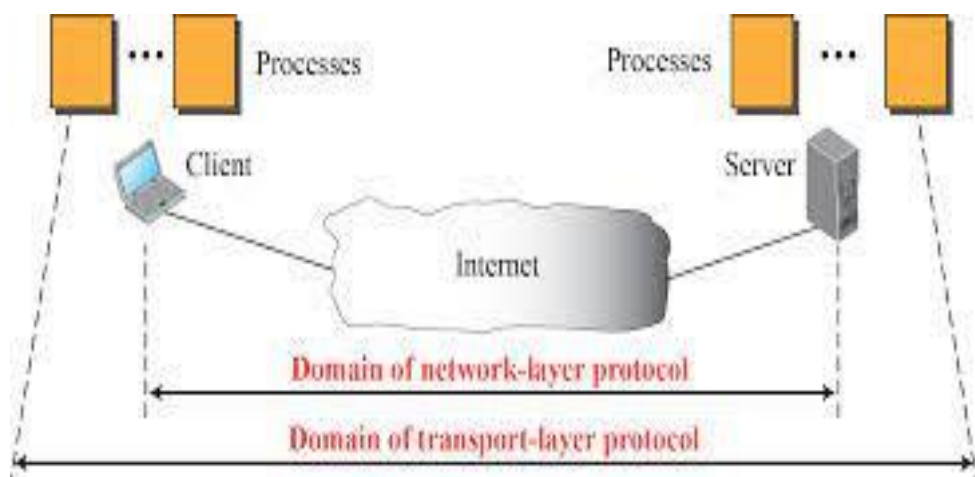


Figure 13.1 Network layer versus transport layer

- Although there are a few ways to achieve process-to-process communication, the most common is through the client-server paradigm.
- A process on the local host, called a client, needs services from a process usually on the remote host, called a server.
- However, operating systems today support both multiuser and multiprogramming environments.
- A remote computer can run several server programs at the same time, just as several local computers can run one or more client programs at the same time.
- For communication, we must define the local host, local process, remote host, and remote process.
- The local host and the remote host are defined using IP addresses.
- To define the processes, we need second identifiers, called port numbers.
- In the TCP/IP protocol suite, the port numbers are integers between 0 and 65,535 (16 bits).
- The client program defines itself with a port number, called the ephemeral port number.
- The word ephemeral means “short-lived” and is used because the life of a client is normally short.
- An ephemeral port number is recommended to be greater than 1023 for some client/server programs to work properly.
- The server process must also define itself with a port number.
- This port number, however, cannot be chosen randomly.
- If the computer at the server site runs a server process and assigns a random number as the port number, the process at the client site that wants to access that server and use its services will not know the port number.
- Of course, one solution would be to send a special packet and request the port number of a specific server, but this creates more overhead.
- TCP/IP has decided to use universal port numbers for servers; these are called well-known port numbers.
- There are some exceptions to this rule; for example, there are clients that are assigned well-known port numbers.
- Every client process knows the well-known port number of the corresponding server process.
- For example, while the daytime client process, a well-known client program, can use an ephemeral (temporary) port number, 52,000, to identify itself, the daytime server process must use the well-known (permanent) port number 13. Figure 13.2 shows this concept.

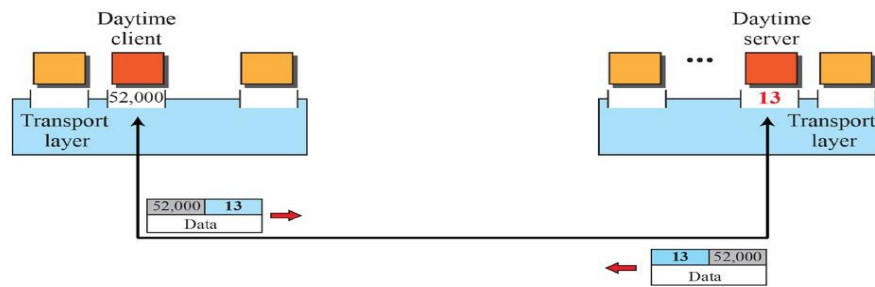


Figure 13.2: Port Numbers

- It should be clear by now that the IP addresses and port numbers play different roles in selecting the final destination of data.
- The destination IP address defines the host among the different hosts in the world.
- After the host has been selected, the port number defines one of the processes on this particular host as shown in figure 13.3.

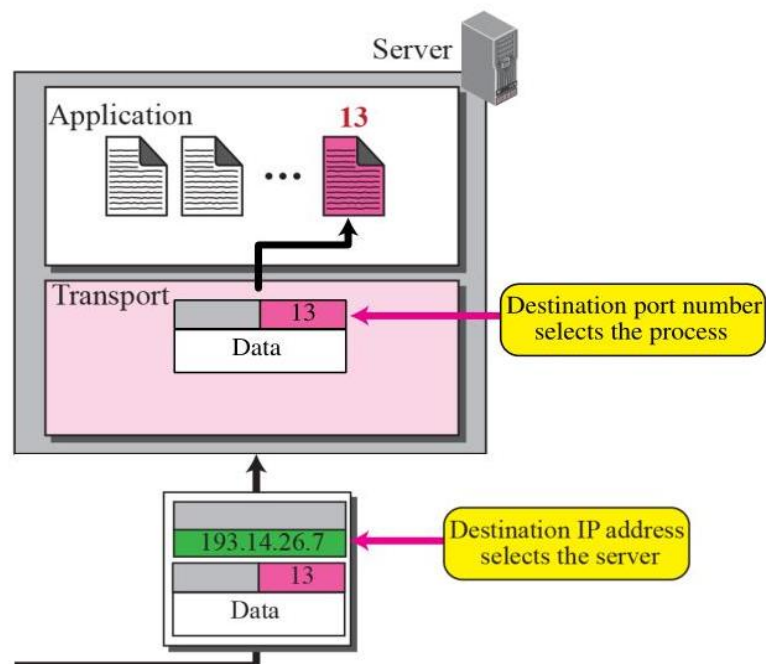
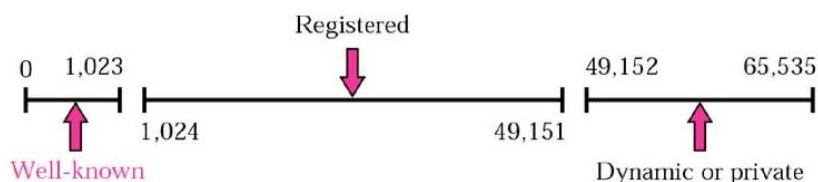


Figure 13.3 IP addresses versus port numbers

ICANN Ranges

ICANN has divided the port numbers into three ranges: well-known, registered, and dynamic (or private), as shown in Figure 13.4



- o **Well-known ports.** The ports ranging from 0 to 1023 are assigned and controlled by ICANN. These are the well-known ports.
- o **Registered ports.** The ports ranging from 1024 to 49,151 are not assigned or controlled by ICANN. They can only be registered with ICANN to prevent duplication.
- o **Dynamic ports.** The ports ranging from 49,152 to 65,535 are neither controlled nor registered. They can be used as temporary or private port numbers.

Socket Addresses:

- A transport-layer protocol in the TCP suite needs both the IP address and the port number, at each end, to make a connection.
- The combination of an IP address and a port number is called a socket address.
- The client socket address defines the client process uniquely just as the server socket address defines the server process uniquely as shown in Figure 13.5)

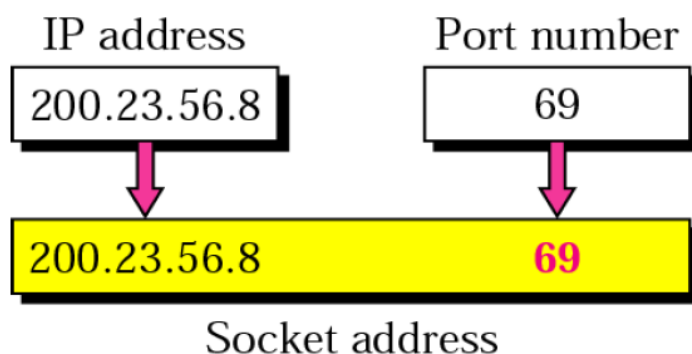


Figure 13.5 Socket address

- To use the services of the transport layer in the Internet, we need a pair of socket addresses: the client socket address and the server socket address.
- These four pieces of information are part of the network-layer packet header and the transport-layer packet header.
- The first header contains the IP addresses; the second header contains the port numbers

Encapsulation and Decapsulation

- To send a message from one process to another, the transport-layer protocol encapsulates and decapsulates messages (Figure 13.6).
- Encapsulation happens at the sender site.
- When a process has a message to send, it passes the message to the transport layer along with a pair of socket addresses and some other pieces of information, which depend on the transport-layer protocol.
- The transport layer receives the data and adds the transport-layer header.
- The packets at the transport layer in the Internet are called user datagrams, segments, or packets, depending on what transport-layer protocol we use.
- In general discussion, we refer to transport-layer payloads as packets.
- Decapsulation happens at the receiver site.
- When the message arrives at the destination transport layer, the header is dropped and the transport layer delivers the message to the process running at the application layer.
- The sender socket address is passed to the process in case it needs to respond to the message received.

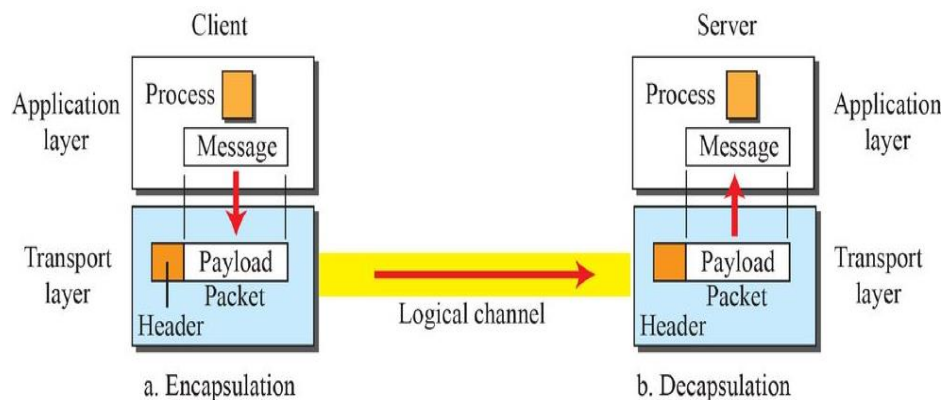


Figure 13.6: Encapsulates and decapsulates messages

Multiplexing and Demultiplexing

- Whenever an entity accepts items from more than one source, this is referred to as multiplexing (many to one); whenever an entity delivers items to more than one source, this is referred to as demultiplexing (one to many).
- The transport layer at the source performs multiplexing; the transport layer at the destination performs demultiplexing (Figure 13.7).

- Figure 13.7 shows communication between a client and two servers.
- Three client processes are running at the client site, P1, P2, and P3.
- The processes P1 and P3 need to send requests to the corresponding server process running in a server.
- The client process P2 needs to send a request to the corresponding server process running at another server.
- The transport layer at the client site accepts three messages from the three processes and creates three packets.
- It acts as a multiplexer.
- The packets 1 and 3 use the same logical channel to reach the transport layer of the first server.
- When they arrive at the server, the transport layer does the job of a demultiplexer and distributes the messages to two different processes.
- The transport layer at the second server receives packet 2 and delivers it to the corresponding process.
- Note that we still have demultiplexing although there is only one message.

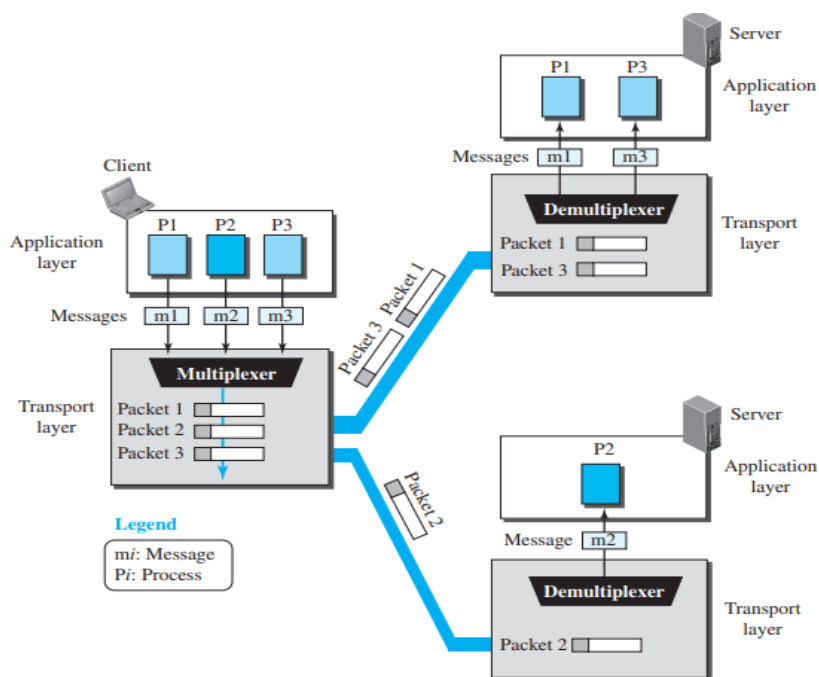


Figure 13.7 Multiplexing and demultiplexing

Flow Control

- Whenever an entity produces items and another entity consumes them, there should be a balance between production and consumption rates.

- If the items are produced faster than they can be consumed, the consumer can be overwhelmed and may need to discard some items.
- If the items are produced more slowly than they can be consumed, the consumer must wait, and the system becomes less efficient.
- Flow control is related to the first issue.
- We need to prevent losing the data items at the consumer site.

Pushing or Pulling

- Delivery of items from a producer to a consumer can occur in one of two ways: pushing or pulling.
- If the sender delivers items whenever they are produced without a prior request from the consumer the delivery is referred to as pushing.
- If the producer delivers the items after the consumer has requested them, the delivery is referred to as pulling.
- Figure 13.8 shows these two types of delivery.

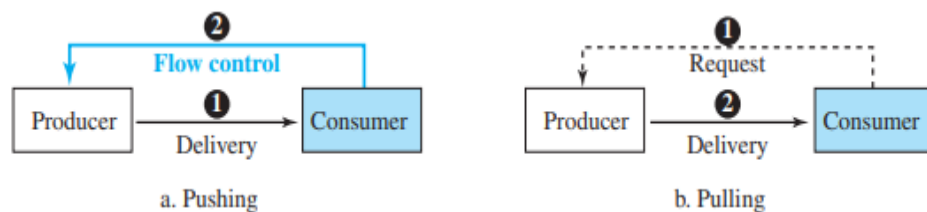


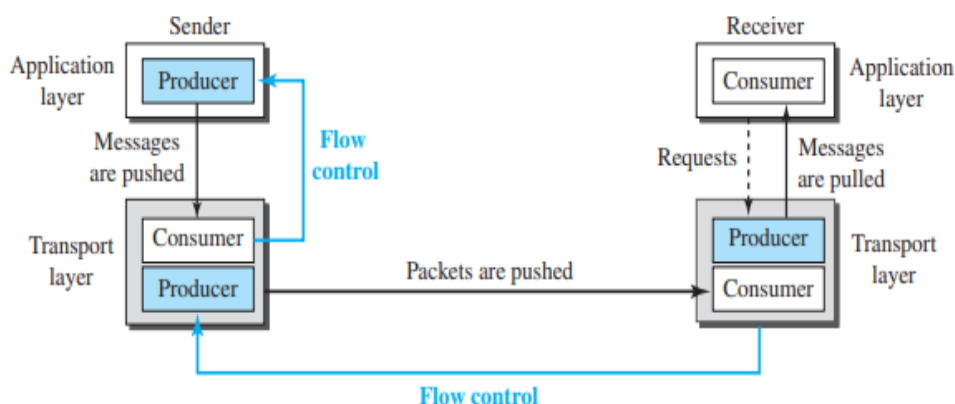
Figure 13.8: Pushing or pulling

- When the producer pushes the items, the consumer may be overwhelmed and there is a need for flow control, in the opposite direction, to prevent discarding of the items.
- In other words, the consumer needs to warn the producer to stop the delivery and to inform the producer when it is again ready to receive the items.
- When the consumer pulls the items, it requests them when it is ready.
- In this case, there is no need for flow control.

Flow Control at Transport Layer

- In communication at the transport layer, we are dealing with four entities: sender process, sender transport layer, receiver transport layer, and receiver process.
- The sending process at the application layer is only a producer.
- It produces message chunks and pushes them to the transport layer.

- The sending transport layer has a double role: it is both a consumer and a producer.
- It consumes the messages pushed by the producer.
- It encapsulates the messages in packets and pushes them to the receiving transport layer.
- The receiving transport layer also has a double role: it is the consumer for the packets received from the sender and the producer that decapsulates the messages and delivers them to the application layer.
- The last delivery, however, is normally a pulling delivery; the transport layer waits until the application-layer process asks for messages.
- Figure 13.9 shows that we need at least two cases of flow control: from the sending transport layer to the sending application layer and from the receiving transport layer to the sending transport layer.



Error Control

- In the Internet, since the underlying network layer (IP) is unreliable, we need to make the transport layer reliable if the application requires reliability.
- Reliability can be achieved to add error control services to the transport layer.
- Error control at the transport layer is responsible for:
 1. Detecting and discarding corrupted packets.
 2. Keeping track of lost and discarded packets and resending them.
 3. Recognizing duplicate packets and discarding them.
 4. Buffering out-of-order packets until the missing packets arrive.

- Error control, unlike flow control, involves only the sending and receiving transport layers.
- We are assuming that the message chunks exchanged between the application and transport layers are error free.
- Figure 23.10 shows the error control between the sending and receiving transport layers.
- As with the case of flow control, the receiving transport layer manages error control, most of the time, by informing the sending transport layer about the problems.

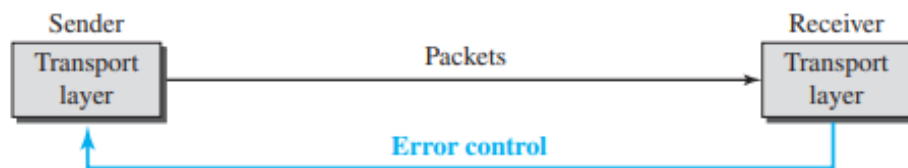


Figure 13.10 Error Control

Congestion Control

- An important issue in a packet-switched network, such as the Internet, is congestion.
- Congestion in a network may occur if the load on the network—the number of packets sent to the network—is greater than the capacity of the network—the number of packets a network can handle.
- Congestion control refers to the mechanisms and techniques that control the congestion and keep the load below the capacity.
- We may ask why there is congestion in a network.
- Congestion happens in any system that involves waiting.
- For example, congestion happens on a freeway because any abnormality in the flow, such as an accident during rush hour, creates blockage.
- Congestion in a network or internetwork occurs because routers and switches have queues—buffers that hold the packets before and after processing.
- A router, for example, has an input queue and an output queue for each interface.
- If a router cannot process the packets at the same rate at which they arrive, the queues become overloaded and congestion occurs.
- Congestion at the transport layer is actually the result of congestion at the network layer, which manifests itself at the transport layer.

- In a connectionless service, the source process (application program) needs to divide its message into chunks of data of the size acceptable by the transport layer and deliver them to the transport layer one by one.
- The transport layer treats each chunk as a single unit without any relation between the chunks.
- When a chunk arrives from the application layer, the transport layer encapsulates it in a packet and sends it.
- To show the independency of packets, assume that a client process has three chunks of messages to send to a server process.
- The chunks are handed over to the connectionless transport protocol in order.
- However, since there is no dependency between the packets at the transport layer, the packets may arrive out of order at the destination and will be delivered out of order to the server process (Figure 13.11).

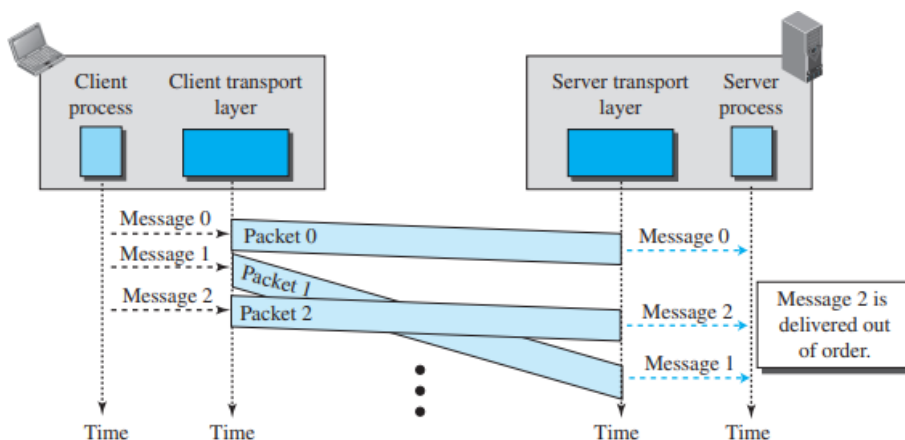


Figure 13.11 Connectionless Service

Connection-Oriented Service

- In a connection-oriented service, the client and the server first need to establish a logical connection between themselves.
- The data exchange can only happen after the connection establishment.
- After data exchange, the connection needs to be torn down (Figure 13.12).
- As we mentioned before, the connection-oriented service at the transport layer is different from the same service at the network layer.

- In the network layer, connection oriented service means a coordination between the two end hosts and all the routers in between.
- At the transport layer, connection-oriented service involves only the two hosts; the service is end to end.
- This means that we should be able to make a connection-oriented protocol at the transport layer over either a connectionless or connection-oriented protocol at the network layer.
- Figure 13.11 shows the connection establishment, data-transfer, and tear-down phases in a connection-oriented service at the transport layer.

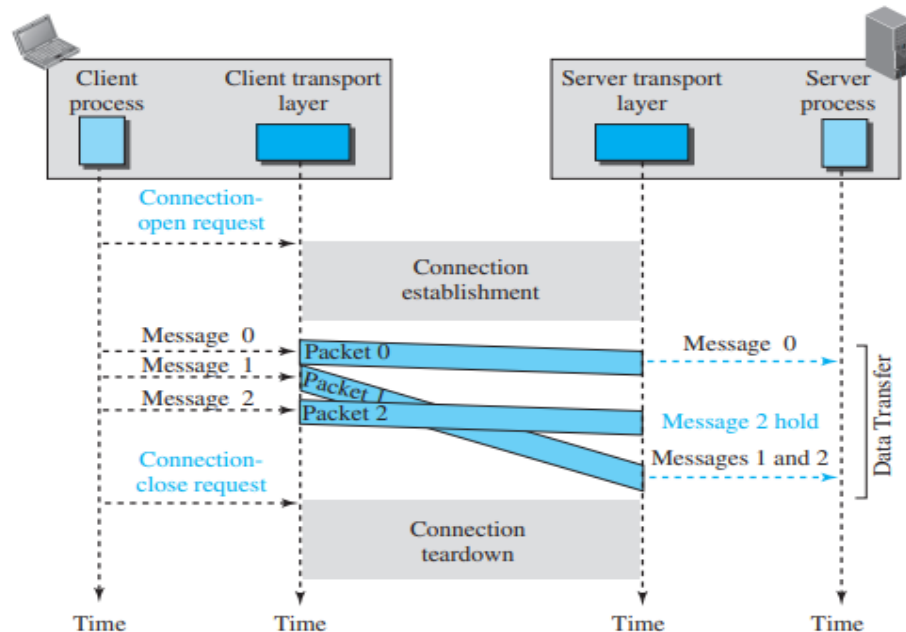


Figure 13.11 Connection oriented

13.1.2 USER DATAGRAM PROTOCOL

- The User Datagram Protocol (UDP) is a connectionless, unreliable transport protocol.
- It does not add anything to the services of IP except for providing process-to-process communication instead of host-to-host communication.
- If UDP is so powerless, why would a process want to use it?
- UDP is a very simple protocol using a minimum of overhead.
- If a process wants to send a small message and does not care much about reliability, it can use UDP.
- Sending a small message using UDP takes much less interaction between the sender and receiver than using TCP.

User Datagram

- UDP packets, called user datagrams, have a fixed-size header of 8 bytes made of four fields, each of 2 bytes (16 bits).
- Figure 13.12 shows the format of a user datagram.
- The first two fields define the source and destination port numbers.
- The third field 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 less because a UDP user datagram is stored in an IP datagram with the total length of 65,535 bytes.
- The last field can carry the optional checksum.

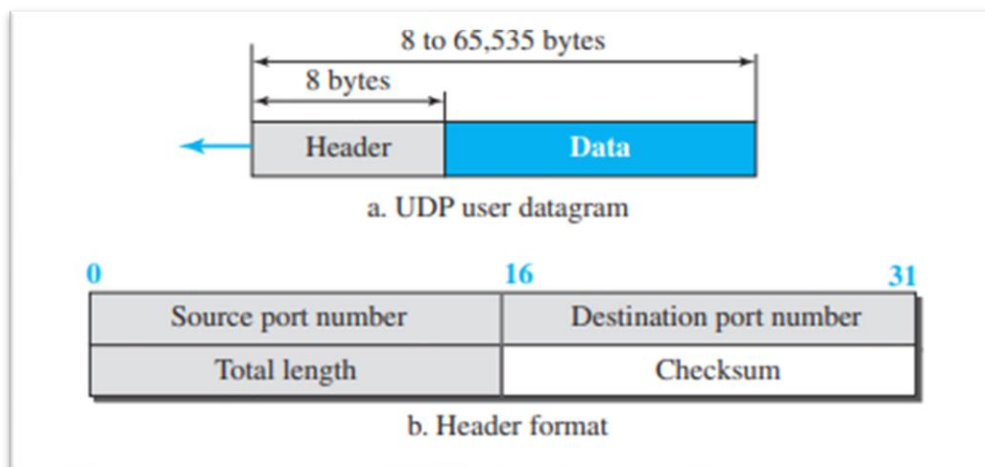


Figure 13.12 UDP datagram

UDP Services

- General services are provided by UDP

Process-to-Process Communication:

- UDP provides process-to-process communication using socket addresses, a combination of IP addresses and port numbers.

Connectionless Services:

- UDP provides a connectionless service.
- This means that each user datagram sent by UDP is an independent datagram.

Flow Control:

- UDP is a very simple protocol.

- There is no flow control, and hence no window mechanism.
- The receiver may overflow with incoming messages.
- The lack of flow control means that the process using UDP should provide for this service, if needed.

Error Control:

- 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

Checksum:

- UDP checksum calculation includes three sections: a pseudo header, the UDP header, and the data coming from the application layer.
- The pseudo header is the part of the header of the IP packet in which the user datagram is to be encapsulated with some fields filled with 0s.

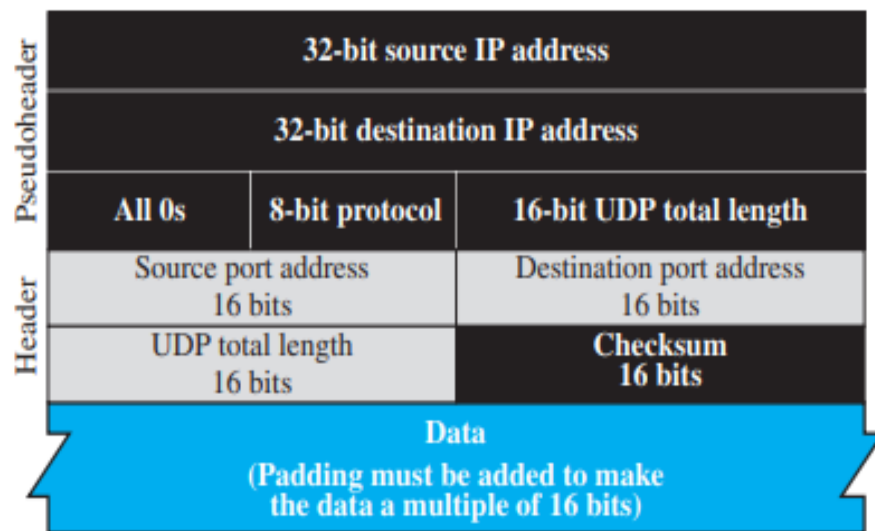


Figure 13.13 Pseudo header

13.1.3 TRANSMISSION CONTROL PROTOCOL

Transmission Control Protocol (TCP) is a connection-oriented, reliable protocol. TCP explicitly defines connection establishment, data transfer, and connection tear down phases to provide a connection-oriented service.

TCP Services

Process-to-Process Communication:

- As with UDP, TCP provides process-to-process communication using port numbers.

Stream Delivery Service:

- TCP, unlike UDP, is a stream-oriented protocol.
- In UDP, a process sends messages with predefined boundaries to UDP for delivery.
- UDP adds its own header to each of these messages and delivers it 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 bytes across the Internet.
- This imaginary environment is depicted in Figure 13.14.
- The sending process produces (writes to) the stream and the receiving process consumes (reads from) it.

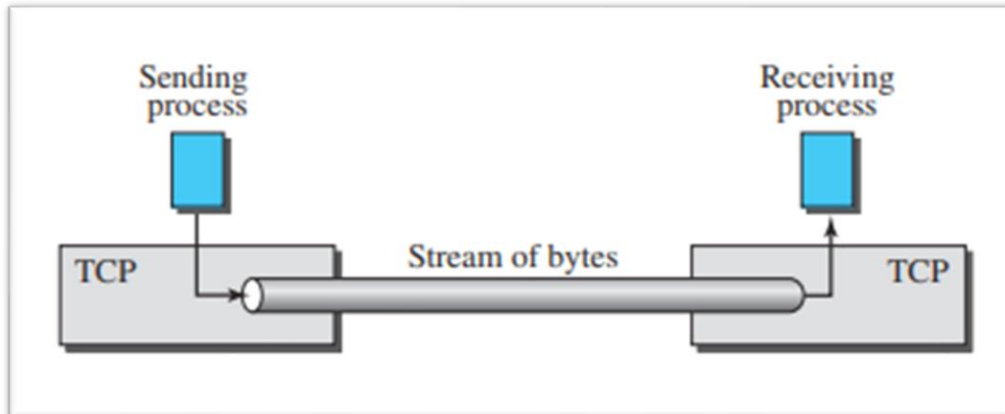


Figure 13.14 Stream Delivery

Full-Duplex Communication

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

Multiplexing and Demultiplexing

- Like UDP, TCP performs multiplexing at the sender and demultiplexing at the receiver.
- However, since TCP is a connection-oriented protocol, a connection needs to be established for each pair of processes.

Connection-Oriented Service

- TCP, unlike UDP, is a connection-oriented protocol.
- When a process at site A wants to send to and receive data from another process at site B, the following three phases occur:
 1. The two TCPs establish a logical connection between them.
 2. Data are exchanged in both directions.
 3. The connection is terminated.

Note that this is a logical connection, not a physical connection.

Reliable Service

- TCP is a reliable transport protocol.
- It uses an acknowledgment mechanism to check the safe and sound arrival of data.

13.1.4 SCTP

Stream Control Transmission Protocol

(SCTP) is a new transport-layer protocol designed to combine some features of UDP and TCP in an effort to create a better protocol for multimedia communication.

SCTP Services

Services offered by SCTP to the application-layer processes.

Process-to-Process Communication

SCTP, like UDP or TCP, provides process-to-process communication.

Multiple Streams We learned that TCP is a stream-oriented protocol.

- Each connection between a TCP client and a TCP server involves a single stream.
- The problem with this approach is that a loss at any point in the stream blocks the delivery of the rest of the data.
- This can be acceptable when we are transferring text; it is not when we are sending real-time data such as audio or video.
- SCTP allows multistream service in each connection, which is called association in SCTP terminology.

- If one of the streams is blocked, the other streams can still deliver their data.
- Figure 13.15 shows the idea of multiple-stream delivery.

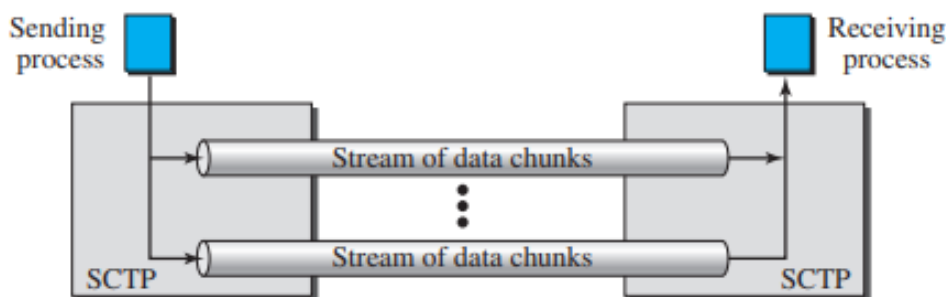


Figure 13.15 Multiple-stream concept

Multihoming

- A TCP connection involves one source and one destination IP address.
- This means that even if the sender or receiver is a multihomed host (connected to more than one physical address with multiple IP addresses), only one of these IP addresses per end can be utilized during the connection.
- An SCTP association, on the other hand, supports multihoming service.
- Figure 13.16 shows the idea of multihoming.
- In the figure, the client is connected to two local networks with two IP addresses.
- The server is also connected to two networks with two IP addresses.
- The client and the server can make an association using four different pairs of IP addresses.
- However, note that in the current implementations of SCTP, only one pair of IP addresses can be chosen for normal communication; the alternative is used if the main choice fails.
- In other words, at present, SCTP does not allow load sharing between different paths.

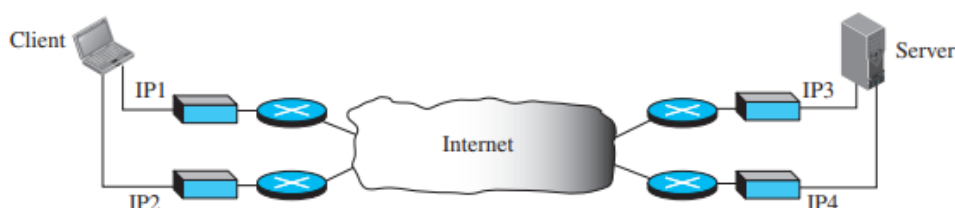


Figure 13.16: Multihoming concept

Full-Duplex Communication

- Like TCP, SCTP offers full-duplex service, where data can flow in both directions at the same time.
- Each SCTP then has a sending and receiving buffer and packets are sent in both directions.

Connection-Oriented Service

- Like TCP, SCTP is a connection-oriented protocol.
- However, in SCTP, a connection is called an association.

Reliable Service

- SCTP, like 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

SCTP Features

The following shows the general features of SCTP.

Transmission Sequence Number (TSN)

- The unit of data in SCTP is a data chunk, which may or may not have a one-to-one relationship with the message coming from the process because of fragmentation.
- Data transfer in SCTP is controlled by numbering the data chunks.
- SCTP uses a transmission sequence number (TSN) to number the data chunks.
- In other words, the TSN in SCTP plays a role analogous to the sequence number in TCP.
- TSNs are 32 bits long and randomly initialized between 0 and $2^{32} - 1$.
- Each data chunk must carry the corresponding TSN in its header.

Stream Identifier (SI)

- In SCTP, there may be several streams in each association.
- Each stream in SCTP needs to be identified using a stream identifier (SI).
- Each data chunk must carry the SI in its header so that when it arrives at the destination, it can be properly placed in its stream.
- The SI is a 16-bit number starting from 0.

- When a data chunk arrives at the destination SCTP, it is delivered to the appropriate stream and in the proper order.
- This means that, in addition to an SI, SCTP defines each data chunk in each stream with a stream sequence number (SSN).

Packets

- In TCP, a segment carries data and control information.
- Data are carried as a collection of bytes; control information is defined by six control flags in the header.
- The design of SCTP is totally different: data are carried as data chunks, control information as control chunks.
- Several control chunks and data chunks can be packed together in a packet.
- A packet in SCTP plays the same role as a segment in TCP.
- Figure 13.17 compares a segment in TCP and a packet in SCTP.

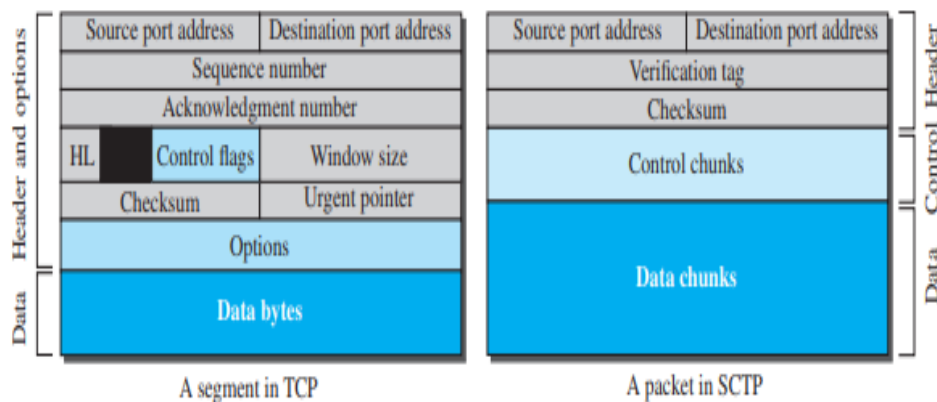


Figure 13.17 segment in TCP and SCTP

Acknowledgment Number

- TCP acknowledgment numbers are byte-oriented and refer to the sequence numbers.
- SCTP acknowledgment numbers are chunk-oriented.
- They refer to the TSN.
- A second difference between TCP and SCTP acknowledgments is the control information.
- Recall that this information is part of the segment header in TCP.
- To acknowledge segments that carry only control information.

An SCTP Association

SCTP, like TCP, is a connection-oriented protocol. However, a connection in SCTP is called an association to emphasize multihoming.

Association Establishment

- Association establishment in SCTP requires a four-way handshake.
- In this procedure, a process, normally a client, wants to establish an association with another process, normally a server, using SCTP as the transport-layer protocol.
- Similar to TCP, the SCTP server needs to be prepared to receive any association (passive open).
- Association establishment, however, is initiated by the client (active open).
- SCTP association establishment is shown in Figure 13.18.

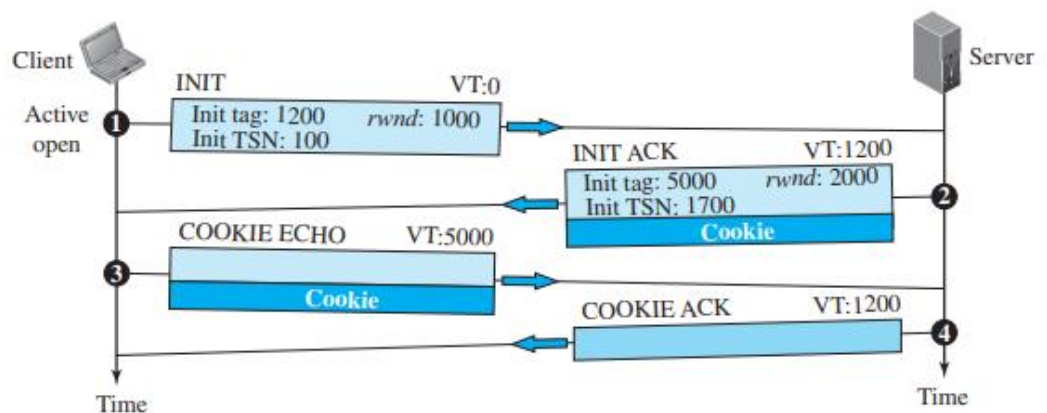


Figure 13.18 SCTP association establishment

Data Transfer:

- The whole purpose of an association is to transfer data between two ends.
- After the association is established, bidirectional data transfer can take place.
- The client and the server can both send data. Like TCP, SCTP supports piggybacking.
- There is a major difference, however, between data transfer in TCP and SCTP.
- TCP receives messages from a process as a stream of bytes without recognizing any boundary between them.
- The process may insert some boundaries for its peer use, but TCP treats that mark as part of the text.

- In other words, TCP takes each message and appends it to its buffer.
- A segment can carry parts of two different messages.
- The only ordering system imposed by TCP is the byte numbers.

Association Termination

- In SCTP, like TCP, either of the two parties involved in exchanging data (client or server) can close the connection.
- However, unlike TCP, SCTP does not allow a “half closed” association.
- If one end closes the association, the other end must stop sending new data.
- If any data are left over in the queue of the recipient of the termination request, they are sent and the association is closed.
- Association termination uses three packets, as shown in Figure 13.19.

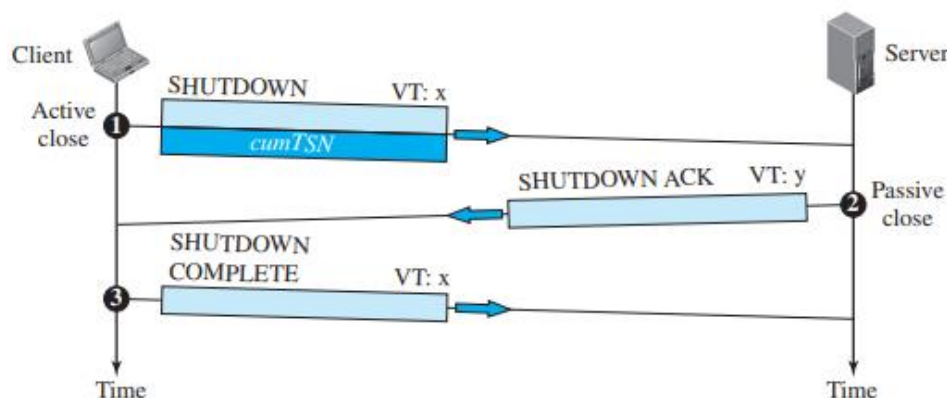


Figure 13.19 Association termination

13.2 LET US SUM UP

- UDP is an unreliable and connectionless transport-layer protocol that creates a process-to-process communication, which means it requires little overhead and offers fast delivery.
- The UDP packet is called a user datagram.
- UDP has no flow- or error-control mechanism; its only attempt at error control is the checksum.
- A user datagram is encapsulated in the data field of an IP datagram.
- UDP uses multiplexing and demultiplexing to handle outgoing and incoming user datagrams.
- Transmission Control Protocol (TCP) is another transport-layer protocol in the TCP/IP protocol suite.

- It provides process-to-process, full-duplex, and connection-oriented service.
- A TCP connection consists of three phases: connection establishment, data transfer, and connection termination.
- SCTP is a message-oriented, reliable protocol that combines the good features of UDP and TCP.

13.3 LIST OF REFERENCES

1. <https://ieeexplore.ieee.org/abstract/document/8786240>
2. <https://ieeexplore.ieee.org/abstract/document/1457052>

13.4 BIBLIOGRAPHY

1. Data Communications and Networking, Behrouz A. Forouzan, Fifth Edition, TMH, 2018.
2. Computer Network, Bhushan Trivedi, Oxford University Press, 2016

13.5 UNIT END EXERCISES

1. What is the source port number?
2. Compare the TCP header and the UDP header.
3. What is the SSN?
4. List and explain features of SCTP.
5. List and explain services of TCP.
6. Compare TCP and UDP.

INTRODUCTION TO THE APPLICATION LAYER

Unit Structure

14.0 Objectives

14.1 Introduction

14.1.1 Client Server Programming

14.1.2 Iterative Programming.

14.2 Let us Sum Up

14.3 List of References

14.4 Bibliography

14.5 Unit End Exercises

14.0 OBJECTIVES

After going through this unit, you will be able to:

- Understand the services provided by the application layer.
- Understand client-server programming
- Understand one-to-many delivery.
- Understand socket programming in C.
- Understand application programming interfaces
- Understand iterative client-server programs

14.1 INTRODUCTION

The application layer provides services to the user. Communication is provided using a logical connection, which means that the two application layers assume that there is an imaginary direct connection through which they can send and receive messages. The communication at the application layer is logical, not physical.

14.1.1 Client Server Programming

- In a client-server paradigm, communication at the application layer is between two running application programs called processes: a client and a server.
- A client is a running program that initializes the communication by sending a request; a server is another application program that waits for a request from a client.

- The server handles the request received from a client, prepares a result, and sends the result back to the client.
- This definition of a server implies that a server must be running when a request from a client arrives, but the client needs to be run only when it is needed.
- This means that if we have two computers connected to each other somewhere, we can run a client process on one of them and the server on the other.
- However, we need to be careful that the server program is started before we start running the client program.
- In other words, the lifetime of a server is infinite: it should be started and run forever, waiting for the clients.
- The lifetime of a client is finite: it normally sends a finite number of requests to the corresponding server, receives the responses, and stops.

Application Programming Interface

- A computer program is normally written in a computer language with a predefined set of instructions that tells the computer what to do.
- A computer language has a set of instructions for mathematical operations, a set of instructions for string manipulation, a set of instructions for input/ output access, and so on.
- If we need a process to be able to communicate with another process, we need a new set of instructions to tell the lowest four layers of the TCP/IP suite to open the connection, send and receive data from the other end, and close the connection.
- A set of instructions of this kind is normally referred to as an application programming interface (API).
- An interface in programming is a set of instructions between two entities.
- In this case, one of the entities is the process at the application layer and the other is the operating system that encapsulates the first four layers of the TCP/IP protocol suite.
- In other words, a computer manufacturer needs to build the first four layers of the suite in the operating system and include an API.
- In this way, the processes running at the application layer are able to communicate with the operating system when sending and receiving messages through the Internet.
- Several APIs have been designed for communication.
- Three among them are common: socket interface, Transport Layer Interface (TLI), and STREAM.

- The socket interface is a set of instructions that provide communication between the application layer and the operating system, as shown in Figure 14.1.
- It is a set of instructions that can be used by a process to communicate with another process.
- The idea of sockets allows us to use the set of all instructions already designed in a programming language for other sources and sinks.
- For example, in most computer languages, like C, C++, or Java, we have several instructions that can read and write data to other sources and sinks such as a keyboard (a source), a monitor (a sink), or a file (source and sink).
- without changing the way, we send data or receive data.

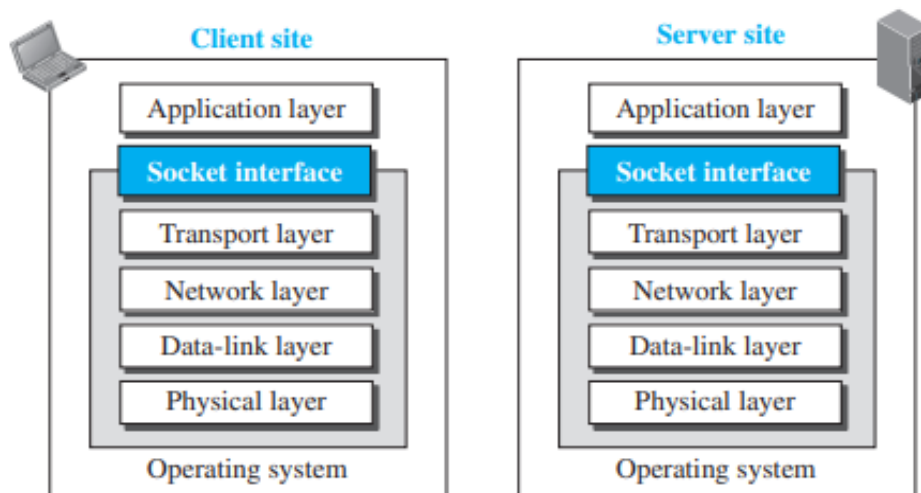


Figure 14.1 Position of the socket interface

Figure 14.2 shows the idea and compares the sockets with other sources and sinks.

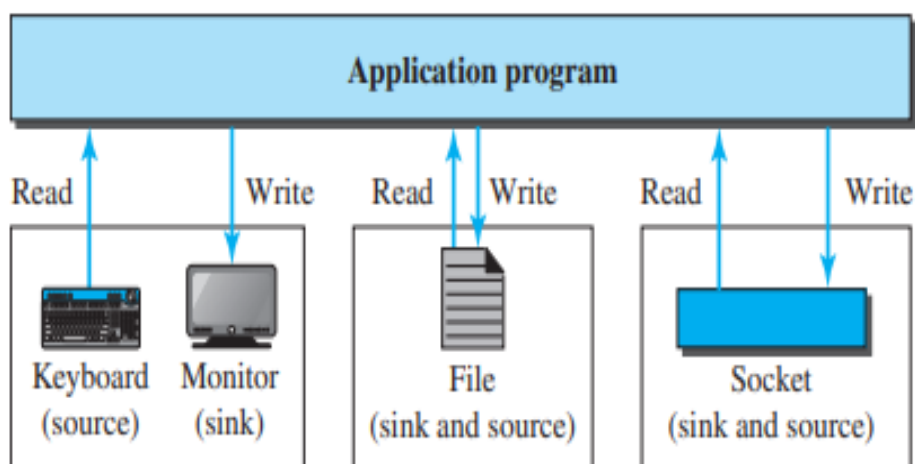


Figure 14.2 Sockets used the same way as other sources and sinks

Sockets

- Although a socket is supposed to behave like a terminal or a file, it is not a physical entity like them; it is an abstraction.
- It is an object that is created and used by the application program.
- We can say that, as far as the application layer is concerned, communication between a client process and a server process is communication between two sockets, created at two ends, as shown in Figure 14.3.
- The client thinks that the socket is the entity that receives the request and gives the response; the server thinks that the socket is the one that has a request and needs the response.
- If we create two sockets, one at each end, and define the source and destination addresses correctly, we can use the available instructions to send and receive data.
- The rest is the responsibility of the operating system and the embedded TCP/IP protocol

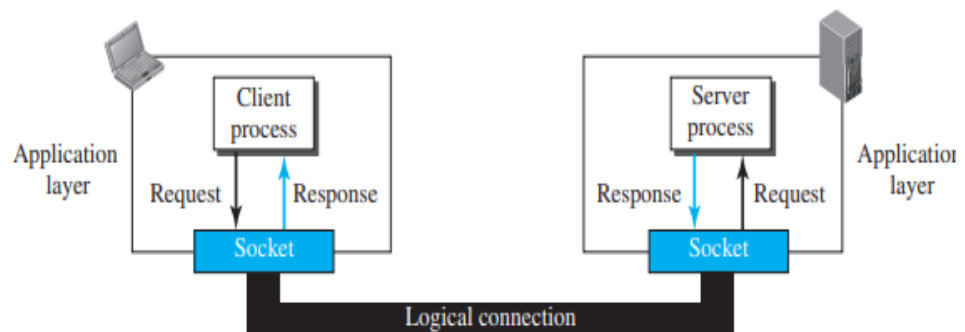


Figure 14.3 Use of sockets in process-to-process communication

Socket Addresses

- The interaction between a client and a server is two-way communication.
- In a two-way communication, we need a pair of addresses: local (sender) and remote (receiver).
- The local address in one direction is the remote address in the other direction and vice versa.
- Since communication in the client-server paradigm is between two sockets, we need a pair of socket addresses for communication: a local socket address and a remote socket address.
- However, we need to define a socket address in terms of identifiers used in the TCP/IP protocol suite.

- A socket address should be a combination of an IP address and a port number as shown in Figure 14.4.
- Since a socket defines the end-point of the communication, we can say that a socket is identified by a pair of socket addresses, a local and a remote.

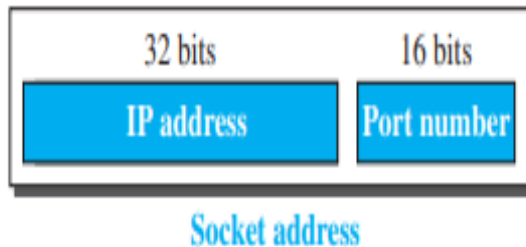


Figure 14.4 A Socket address

Finding Socket Addresses

- How can a client or a server find a pair of socket addresses for communication?
- The situation is different for each site.

Server Site

- The server needs a local (server) and a remote (client) socket address for communication.

Local Socket Address

- The local (server) socket address is provided by the operating system.
- The operating system knows the IP address of the computer on which the server process is running.
- The port number of a server process, however, needs to be assigned.
- If the server process is a standard one defined by the Internet authority, a port number is already assigned to it.

Remote Socket Address

- The remote socket address for a server is the socket address of the client that makes the connection.
- Since the server can serve many clients, it does not know beforehand the remote socket address for communication.
- The server can find this socket address when a client tries to connect to the server.

- The client socket address, which is contained in the request packet sent to the server, becomes the remote socket address that is used for responding to the client.
- In other words, although the local socket address for a server is fixed and used during its lifetime, the remote socket address is changed in each interaction with a different client.

Client Site

The client also needs a local (client) and a remote (server) socket address for communication.

- Local Socket Address
- The local (client) socket address is also provided by the operating system.
- The operating system knows the IP address of the computer on which the client is running.
- The port number, however, is a 16-bit temporary integer that is assigned to a client process each time the process needs to start the communication.
- integers defined by the Internet authority and called the ephemeral (temporary) port numbers.
- The operating system, however, needs to guarantee that the new port number is not used by any other running client process.
- The operating system needs to remember the port number to be able to redirect the response received from the server process to the client process that sent the request.

Remote Socket Address

- Finding the remote (server) socket address for a client, however, needs more work.
- When a client process starts, it should know the socket address of the server it wants to connect to.
- We will have two situations in this case
- Sometimes, the user who starts the client process knows both the server port number and IP address of the computer on which the server is running.
- This usually occurs in situations when we have written client and server applications and we want to test them.
- For example, at the end of this chapter we write some simple client and server programs and we test them using this approach.

- In this situation, the programmer can provide these two pieces of information when he runs the client program.
- Although each standard application has a well-known port number, most of the time, we do not know the IP address.
- This happens in situations such as when we need to contact a web page, send an e-mail to a friend, copy a file from a remote site, and so on.
- In these situations, the server has a name, an identifier that uniquely defines the server process.

Iterative Communication Using UDP

- Communication between a client program and a server program can occur iteratively or concurrently.
- Although several client programs can access the same server program at the same time, the server program can be designed to respond iteratively or concurrently.
- An iterative server can process one client request at a time; it receives a request, processes it, and sends the response to the requestor before handling another request.
- When the server is handling the request from a client, the requests from other clients, and even other requests from the same client, need to be queued at the server site and wait for the server to be freed.
- The received and queued requests are handled in the first-in, first-out fashion.
- In this section, we discuss iterative communication using UDP.

Sockets Used for UDP

- In UDP communication, the client and server use only one socket each.
- The socket created at the server site lasts forever; the socket created at the client site is closed (destroyed) when the client process terminates.
- Figure 14.5 shows the lifetime of the sockets in the server and client processes.
- In other words, different clients use different sockets, but the server creates only one socket and changes only the remote socket address each time a new client makes a connection.
- This is logical, because the server does know its own socket address, but does not know the socket addresses of the clients who need its services; it needs to wait for the client to connect before filling this part of the socket address.

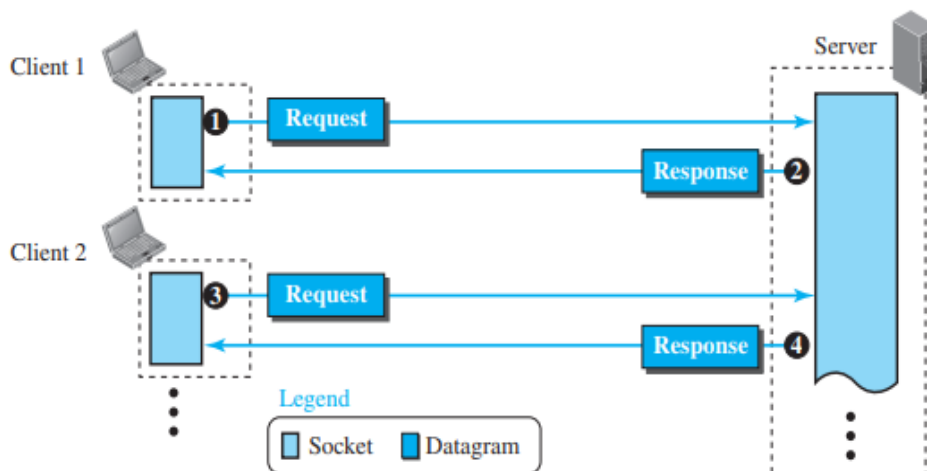


Figure 14.5 Sockets for UDP communication

Sockets Used in TCP

The TCP server uses two different sockets, one for connection establishment and the other for data transfer.

We call the first one the listen socket and the second the socket.

The reason for having two types of sockets is to separate the connection phase from the data exchange phase.

A server uses a listen socket to listen for a new client trying to establish connection.

After the connection is established, the server creates a socket to exchange data with the client and finally to terminate the connection.

The client uses only one socket for both connection establishment and data exchange (see Figure 14.6).

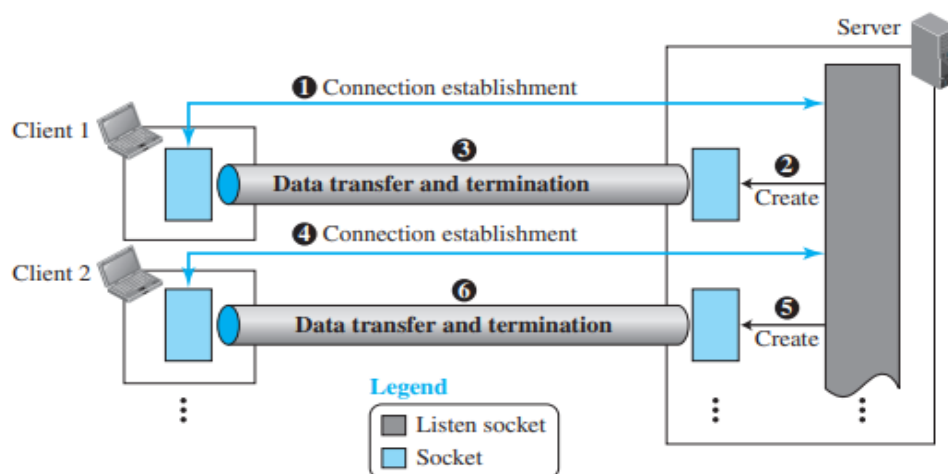


Figure 14.6 Sockets used in TCP communication

Concurrent Communication

- A concurrent server can process several client requests at the same time.
- This can be done using the available provisions in the underlying programming language.
- In C, a server can create several child processes, in which a child can handle a client.
- In Java, threading allows several clients to be handled by each thread.

14.2 ITERATIVE PROGRAMMING

ITERATIVE PROGRAMMING IN C

- Some simple iterative client-server programs using C, a procedural programming language.
- Socket programming traditionally started in the C language.
- The low-level feature of the C language better reveals some subtleties in this type of programming

General Issues

- The important issue in socket interface is to understand the role of a socket in communication.
- The socket has no buffer to store data to be sent or received.
- It is capable of neither sending nor receiving data.
- The socket just acts as a reference or a label.
- The buffers and necessary variables are created inside the operating system.

Socket Structure in C

- The C language defines a socket as a structure (struct).
- The socket structure is made of five fields; each socket address itself is a structure made of five fields, as shown in Figure 14.7.
- Note that the programmer should not redefine this structure; it is already defined in the header files.
- We briefly discuss the five fields in a socket structure.

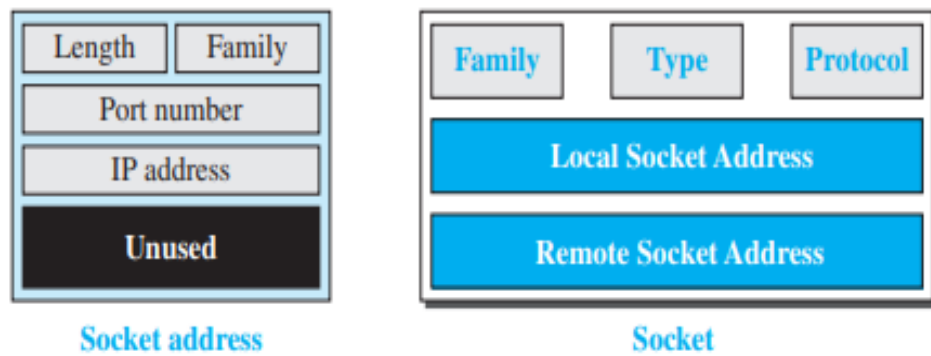


Figure 14.7 Socket data structure

Family:

- This field defines the family protocol (how to interpret the addresses and port number).
- The common values are `PF_INET` (for current Internet), `PF_INET6` (for next-generation Internet), and so on. We use `PF_INET` for this section.

Type

- This field defines four types of sockets: `SOCK_STREAM` (for TCP), `SOCK_DGRAM` (for UDP), `SOCK_SEQPACKET` (for SCTP), and `SOCK_RAW` (for applications that directly use the services of IP).

Protocol

- This field defines the specific protocol in the family.
- It is set to 0 for TCP/IP protocol suite because it is the only protocol in the family.

Local socket address.

- This field defines the local socket address.
- A socket address is itself a structure made of the length field, the family field (which is set to the constant `AF_INET` for TCP/IP protocol suite), the port number field (which defines the process), and the IP address field (which defines the host on which the process is running).
- It also contains an unused field.

Remote socket address

- This field defines the remote socket address.
- Its structure is the same as the local socket address.

Header Files

- To be able to use the definition of the socket and all procedures (functions) defined in the interface, we need a set of header files.
- We have collected all of these header files in a file named header Files .h.
- This file needs to be created in the same directory as the programs and its name should be included in all programs.

```
// "headerFiles.h"

#include <stdio.h>

#include <stdlib.h>

#include <sys/types.h>

#include <sys/socket.h>

#include <netinet/in.h>

#include <netdb.h>

#include <errno.h>

#include <signal.h>

#include <unistd.h>

#include <string.h>

#include <arpa/inet.h>

#include <sys/wait.h>
```

Iterative Programming Using UDP

- UDP provides a connectionless server, in which a client sends a request and the server sends back a response.

Programming Examples

- The way to write client and server programs to simulate the standard echo application using TCP.
- The client program sends a short string of characters to the server; the server echoes back the same string to the client.
- However, before we do so, we need to provide the flow diagram for the client and server data-transfer boxes, which is shown in Figure 14.8.

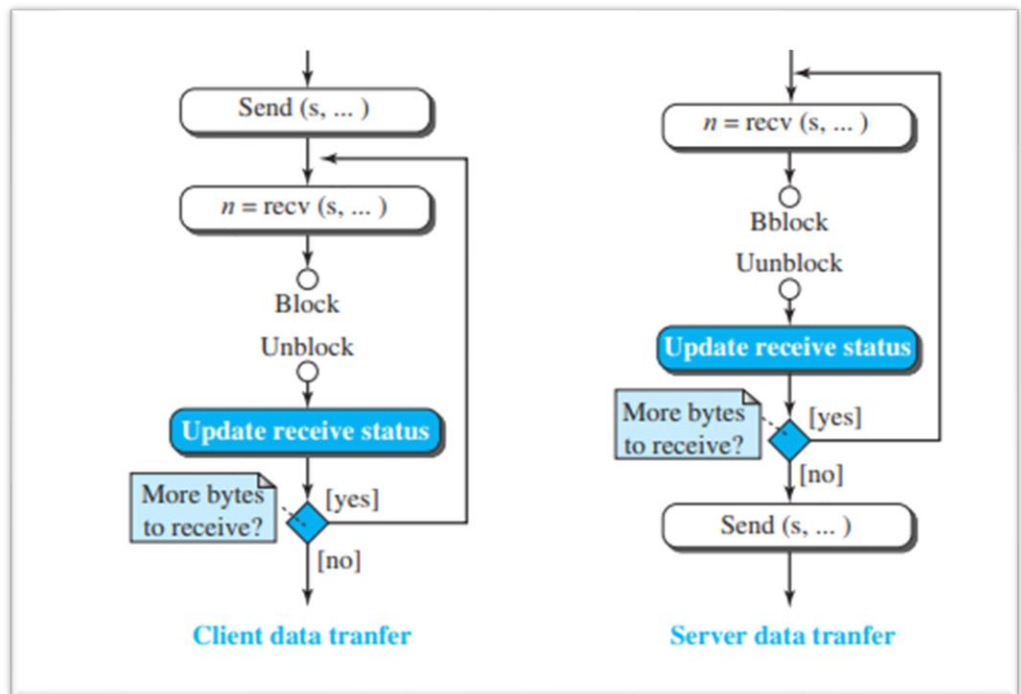


Figure 14.8 Flow diagram for the client and server data-transfer boxes

- For this special case, since the size of the string to be sent is small (less than a few words), we can do it in one call to the send procedure in the client.
- However, it is not guaranteed that the TCP will send the whole message in one segment.
- Therefore, we need to use a set of receive calls in the server site (in a loop), to receive all the segments and collect them in the buffer to be sent back in one shot.
- When the server is sending back the echo message, it may also use several segments to do so, which means the rcv procedure in the client needs to be called as many times as needed.
- Another issue to be solved is setting the buffers that hold data at each site.
- We need to control how many bytes of data we have received and where the next chunk of data is stored.
- The program sets some variables to control the situation, as shown in Figure 14.8.
- In each iteration, the pointer (ptr) moves ahead to point to the next bytes to receive, the length of received bytes (len) is increased, and the maximum number of bytes to be received (maxLen) is decreased.

ITERATIVE PROGRAMMING IN JAVA

- The entities defined in the C language are redefined in an object-programming language.
- We have chosen the Java language because many aspects of programming can be easily shown using the powerful classes available in Java.
- We touch the main issues in programming using the traditional socket interface API, but they can be extended to other areas of network programming without difficulty.
- We assume that the reader is familiar with the basics of Java programming.

Addresses and Ports

- Network programming in any language definitely needs to deal with IP addresses and port numbers.
- We briefly introduce how addresses and ports are represented in Java.
- We recommend that the reader compare the representations of these two entities in C and Java.

IP Addresses

- There are two types of IP addresses used in the Internet: IPv4 addresses (32 bits) and IPv6 addresses (128 bits).
- In Java, an IP address is defined as an object, the instance of InetAddress class.
- The class was originally defined as a final class, which means it was not inheritable.
- Later Java changed the class and defined two subclasses inherited from this class: Inet4Address and Inet6Address.
- However, most of the time we use only the InetAddress class to create both IPv4 and IPv6 addresses.
- There is no public constructor in the InetAddress class, but we can use any of the static methods in this class to return an instance of InetAddress.
- The class has also some instance methods that can be used to change the format of the address object or get some information about the object.

Port Numbers

- A port number in the TCP/IP protocol suite is an unsigned 16-bit integer.
- However, since Java does not define an unsigned numeric data type, a port number in Java is defined as an integer data type (32-bit int) in which the left 16 bits are set to zeros.
- This prevents a large port number from being interpreted as a negative number.

InetAddress

- A socket address is a combination of an IP address and a port number.
- In Java, there is an abstract class named `Socket`, but the class used in Java network programming is the `InetAddress` class that inherits from the `Socket` class.

Iterative Programming Using UDP

- To be consistent with the C socket programming section above, first to discuss Java programming using the service of UDP, a connectionless service.

Two Classes Designed for UDP

- There are two classes designed to be used with UDP: `DatagramSocket` class and `DatagramPacket` class.

DatagramSocket Class

- The `DatagramSocket` class is used to create sockets in the client and server.
- It also provides methods to send a datagram, to receive a datagram, and to close the socket

DatagramPacket Class

- The `DatagramPacket` class is used to create datagram packets.

14.2 LET US SUM UP

- Applications in the Internet are designed using either a client-server paradigm or a peer-to-peer paradigm.
- In a client-server paradigm, an application program, called a server, provides services, and another application program, called a client, receives services.
- A server program is an infinite program; a client program is finite.
- In a peer-to-peer paradigm, a peer can be both a client and a server.

- A server in a client-server paradigm can be designed either as an iterative server or as a concurrent server.
- An iterative server handles the clients one by one.
- A concurrent server can simultaneously serve as many clients as the computer resources permit.
- A client-server pair that uses the services of a connectionless transport layer, such as UDP, should be designed as connectionless programs.
- A client-server pair that uses the services of a connection-oriented transport layer, such as TCP, should be designed as connection-oriented programs.

14.3 LIST OF REFERENCES

<https://ieeexplore.ieee.org/document/7740559>

<https://ieeexplore.ieee.org/document/8394143>

14.4 BIBLIOGRAPHY

1. Data Communications and Networking, Behrouz A. Forouzan, Fifth Edition, TMH, 2018.
2. Computer Network, Bhushan Trivedi, Oxford University Press, 2016

14.5 UNIT END EXERCISES

1. How is an IP address represented in Java?
2. What is the difference between the DatagramSocket class and the Socket class in Java?
3. Write a note on Iterative programming.
4. Write note on TCP protocol.
5. Application Programming Interface.
6. Explain Client server programming.
7. Application-Layer Paradigms.

STANDARD CLIENT-SERVER PROTOCOLS

Unit Structure

- 15.0 Objectives
- 15.1 Introduction
 - 15.1.1 WWW
 - 15.1.2 HTTP
 - 15.1.3 FTP
 - 15.1.4 Electronic Mail
 - 15.1.5 TELNET
 - 15.1.6 Secure Cell
 - 15.1.7 DNS
 - 15.1.8 SNMP
- 15.2 Let us Sum Up
- 15.3 List of References
- 15.4 Bibliography
- 15.5 Unit End Exercises

15.0 OBJECTIVES

- After going through this unit, you will be able to:
- To understand the World Wide Web.
- Understand Hyper Text Transfer Protocol
- Understand client-server application program.
- Understand two protocols: SMPT and POP.
- Understand TELNET.
- Understand Secure Shell
- Understand Domain Name System.

15.1 INTRODUCTION

During the lifetime of the Internet, several client-server application programs have been developed. We do not have to redefine them, but we need to understand what they do. For each application, we also need to know the options available to us. The study of these applications and the ways they provide different services can help us to create customized applications in the future.

15.1.1 WWW (World Wide Web)

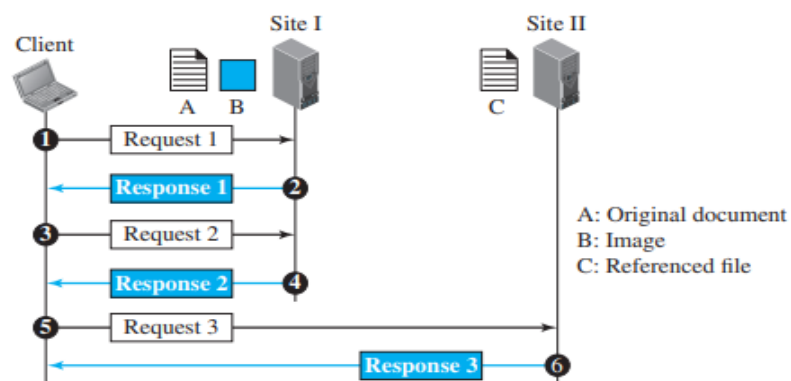
- The idea of the Web was first proposed by Tim Berners-Lee in 1989 at CERN, the European Organization for Nuclear Research, to allow several researchers at different locations throughout Europe to access each other's' researches.
- The commercial Web started in the early 1990s.
- The Web today is a repository of information in which the documents, called web pages, are distributed all over the world and related documents are linked together.
- The popularity and growth of the Web can be related to two terms in the above statement: distributed and linked.
- Distribution allows the growth of the Web.
- Each web server in the world can add a new web page to the repository and announce it to all Internet users without overloading a few servers.
- Linking allows one web page to refer to another web page stored in another server somewhere else in the world.
- The linking of web pages was achieved using a concept called hypertext, which was introduced many years before the advent of the Internet.
- The idea was to use a machine that automatically retrieved another document stored in the system when a link to it appeared in the document.
- The Web implemented this idea electronically to allow the linked document to be retrieved when the link was clicked by the user.
- Today, the term hypertext, coined to mean linked text documents, has been changed to hypermedia, to show that a web page can be a text document, an image, an audio file, or a video file.
- The purpose of the Web has gone beyond the simple retrieving of linked documents.
- Today, the Web is used to provide electronic shopping and gaming.
- One can use the Web to listen to radio programs or view television programs whenever one desires without being forced to listen to or view these programs when they are broadcast.

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 web pages.
- Each web page, however, can contain some links to other web pages in the same or other sites.
- In other words, a web page can be simple or composite.
- A simple web page has no links to other web pages; a composite web page has one or more links to other web pages.
- Each web page is a file with a name and address.

Example:

- Assume we need to retrieve a scientific document that contains one reference to another text file and one reference to a large image.
- Figure 15.1 shows the situation.
- The main document and the image are stored in two separate files (file A and file B) in the same site; the referenced text file (file C) is stored in another site.
- Since we are dealing with three different files, we need three transactions if we want to see the whole document.
- The first transaction (request/response) retrieves a copy of the main document (file A), which has references (pointers) to the second and third files.

**Figure 15.1 Example**

When a copy of the main document is retrieved and browsed, the user can click on the reference to the image to invoke the second transaction and retrieve a copy of the image (file B).

If the user needs to see the contents of the referenced text file, she can click on its reference (pointer) invoking the third transaction and retrieving a copy of file C.

Note that although files A and B both are stored in site I, they are independent files with different names and addresses.

Two transactions are needed to retrieve them.

A very important point we need to remember is that file A, file B, and file C in Figure 15.1 are independent web pages, each with independent names and addresses.

Although references to file B or C are included in file A, it does not mean that each of these files cannot be retrieved independently.

A second user can retrieve file B with one transaction.

A third user can retrieve file C with one transaction.

Web Client (Browser)

A variety of vendors offer commercial browsers that interpret and display a web page, and all of them use nearly the same architecture.

Each browser usually consists of three parts: a controller, client protocols, and interpreters.

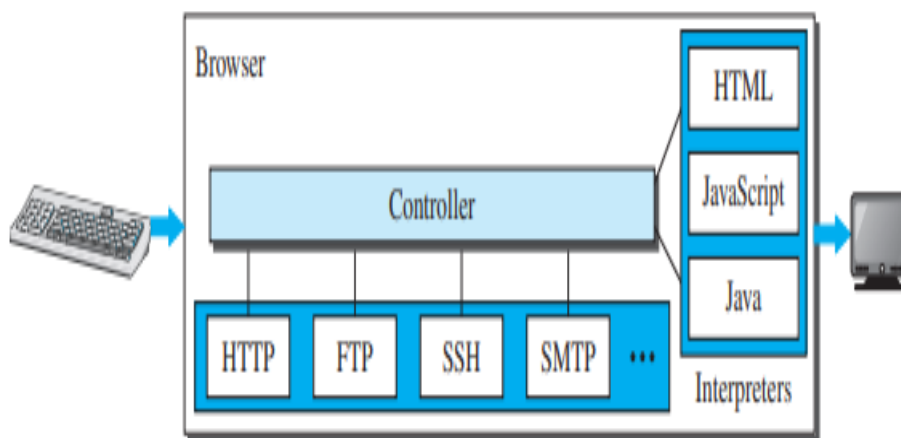


Figure 15.2 Browser

- 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 later, such as HTTP or FTP.
- The interpreter can be HTML, Java, or JavaScript, depending on the type of document.
- Some commercial browsers include Internet Explorer, Netscape Navigator, and Firefox.

Web Server

- **The web page is stored at the server.**
- **Each time a 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 a 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.**
- **Some popular web servers include Apache and Microsoft Internet Information Server.**

Uniform Resource Locator (URL)

- A web page, as a file, needs to have a unique identifier to distinguish it from other web pages.
- To define a web page, we need three identifiers: host, port, and path.
- However, before defining the web page, we need to tell the browser what client server application we want to use, which is called the protocol.
- This means we need four identifiers to define the web page.
- The first is the type of vehicle to be used to fetch the web page; the last three make up the combination that defines the destination object (web page).

Protocol:

- **The first identifier is the abbreviation for the client-server program that we need in order to access the web page.**
- **Although most of the time the protocol is HTTP (HyperText Transfer Protocol) can also use other protocols such as FTP (File Transfer Protocol).**

Host:

- The host identifier can be the IP address of the server or the unique name given to the server.
- IP addresses can be defined in dotted decimal notation, (such as 64.23.56.17); the name is normally the domain name that uniquely defines the host, such as forouzan.com, Domain Name System (DNS).

Port:

- The port, a 16-bit integer, is normally predefined for the client-server application.
- For example, if the HTTP protocol is used for accessing the web page, the well-known port number is 80.
- However, if a different port is used, the number can be explicitly given.

Path:

- The path identifies the location and the name of the file in the underlying operating system.
- The format of this identifier normally depends on the operating system.
- In UNIX, a path is a set of directory names followed by the file name, all separated by a slash.
- For example, /top/next/last/myfile is a path that uniquely defines a file named myfile, stored in the directory last, which itself is part of the directory next, which itself is under the directory top.
- In other words, the path lists the directories from the top to the bottom, followed by the file name.

To combine these four pieces together, the uniform resource locator (URL) has been designed; it uses three different separators between the four pieces as shown below:

- protocol://host/path Used most of the time
- protocol://host:port/path Used when port number is needed

15.1.2 Hyper Text Transfer Protocol (HTTP)

- **The HyperText Transfer Protocol (HTTP) is used to define how the client-server programs can be written to retrieve web pages from the Web.**

- **An HTTP client sends a request; an HTTP server returns a response.**
- **The server uses the port number 80; the client uses a temporary port number.**
- **HTTP uses the services of TCP, which, is a connection-oriented and reliable protocol.**
- **This means that, before any transaction between the client and the server can take place, a connection needs to be established between them.**
- **After the transaction, the connection should be terminated.**
- **The client and server, however, do not need to worry about errors in messages exchanged or loss of any message, because the TCP is reliable.**

Nonpersistent versus Persistent Connections

- The hypertext concept embedded in web page documents may require several requests and responses.
- If the web pages, objects to be retrieved, are located on different servers, we do not have any other choice than to create a new TCP connection for retrieving each object.
- However, if some of the objects are located on the same server, we have two choices: to retrieve each object using a new TCP connection or to make a TCP connection and retrieve them all.
- The first method is referred to as a nonpersistent connection, the second as a persistent connection.
- HTTP, prior to version 1.1, specified nonpersistent connections, while persistent connections are the default in version 1.1, but it can be changed by the user.

Nonpersistent Connections

- In a nonpersistent connection, one TCP connection is made for each request/response.
- The following lists the steps in this strategy:
 1. The client opens a TCP connection and sends a request.
 2. The server sends the response and closes the connection.
 3. The client reads the data until it encounters an end-of-file marker; it then closes the connection.
- **In this strategy, if a file contains links to N different pictures in different files (all located on the same server), the connection must be opened and closed N + 1 times.**

- The nonpersistent strategy imposes high overhead on the server because the server needs $N + 1$ different buffers each time a connection is opened.

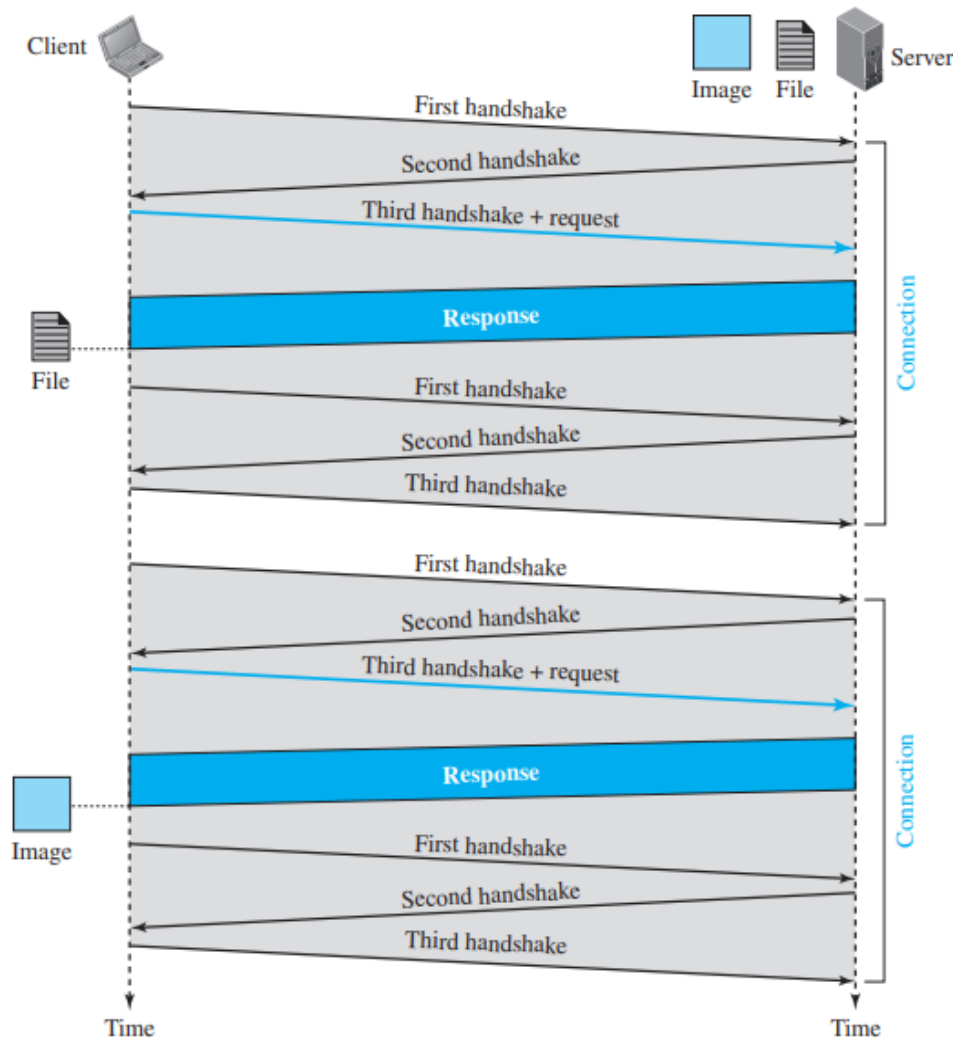


Figure 15.3 Non persistent connection

Persistent Connections

- HTTP version 1.1 specifies a persistent connection by default.
- In a persistent connection, the server leaves the connection open for more requests after sending a response.
- The server can close the connection at the request of a client or if a time-out has been reached.
- The sender usually sends the length of the data with each response.
- However, there are some occasions when the sender does not know the length of the data.
- This is the case when a document is created dynamically or actively.

- In these cases, the server informs the client that the length is not known and closes the connection after sending the data so the client knows that the end of the data has been reached.
- Time and resources are saved using persistent connections.
- Only one set of buffers and variables needs to be set for the connection at each site.
- The round-trip time for connection establishment and connection termination is saved.

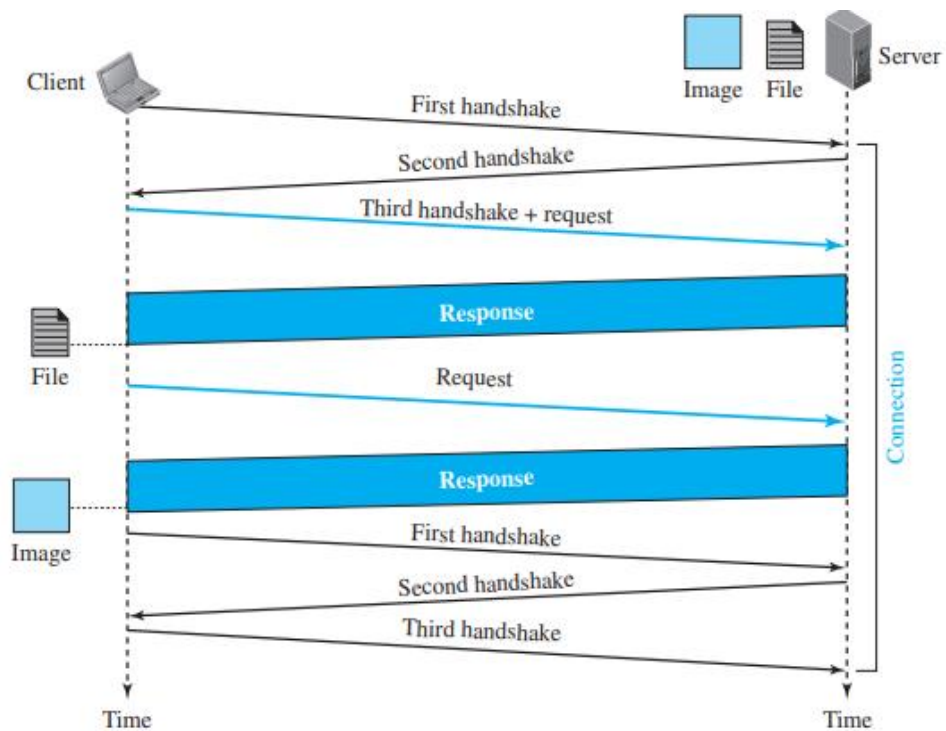


Figure 15.4 Persistent Connection

Web Caching: Proxy Servers

- HTTP supports proxy servers.
- A proxy server is a computer that keeps copies of responses to recent requests.
- The HTTP client sends a request to the proxy server.
- The proxy server checks its cache.
- If the response is not stored in the cache, the proxy server sends the request to the corresponding server.
- Incoming responses are sent to the proxy server and stored for future requests from other clients.

- The proxy server reduces the load on the original server, decreases traffic, and improves latency.
- However, to use the proxy server, the client must be configured to access the proxy instead of the target server.
- Note that the proxy server acts as both server and client.
- When it receives a request from a client for which it has a response, it acts as a server and sends the response to the client.
- When it receives a request from a client for which it does not have a response, it first acts as a client and sends a request to the target server.
- When the response has been received, it acts again as a server and sends the response to the client.

Proxy Server Location

- The proxy servers are normally located at the client site. This means that we can have a hierarchy of proxy servers, as shown below:
 1. A client computer can also be used as a proxy server, in a small capacity, that stores responses to requests often invoked by the client.
 2. In a company, a proxy server may be installed on the computer LAN to reduce the load going out of and coming into the LAN.
 3. An ISP with many customers can install a proxy server to reduce the load going out of and coming into the ISP network.

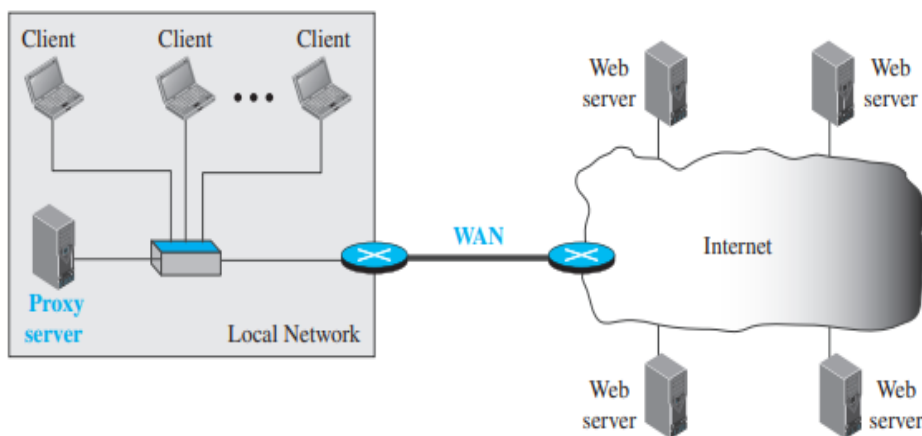


Figure 15.5 Proxy Server

- The proxy server is installed in the local network.
- When an HTTP request is created by any of the clients (browsers), the request is first directed to the proxy server.
- If the proxy server already has the corresponding web page, it sends the response to the client.

- Otherwise, the proxy server acts as a client and sends the request to the web server in the Internet.
- When the response is returned, the proxy server makes a copy and stores it in its cache before sending it to the requesting client.

HTTP Security

- HTTP per se does not provide security.
- However, HTTP can be run over the Secure Socket Layer (SSL).
- In this case, HTTP is referred to as HTTPS.
- HTTPS provides confidentiality, client and server authentication, and data integrity.

15.1.3 FTP (File Transfer Protocol)

- **File Transfer Protocol (FTP) is the standard protocol provided by TCP/IP for copying a file from one host to another.**
- **Although transferring files from one system to another seems simple and straightforward, some problems must be dealt with first.**
- **For example, two systems may use different file name conventions.**
- **Two systems may have different ways to represent data.**
- **Two systems may have different directory structures.**
- **All of these problems have been solved by FTP in a very simple and elegant approach.**
- **Although we can transfer files using HTTP, FTP is a better choice to transfer large files or to transfer files using different formats.**

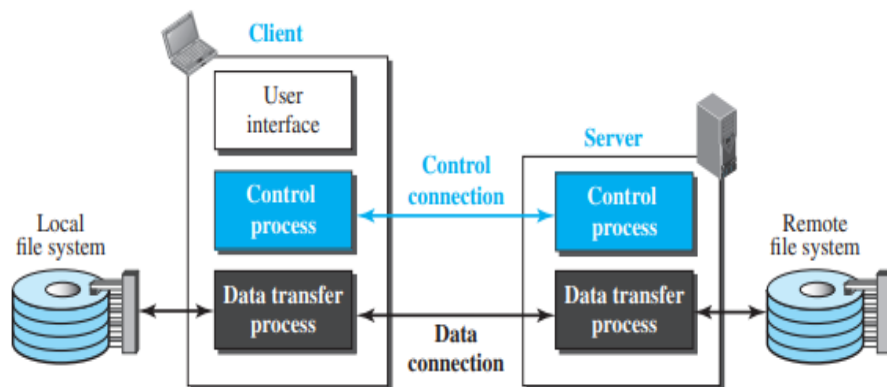


Figure 15.6 FTP

Figure 15.6 shows the basic model of FTP.

- The client has three components: the user interface, the client control process, and the client data transfer process.
- The server has two components: the server control process and the server data transfer process.
- The control connection is made between the control processes.
- The data connection is made between the data transfer processes.
- Separation of commands and data transfer makes FTP more efficient.
- The control connection uses very simple rules of communication.
- We need to transfer only a line of command or a line of response at a time.
- The data connection, on the other hand, needs more complex rules due to the variety of data types transferred.

Two Connections

- The two connections in FTP have different lifetimes.
- The control connection remains connected during the entire interactive FTP session.
- The data connection is opened and then closed for each file transfer activity.
- It opens each time commands that involve transferring files are used, and it closes when the file is transferred.
- In other words, when a user starts an FTP session, the control connection opens.
- While the control connection is open, the data connection can be opened and closed multiple times if several files are transferred.
- FTP uses two well-known TCP ports: port 21 is used for the control connection, and port 20 is used for the data connection.

Control Connection

- For control communication, FTP uses the same approach as TELNET.
- It uses the NVT ASCII character set as used by TELNET.
- Communication is achieved through commands and responses.
- This simple method is adequate for the control connection because we send one command (or response) at a time.
- Each line is terminated with a two-character (carriage return and line feed) end-of-line token.

- During this control connection, commands are sent from the client to the server and responses are sent from the server to the client.
- Commands, which are sent from the FTP client control process, are in the form of ASCII uppercase, which may or may not be followed by an argument.

Some of the most common commands are shown in Table 15.1

<i>Command</i>	<i>Argument(s)</i>	<i>Description</i>
ABOR		Abort the previous command
CDUP		Change to parent directory
CWD	Directory name	Change to another directory
DELE	File name	Delete a file
LIST	Directory name	List subdirectories or files
MKD	Directory name	Create a new directory
PASS	User password	Password
PASV		Server chooses a port
PORT	Port identifier	Client chooses a port
PWD		Display name of current directory
QUIT		Log out of the system
RETR	File name(s)	Retrieve files; files are transferred from server to client
RMD	Directory name	Delete a directory
RNFR	File name (old)	Identify a file to be renamed
RNTO	File name (new)	Rename the file
STOR	File name(s)	Store files; file(s) are transferred from client to server
STRU	F, R, or P	Define data organization (F : file, R : record, or P : page)
TYPE	A, E, I	Default file type (A : ASCII, E : EBCDIC, I : image)
USER	User ID	User information
MODE	S, B, or C	Define transmission mode (S : stream, B : block, or C : compressed)

Table 15.1 Some FTP Commands

- Every FTP command generates at least one response.
- A response has two parts: a three-digit number followed by text.
- The numeric part defines the code; the text part defines needed parameters or further explanations.
- The first digit defines the status of the command.
- The second digit defines the area in which the status applies.
- The third digit provides additional information.

Table 15.2 shows some common responses.

<i>Code</i>	<i>Description</i>	<i>Code</i>	<i>Description</i>
125	Data connection open	250	Request file action OK
150	File status OK	331	User name OK; password is needed
200	Command OK	425	Cannot open data connection
220	Service ready	450	File action not taken; file not available
221	Service closing	452	Action aborted; insufficient storage
225	Data connection open	500	Syntax error; unrecognized command
226	Closing data connection	501	Syntax error in parameters or arguments
230	User login OK	530	User not logged in

Table 15.2 Some responses in FTP

Data Connection

- The data connection uses the well-known port 20 at the server site.
- However, the creation of a data connection is different from the control connection.
- The following shows the steps:
 1. The client, not the server, issues a passive open using an ephemeral port. This must be done by the client because it is the client that issues the commands for transferring files.
 2. Using the PORT command the client sends this port number to the server.
 3. The server receives the port number and issues an active open using the well-known port 20 and the received ephemeral port number.

Communication over Data Connection

- The purpose and implementation of the data connection are different from those of the control connection.
- We want to transfer files through the data connection.
- The client must define the type of file to be transferred, the structure of the data, and the transmission mode.
- Before sending the file through the data connection, we prepare for transmission through the control connection.
- The heterogeneity problem is resolved by defining three attributes of communication: file type, data structure, and transmission mode.

File Type

FTP can transfer one of the following file types across the data connection: ASCII file, EBCDIC file, or image file.

Data Structure

- FTP can transfer a file across the data connection using one of the following interpretations of the structure of the data: file structure, record structure, or page structure.
- The file structure format (used by default) has no structure.
- It is a continuous stream of bytes.
- In the record structure, the file is divided into records.
- This can be used only with text files.
- In the page structure, the file is divided into pages, with each page having a page number and a page header.
- The pages can be stored and accessed randomly or sequentially.

Transmission Mode

- FTP can transfer a file across the data connection using one of the following three transmission modes: stream mode, block mode, or compressed mode.
- The stream mode is the default mode; data are delivered from FTP to TCP as a continuous stream of bytes.
- In the block mode, data can be delivered from FTP to TCP in blocks.
- In this case, each block is preceded by a 3-byte header.
- The first byte is called the block descriptor; the next two bytes define the size of the block in bytes.

File Transfer

- File transfer occurs over the data connection under the control of the commands sent over the control connection.
- However, we should remember that file transfer in FTP means one of three things: retrieving a file (server to client), storing a file (client to server), and directory listing (server to client).

Security for FTP

- The FTP protocol was designed when security was not a big issue.
- Although FTP requires a password, the password is sent in plaintext (unencrypted), which means it can be intercepted and used by an attacker.
- The data transfer connection also transfers data in plaintext, which is insecure.
- To be secure, one can add a Secure Socket Layer between the FTP application layer and the TCP layer.
- In this case FTP is called SSL-FTP.

15.1.4 ELECTRONIC MAIL

- Electronic mail (or e-mail) allows users to exchange messages.
- The nature of this application, however, is different from other applications discussed so far.
- In an application such as HTTP or FTP, the server program is running all the time, waiting for a request from a client.
- When the request arrives, the server provides the service.
- There is a request and there is a response.
- In the case of electronic mail, the situation is different.
- First, e-mail is considered a one-way transaction.
- When Alice sends an email to Bob, she may expect a response, but this is not a mandate. Bob may or may not respond.
- If he does respond, it is another one-way transaction.
- Second, it is neither feasible nor logical for Bob to run a server program and wait until someone sends an e-mail to him.
- Bob may turn off his computer when he is not using it.
- This means that the idea of client/server programming should be implemented in another way: using some intermediate computers (servers).
- The users run only client programs when they want and the intermediate servers apply the client/server paradigm.

Architecture

- To explain the architecture of e-mail, we give a common scenario, as shown in Figure 15.7.
- Another possibility is the case in which Alice or Bob is directly connected to the corresponding mail server, in which LAN or WAN connection is not required, but this variation in the scenario does not affect our discussion.

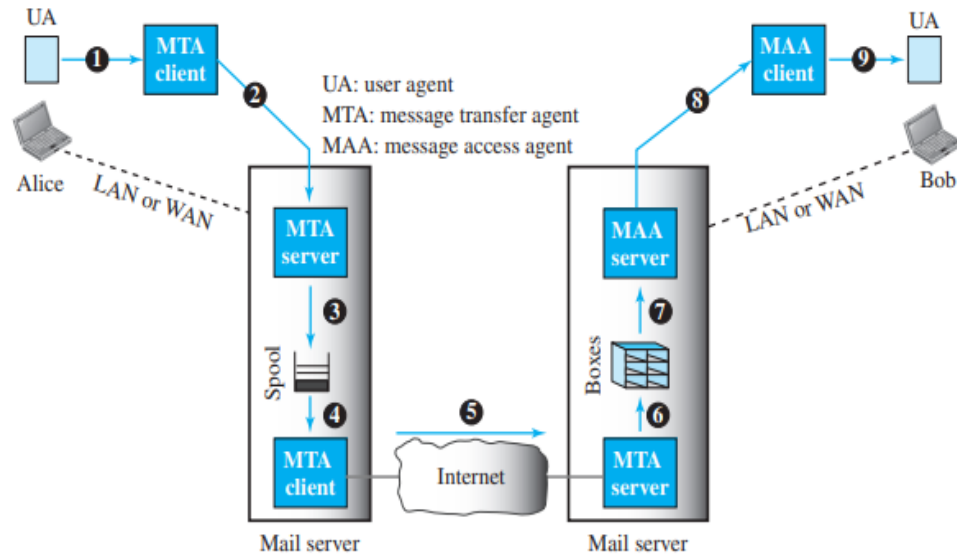


Figure 15.7 common scenario of architecture

- In the common scenario, the sender and the receiver of the e-mail, Alice and Bob respectively, are connected via a LAN or a WAN to two mail servers.
- The administrator has created one mailbox for each user where the received messages are stored.
- A mail box is part of a server hard drive, a special file with permission restrictions.
- Only the owner of the mailbox has access to it.
- The administrator has also created a queue (spool) to store messages waiting to be sent.
- A simple e-mail from Alice to Bob takes nine different steps, as shown in the figure 15.7.
- Alice and Bob use three different agents: a user agent (UA), a message transfer agent (MTA), and a message access agent (MAA).
- When Alice needs to send a message to Bob, she runs a UA program to prepare the message and send it to her mail server.
- The mail server at her site uses a queue (spool) to store messages waiting to be sent.
- The message, however, needs to be sent through the Internet from Alice's site to Bob's site using an MTA.
- Here two message transfer agents are needed: one client and one server.
- Like most client-server programs on the Internet, the server needs to run all the time because it does not know when a client will ask for a connection.
- The client, on the other hand, can be triggered by the system when there is a message in the queue to be sent.

- The user agent at the Bob site allows Bob to read the received message.
- Bob later uses an MAA client to retrieve the message from an MAA server running on the second server.

There are two important points we need to emphasize here.

- First, Bob cannot bypass the mail server and use the MTA server directly.
- To use the MTA server directly, Bob would need to run the MTA server all the time because he does not know when a message will arrive.
- This implies that Bob must keep his computer on all the time if he is connected to his system through a LAN.
- If he is connected through a WAN, he must keep the connection up all the time.
- Neither of these situations is feasible today.
- Second, note that Bob needs another pair of client-server programs: message access programs.
- This is because an MTA client-server program is a push program: the client pushes the message to the server.
- Bob needs a pull program.
- The client needs to pull the message from the server.

User Agent

- The first component of an electronic mail system is the user agent (UA).
- It provides service to the user to make the process of sending and receiving a message easier.
- A user agent is a software package (program) that composes, reads, replies to, and forwards messages.
- It also handles local mailboxes on the user computers.
- There are two types of user agents: command-driven and GUI-based.
- Command driven user agents belong to the early days of electronic mail.
- They are still present as the underlying user agents.
- A command-driven user agent normally accepts a onecharacter command from the keyboard to perform its task.
- For example, a user can type the character r, at the command prompt, to reply to the sender of the message, or type the character R to reply to the sender and all recipients.

- Some examples of command driven user agents are mail, pine, and elm.
- Modern user agents are GUI-based.
- They contain graphical user interface (GUI) components that allow the user to interact with the software by using both the keyboard and the mouse.
- They have graphical components such as icons, menu bars, and windows that make the services easy to access.
- Some examples of GUI-based user agents are Eudora and Outlook.

Sending Mail

- To send mail, the user, through the UA, creates mail that looks very similar to postal mail.
- It has an envelope and a message.
- The envelope usually contains the sender address, the receiver address, and other information.
- The message contains the header and the body.
- The header of the message defines the sender, the receiver, the subject of the message, and some other information.
- The body of the message contains the actual information to be read by the recipient.

Receiving Mail

- The user agent is triggered by the user (or a timer).
- If a user has mail, the UA informs the user with a notice.
- If the user is ready to read the mail, a list is displayed in which each line contains a summary of the information about a particular message in the mail box.
- The summary usually includes the sender mail address, the subject, and the time the mail was sent or received.
- The user can select any of the messages and display its contents on the screen.

Addresses

- To deliver mail, a mail handling system must use an addressing system with unique addresses.
- In the Internet, the address consists of two parts: a local part and a domain name, separated by an @ sign

Message Transfer Agent: SMTP

- Based on the common scenario (Figure 15.7), we can say that the e-mail is one of those applications that needs three uses of client-server paradigms to accomplish its task.
- It is important that we distinguish these three when we are dealing with e-mail.
- Figure 15.8 shows these three client-server applications.
- We refer to the first and the second as Message Transfer Agents (MTAs), the third as Message Access Agent (MAA).

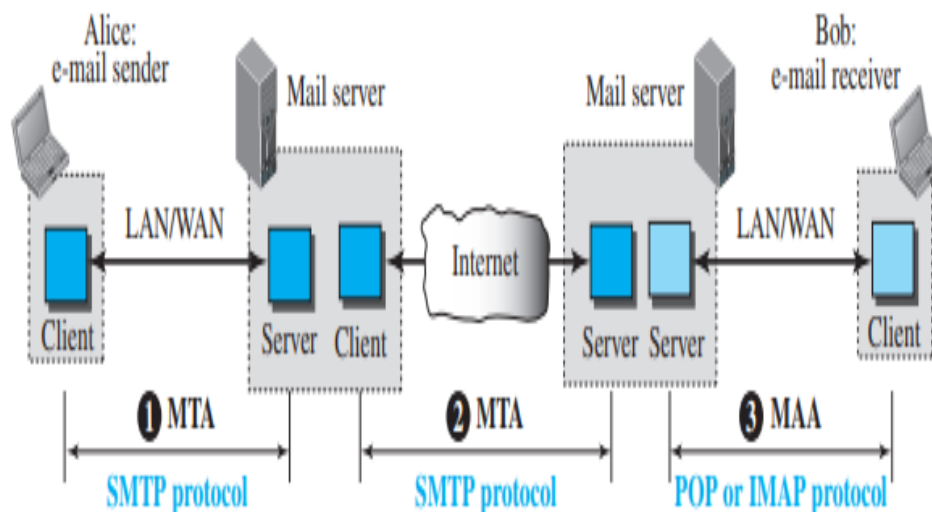


Figure 15.8 three client-server applications

- The formal protocol that defines the MTA client and server in the Internet is called Simple Mail Transfer Protocol (SMTP).
- SMTP is used two times, between the sender and the sender's mail server and between the two mail servers.
- As we will see shortly, another protocol is needed between the mail server and the receiver.
- SMTP simply defines how commands and responses must be sent back and forth.

Commands and Responses

- SMTP uses commands and responses to transfer messages between an MTA client and an MTA server.
- The command is from an MTA client to an MTA server; the response is from an MTA server to the MTA client.
- Each command or reply is terminated by a two character (carriage return and line feed) end-of-line token.

Commands

- Commands are sent from the client to the server.
- The format of a command is shown below:

Keyword: argument(s)

- It consists of a keyword followed by zero or more arguments.
- SMTP defines 14 commands, listed in Table 15.3

<i>Keyword</i>	<i>Argument(s)</i>	<i>Description</i>
HELO	Sender's host name	Identifies itself
MAIL FROM	Sender of the message	Identifies the sender of the message
RCPT TO	Intended recipient	Identifies the recipient of the message
DATA	Body of the mail	Sends the actual message
QUIT		Terminates the message
RSET		Aborts the current mail transaction
VRFY	Name of recipient	Verifies the address of the recipient
NOOP		Checks the status of the recipient
TURN		Switches the sender and the recipient
EXPN	Mailing list	Asks the recipient to expand the mailing list
HELP	Command name	Asks the recipient to send information about the command sent as the argument
SEND FROM	Intended recipient	Specifies that the mail be delivered only to the terminal of the recipient, and not to the mailbox
SMOL FROM	Intended recipient	Specifies that the mail be delivered to the terminal <i>or</i> the mailbox of the recipient
SMAL FROM	Intended recipient	Specifies that the mail be delivered to the terminal <i>and</i> the mailbox of the recipient

Table 15.3 SMTP commands

Responses:

- Responses are sent from the server to the client.
- A response is a three digit code that may be followed by additional textual information.
- Table 26.4 shows the most common response types.

<i>Code</i>	<i>Description</i>
Positive Completion Reply	
211	System status or help reply
214	Help message
220	Service ready
221	Service closing transmission channel
250	Request command completed
251	User not local; the message will be forwarded
Positive Intermediate Reply	
354	Start mail input
Transient Negative Completion Reply	
421	Service not available
450	Mailbox not available
451	Command aborted: local error
452	Command aborted; insufficient storage
Permanent Negative Completion Reply	
500	Syntax error; unrecognized command

<i>Code</i>	<i>Description</i>
501	Syntax error in parameters or arguments
502	Command not implemented
503	Bad sequence of commands
504	Command temporarily not implemented
550	Command is not executed; mailbox unavailable
551	User not local
552	Requested action aborted; exceeded storage location
553	Requested action not taken; mailbox name not allowed
554	Transaction failed

Table 15.4 Responses**15.1.5 TELNET**

- A server program can provide a specific service to its corresponding client program.
- For example, the FTP server is designed to let the FTP client store or retrieve files on the server site.

- However, it is impossible to have a client/server pair for each type of service we need; the number of servers soon becomes intractable.
- The idea is not scalable.
- Another solution is to have a specific client/server program for a set of common scenarios, but to have some generic client/server programs that allow a user on the client site to log into the computer at the server site and use the services available there.
- For example, if a student needs to use the Java compiler program at her university lab, there is no need for a Java compiler client and a Java compiler server.
- The student can use a client logging program to log into the university server and use the compiler program at the university.
- We refer to these generic client/server pairs as remote logging applications.
- One of the original remote logging protocols is TELNET, which is an abbreviation for TErminaL NETwork.
- Although TELNET requires a logging name and password, it is vulnerable to hacking because it sends all data including the password in plaintext (not encrypted).
- A hacker can eavesdrop and obtain the logging name and password.
- Because of this security issue, the use of TELNET has diminished in favor of another protocol, Secure Shell (SSH), which we describe in the next section.
- Although TELNET is almost replaced by SSH, we briefly discuss TELNET here for two reasons:
 1. The simple plaintext architecture of TELNET allows us to explain the issues and challenges related to the concept of remote logging, which is also used in SSH when it serves as a remote logging protocol.
 2. Network administrators often use TELNET for diagnostic and debugging purposes.

Local versus Remote Logging

The concept of local and remote logging as shown in Figure 15.9

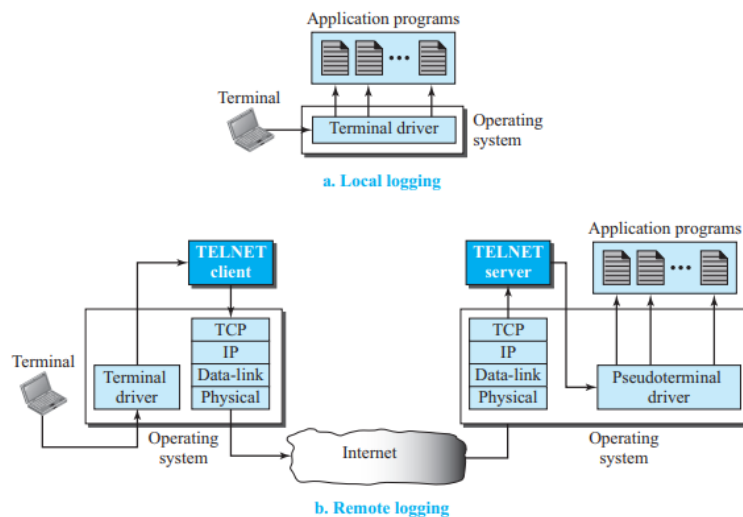


Figure 15.9 local and remote logging

When a user logs into a local system, it is called local logging.

As a user types at a terminal or at a workstation running a terminal emulator, the keystrokes are accepted by the terminal driver.

The terminal driver passes the characters to the operating system.

The operating system, in turn, interprets the combination of characters and invokes the desired application program or utility.

However, when a user wants to access an application program or utility located on a remote machine, she performs remote logging.

Here the TELNET client and server programs come into use.

The user sends the keystrokes to the terminal driver where the local operating system accepts the characters but does not interpret them.

The characters are sent to the TELNET client, which transforms the characters into a universal character set called Network Virtual Terminal (NVT) characters (discussed below) and delivers them to the local TCP/IP stack.

- **Network Virtual Terminal (NVT)**
 - The mechanism to access a remote computer is complex.
 - This is because every computer and its operating system accept a special combination of characters as tokens.

- For example, the end-of-file token in a computer running the DOS operating system is Ctrl+z, while the UNIX operating system recognizes Ctrl+d.
- We are dealing with heterogeneous systems.
- If we want to access any remote computer in the world, we must first know what type of computer we will be connected to, and we must also install the specific terminal emulator used by that computer.
- TELNET solves this problem by defining a universal interface called the Network Virtual Terminal (NVT) character set.
- Via this interface, the client TELNET translates characters (data or commands) that come from the local terminal into NVT form and delivers them to the network.
- The server TELNET, on the other hand, translates data and commands from NVT form into the form acceptable by the remote computer.
- Figure 15.10 shows the concept of NVT.
- NVT uses two sets of characters, one for data and one for control. Both are 8-bit bytes as shown in Figure 15.10.
- For data, NVT normally uses what is called NVT ASCII.
- This is an 8-bit character set in which the seven lowest order bits are the same as US ASCII and the highest order bit is 0.
- To send control characters between computers (from client to server or vice versa), NVT uses an 8-bit character set in which the highest order bit is set to 1.

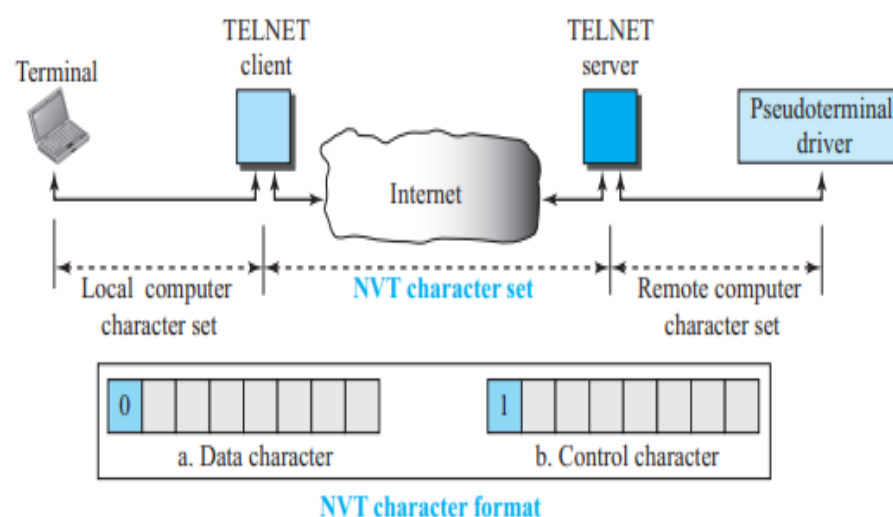


Figure 15.10 Concept of NVT

- **Options**
- TELNET lets the client and server negotiate options before or during the use of the service.
- Options are extra features available to a user with a more sophisticated terminal.
- Users with simpler terminals can use default features.
- **User Interface**
- The operating system (UNIX, for example) defines an interface with user-friendly commands.
- An example of such a set of commands can be found in Table 15.5.

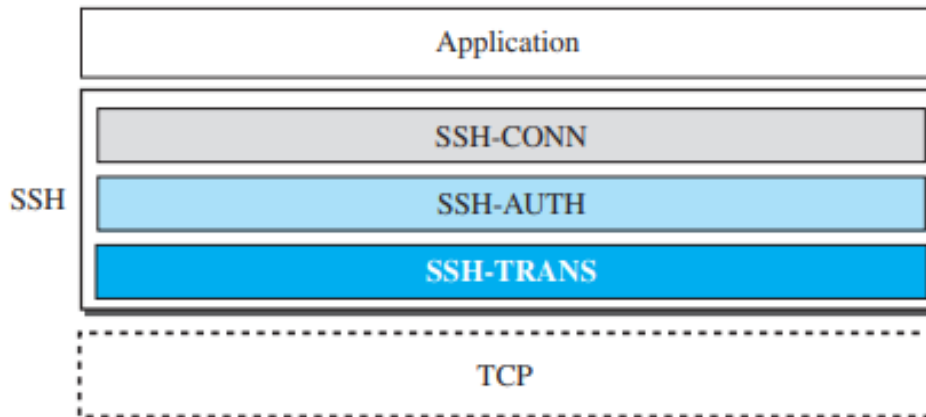
<i>Command</i>	<i>Meaning</i>	<i>Command</i>	<i>Meaning</i>
open	Connect to a remote computer	set	Set the operating parameters
close	Close the connection	status	Display the status information
display	Show the operating parameters	send	Send special characters
mode	Change to line or character mode	quit	Exit TELNET

Table 15.5 Examples of interface commands

15.1.6 SECURE SHELL (SSH)

- Although Secure Shell (SSH) is a secure application program that can be used today for several purposes such as remote logging and file transfer, it was originally designed to replace TELNET.
- There are two versions of SSH: SSH-1 and SSH-2, which are totally incompatible.
- The first version, SSH-1, is now deprecated because of security flaws in it.
- In this section, we discuss only SSH-2
- **Components:**

SSH is an application-layer protocol with three components, as shown in Figure 15.11.



- **SSH Transport-Layer Protocol (SSH-TRANS)**
 - Since TCP is not a secured transport-layer protocol, SSH first uses a protocol that creates a secured channel on top of the TCP.
 - This new layer is an independent protocol referred to as SSH-TRANS.
 - When the procedure implementing this protocol is called, the client and server first use the TCP protocol to establish an insecure connection.
 - Then they exchange several security parameters to establish a secure channel on top of the TCP
 - list of the services provided by this protocol:
 1. Privacy or confidentiality of the message exchanged
 2. Data integrity, which means that it is guaranteed that the messages exchanged between the client and server are not changed by an intruder
 3. Server authentication, which means that the client is now sure that the server is the one that it claims to be
 4. Compression of the messages, which improves the efficiency of the system and makes attack more difficult
- **SSH Authentication Protocol (SSH-AUTH)**
 - After a secure channel is established between the client and the server and the server is authenticated for the client, SSH can call another procedure that can authenticate the client for the server.
 - The client authentication process in SSH is very similar to what is done in Secure Socket Layer (SSL).

- This layer defines a number of authentication tools similar to the ones used in SSL.
- Authentication starts with the client, which sends a request message to the server.
- The request includes the user name, server name, the method of authentication, and the required data.
- The server responds with either a success message, which confirms that the client is authenticated, or a failed message, which means that the process needs to be repeated with a new request message.
- **SSH Connection Protocol (SSH-CONN)**
 - After the secured channel is established and both server and client are authenticated for each other, SSH can call a piece of software that implements the third protocol, SSH CONN.
 - One of the services provided by the SSH-CONN protocol is multiplexing. SSH-CONN takes the secure channel established by the two previous protocols and lets the client create multiple logical channels over it.
 - Each channel can be used for a different purpose, such as remote logging, file transfer, and so on.
- **Applications**

Although SSH is often thought of as a replacement for TELNET, SSH is, in fact, a general-purpose protocol that provides a secure connection between a client and server
- **SSH for Remote Logging**
 - Several free and commercial applications use SSH for remote logging.
 - Among them, we can mention PuTTY, by Simon Tatham, which is a client SSH program that can be used for remote logging.
 - Another application program is Tectia, which can be used on several platforms.
- **SSH for File Transfer**
 - One of the application programs that is built on top of SSH for file transfer is the Secure File Transfer Program (sftp).
 - The sftp application program uses one of the channels provided by the SSH to transfer files.
 - Another common application is called Secure Copy (scp).

- This application uses the same format as the UNIX copy command, cp, to copy files.
- **Port Forwarding**
 - One of the interesting services provided by the SSH protocol is port forwarding.
 - We can use the secured channels available in SSH to access an application program that does not provide security services.
 - Applications such as TELNET and Simple Mail Transfer Protocol (SMTP), can use the services of the SSH port forwarding mechanism.
 - The SSH port forwarding mechanism creates a tunnel through which the messages belonging to other protocols can travel.
 - For this reason, this mechanism is sometimes referred to as SSH tunneling.
- **Format of the SSH Packets**
 - Figure 15.12 shows the format of packets used by the SSH protocols.

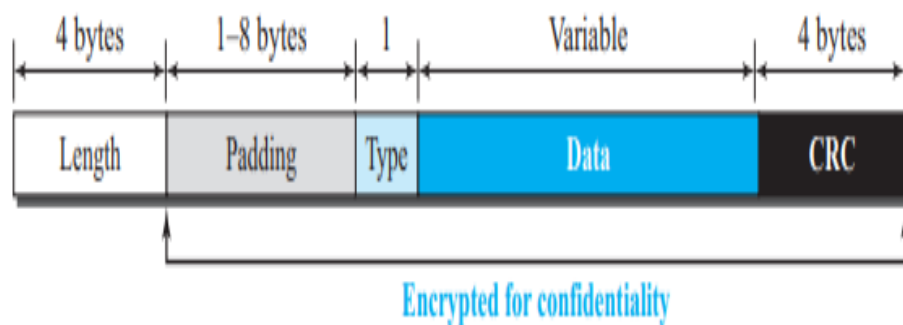


Figure 15.12 SSH Packets

- The length field defines the length of the packet but does not include the padding.
- One to eight bytes of padding is added to the packet to make the attack on the security provision more difficult.
- The cyclic redundancy check (CRC) field is used for error detection.
- The type field designates the type of the packet used in different SSH protocols.
- The data field is the data transferred by the packet in different protocols.

15.1.7 DOMAIN NAME SYSTEM (DNS)

- The last client-server application program we discuss has been designed to help other application programs.
- To identify an entity, TCP/IP protocols use the IP address, which uniquely identifies the connection of a host to the Internet.
- However, people prefer to use names instead of numeric addresses.
- Therefore, the Internet needs to have a directory system that can map a name to an address.
- This is analogous to the telephone network.
- A telephone network is designed to use telephone numbers, not names.
- People can either keep a private file to map a name to the corresponding telephone number or can call the telephone directory to do so.
- The directory system in the Internet can map names to IP addresses.
- Since the Internet is so huge today, a central directory system cannot hold all the mapping.
- In addition, if the central computer fails, the whole communication network will collapse.
- A better solution is to distribute the information among many computers in the world.
- In this method, the host that needs mapping can contact the closest computer holding the needed information.
- This method is used by the Domain Name System (DNS).
- Figure 15.13 shows how TCP/IP uses a DNS client and a DNS server to map a name to an address.
- A user wants to use a file transfer client to access the corresponding file transfer server running on a remote host.
- The user knows only the file transfer server name, such as abcdefxyz.com.
- However, the TCP/IP suite needs the IP address of the file transfer server to make the connection.

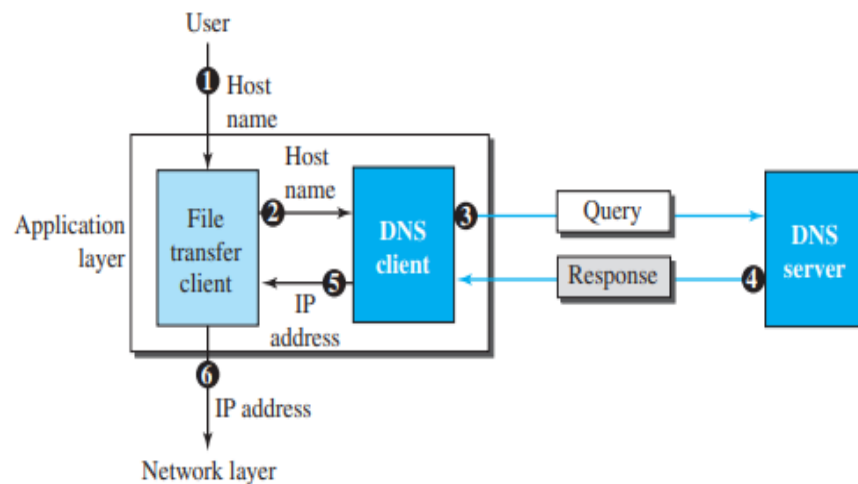


Figure 15.13 Purpose of DNS

- The following six steps map the host name to an IP address:
 1. The user passes the host name to the file transfer client.
 2. The file transfer client passes the host name to the DNS client.
 3. Each computer, after being booted, knows the address of one DNS server. The DNS client sends a message to a DNS server with a query that gives the file transfer server name using the known IP address of the DNS server.
 4. The DNS server responds with the IP address of the desired file transfer server.
 5. The DNS server passes the IP address to the file transfer client.
 6. The file transfer client now uses the received IP address to access the file transfer server.
- Note that the purpose of accessing the Internet is to make a connection between the file transfer client and server, but before this can happen, another connection needs to be made between the DNS client and DNS server.
- In other words, we need at least two connections in this case.
- The first is for mapping the name to an IP address; the second is for transferring files.
- **Name Space**
 - To be unambiguous, the names assigned to machines must be carefully selected from a name space with complete control over the binding between the names and IP addresses.
 - In other words, the names must be unique because the addresses are unique.

- A name space that maps each address to a unique name can be organized in two ways: flat or hierarchical.
- **In a flat name space, a name is assigned to an address.**
 - A name in this space is a sequence of characters without structure.
 - The names may or may not have a common section; if they do, it has no meaning.
 - The main disadvantage of a flat name space is that it cannot be used in a large system such as the Internet because it must be centrally controlled to avoid ambiguity and duplication.
 - In a hierarchical name space, each name is made of several parts.
 - The first part can define the nature of the organization, the second part can define the name of an organization, the third part can define departments in the organization, and so on.
 - In this case, the authority to assign and control the name spaces can be decentralized.
 - A central authority can assign the part of the name that defines the nature of the organization and the name of the organization.
 - The responsibility for the rest of the name can be given to the organization itself.
 - The organization can add suffixes (or prefixes) to the name to define its host or resources.
 - The management of the organization need not worry that the prefix chosen for a host is taken by another organization because, even if part of an address is the same, the whole address is different.
 - For example, assume two organizations call one of their computers caesar.
 - The first organization is given a name by the central authority, such as first.com, the second organization is given the name second.com.
 - When each of these organizations adds the name caesar to the name they have already been given, the end result is two distinguishable names: caesar.first.com and caesar.second.com.
 - The names are unique.
- **Domain Name Space**
 - To have a hierarchical name space, a domain name space was designed.

- In this design the names are defined in an inverted-tree structure with the root at the top.
- The tree can have only 128 levels: level 0 (root) to level 127 (see Figure 15.14).

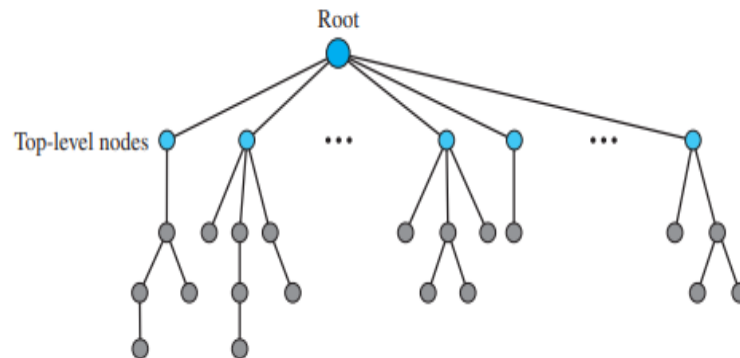


Figure 15.14 Domain name space

- **Label**
 - Each node in the tree has a label, which is a string with a maximum of 63 characters.
 - The root label is a null string (empty string).
 - DNS requires that children of a node (nodes that branch from the same node) have different labels, which guarantees the uniqueness of the domain names.
- **Domain Name**
 - Each node in the tree has a domain name.
 - A full domain name is a sequence of labels separated by dots (.).
 - The domain names are always read from the node up to the root.
 - The last label is the label of the root (null).
 - This means that a full domain name always ends in a null label, which means the last character is a dot because the null string is nothing.
 - Figure 15.15 shows some domain names.
 - If a label is terminated by a null string, it is called a fully qualified domain name (FQDN).
 - The name must end with a null label, but because null means nothing, the label ends with a dot.
 - If a label is not terminated by a null string, it is called a partially qualified domain name (PQDN).

- A PQDN starts from a node, but it does not reach the root.
- It is used when the name to be resolved belongs to the same site as the client.
- Here the resolver can supply the missing part, called the suffix, to create an FQDN.

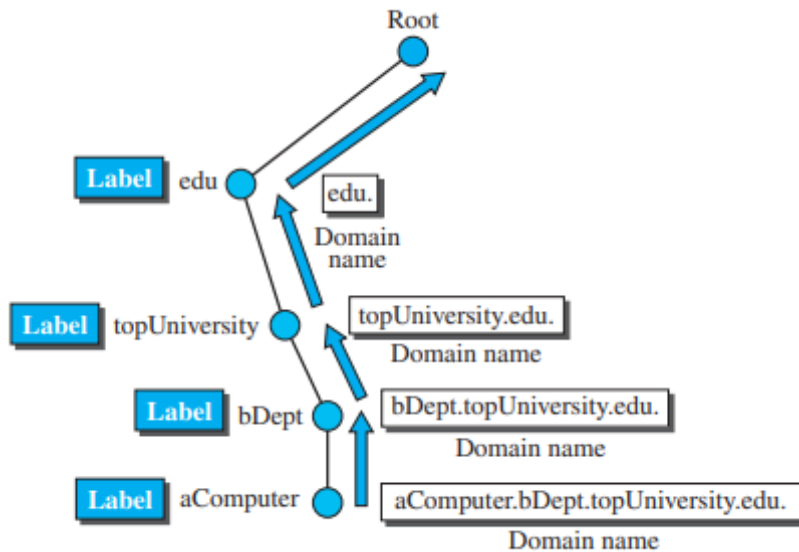


Figure 15.15 Domain names and labels

- **Domain**
 - A domain is a subtree of the domain name space.
 - The name of the domain is the name of the node at the top of the subtree. Figure 15.16 shows some domains.
 - Note that a domain may itself be divided into domains.

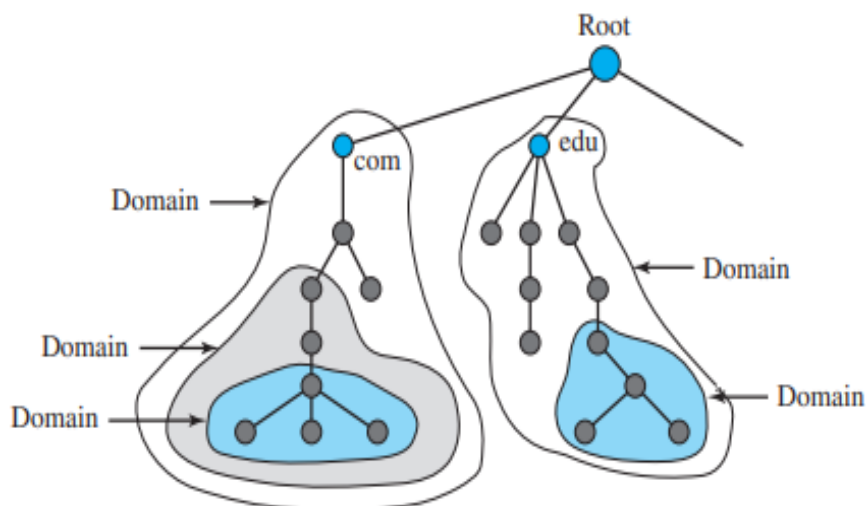


Figure 15.16 Domains

- **Distribution of Name Space**
 - The information contained in the domain name space must be stored.
 - However, it is very inefficient and also not reliable to have just one computer store such a huge amount of information.
 - It is inefficient because responding to requests from all over the world places a heavy load on the system.
- It is not reliable because any failure makes the data inaccessible.
- **Hierarchy of Name Servers**
 - The solution to these problems is to distribute the information among many computers called DNS servers.
 - One way to do this is to divide the whole space into many domains based on the first level.
 - In other words, we let the root stand alone and create as many domains (subtrees) as there are first-level nodes.
 - Because a domain created this way could be very large, DNS allows domains to be divided further into smaller domains (subdomains).
 - Each server can be responsible (authoritative) for either a large or small domain.
 - In other words, we have a hierarchy of servers in the same way that we have a hierarchy of names (see Figure 15.17)

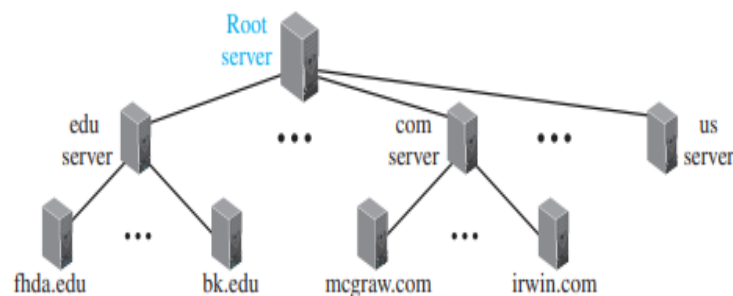


Figure 15.17 Hierarchy of name servers

- **Zone**
 - Since the complete domain name hierarchy cannot be stored on a single server, it is divided among many servers.
 - What a server is responsible for or has authority over is called a zone.
 - We can define a zone as a contiguous part of the entire tree.

- If a server accepts responsibility for a domain and does not divide the domain into smaller domains, the “domain” and the “zone” refer to the same thing.
- The server makes a data base called a zone file and keeps all the information for every node under that domain.
- However, if a server divides its domain into subdomains and delegates part of its authority to other servers, “domain” and “zone” refer to different things.
- The information about the nodes in the subdomains is stored in the servers at the lower levels, with the original server keeping some sort of reference to these lower-level servers.
- Of course, the original server does not free itself from responsibility totally.
- It still has a zone, but the detailed information is kept by the lower-level servers (see Figure 15.18)

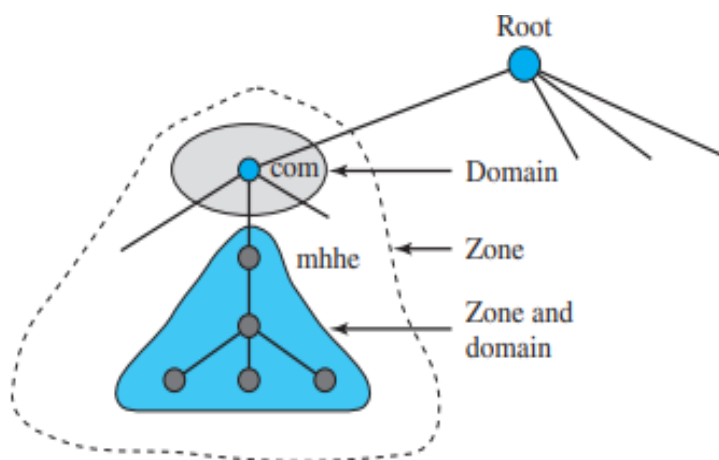


Figure 15.18 Zone

- **Root Server**
 - A root server is a server whose zone consists of the whole tree.
 - A root server usually does not store any information about domains but delegates its authority to other servers, keeping references to those servers.
 - There are several root servers, each covering the whole domain name space.
 - The root servers are distributed all around the world
- **Primary and Secondary Servers**
 - DNS defines two types of servers: primary and secondary.

- A primary server is a server that stores a file about the zone for which it is an authority.
- It is responsible for creating, maintaining, and updating the zone file.
- It stores the zone file on a local disk.
- A secondary server is a server that transfers the complete information about a zone from another server (primary or secondary) and stores the file on its local disk.
- The secondary server neither creates nor updates the zone files.
- If updating is required, it must be done by the primary server, which sends the updated version to the secondary.
- The primary and secondary servers are both authoritative for the zones they serve.
- The idea is not to put the secondary server at a lower level of authority but to create redundancy for the data so that if one server fails, the other can continue serving clients.
- Note also that a server can be a primary server for a specific zone and a secondary server for another zone.
- Therefore, when we refer to a server as a primary or secondary server, we should be careful about which zone we refer to.
- **DNS in the Internet**
 - DNS is a protocol that can be used in different platforms.
 - In the Internet, the domain name space (tree) was originally divided into three different sections: generic domains, country domains, and the inverse domains.
 - However, due to the rapid growth of the Internet, it became extremely difficult to keep track of the inverse domains, which could be used to find the name of a host when given the IP address.
 - The inverse domains are now deprecated.
- **Generic Domains**
 - The generic domains define registered hosts according to their generic behavior.
 - Each node in the tree defines a domain, which is an index to the domain name space database (see Figure 15.19).

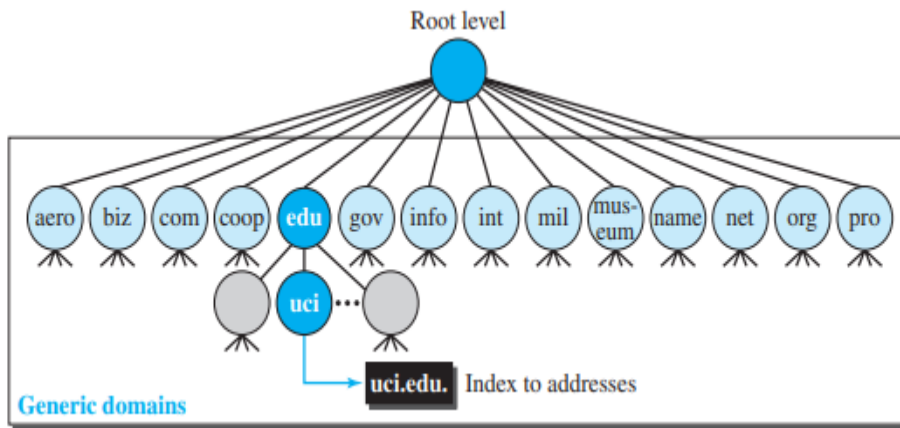


Figure 15.19 Generic domains

- Looking at the tree, we see that the first level in the generic domains section allows 14 possible labels.
- These labels describe the organization types as listed in Table 15.6.

<i>Label</i>	<i>Description</i>	<i>Label</i>	<i>Description</i>
aero	Airlines and aerospace	int	International organizations
biz	Businesses or firms	mil	Military groups
com	Commercial organizations	museum	Museums
coop	Cooperative organizations	name	Personal names (individuals)
edu	Educational institutions	net	Network support centers
gov	Government institutions	org	Nonprofit organizations
info	Information service providers	pro	Professional organizations

Table 15.6 Generic domain labels

- **Country Domains**
 - The country domains section uses two-character country abbreviations (e.g., us for United States).
 - Second labels can be organizational, or they can be more specific national designations.
 - The United States, for example, uses state abbreviations as a sub division of us (e.g., ca.us.).
 - Figure 15.20 shows the country domains section.
 - The address uci.ca.us. can be translated to University of California, Irvine, in the state of California in the United States

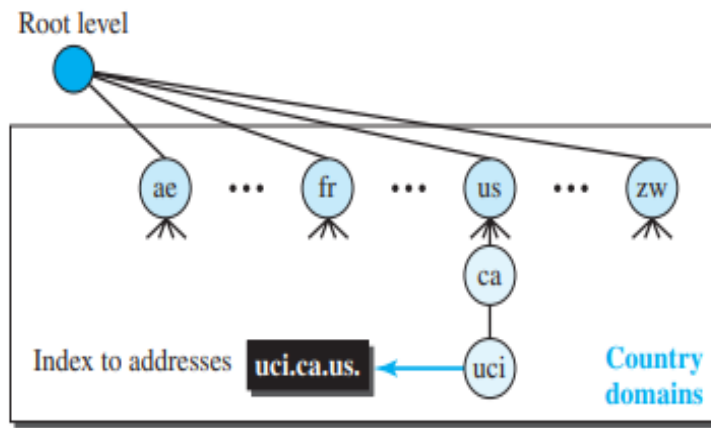


Figure 15.20 Country domains

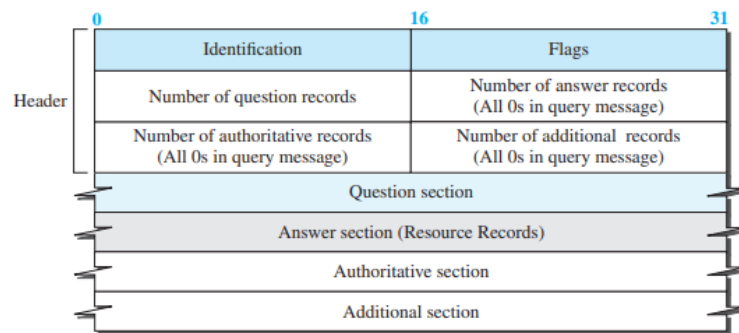
- **Resolution**
 - Mapping a name to an address is called name-address resolution.
 - DNS is designed as a client-server application.
 - A host that needs to map an address to a name or a name to an address calls a DNS client called a resolver.
 - The resolver accesses the closest DNS server with a mapping request.
 - If the server has the information, it satisfies the resolver; otherwise, it either refers the resolver to other servers or asks other servers to provide the information.
 - After the resolver receives the mapping, it interprets the response to see if it is a real resolution or an error, and finally delivers the result to the process that requested it.
 - A resolution can be either recursive or iterative.
- **Caching**
 - Each time a server receives a query for a name that is not in its domain, it needs to search its database for a server IP address.
 - Reduction of this search time would increase efficiency.
 - DNS handles this with a mechanism called caching.
 - When a server asks for a mapping from another server and receives the response, it stores this information in its cache memory before sending it to the client.
 - If the same or another client asks for the same mapping, it can check its cache memory and resolve the problem.

- However, to inform the client that the response is coming from the cache memory and not from an authoritative source, the server marks the response as unauthoritative.
- Caching speeds up resolution, but it can also be problematic.
- If a server caches a mapping for a long time, it may send an outdated mapping to the client.
- To counter this, two techniques are used.
- First, the authoritative server always adds information to the mapping called time to live (TTL).
- It defines the time in seconds that the receiving server can cache the information.
- After that time, the mapping is invalid and any query must be sent again to the authoritative server.
- Second, DNS requires that each server keep a TTL counter for each mapping it caches.
- The cache memory must be searched periodically and those mappings with an expired TTL must be purged
- **Resource Records**
 - The zone information associated with a server is implemented as a set of resource records.
 - In other words, a name server stores a database of resource records.
 - A resource record is a 5-tuple structure, as shown below:
(Domain Name, Type, Class, TTL, Value)
- The domain name field is what identifies the resource record. The value defines the information kept about the domain name.
- The TTL defines the number of seconds for which the information is valid.
- The class defines the type of network; we are only interested in the class IN (Internet).
- The type defines how the value should be interpreted.
- Table 15.7 lists the common types and how the value is interpreted for each type.

Type	Interpretation of value
A	A 32-bit IPv4 address (see Chapter 18)
NS	Identifies the authoritative servers for a zone
CNAME	Defines an alias for the official name of a host
SOA	Marks the beginning of a zone
MX	Redirects mail to a mail server
AAAA	An IPv6 address (see Chapter 22)

Table 15.7 Types

- DNS Messages
- To retrieve information about hosts, DNS uses two types of messages: query and response.
- Both types have the same format as shown in Figure 15.21.

**Figure 15.21 DNS message**

- The identification field is used by the client to match the response with the query.
- The flag field defines whether the message is a query or response.
- It also includes status of error.
- The next four fields in the header define the number of each record type in the message.
- The question section consists of one or more question records.
- It is present in both query and response messages.
- The answer section consists of one or more resource records.
- It is present only in response messages.
- The authoritative section gives information (domain name) about one or more authoritative servers for the query.

- The additional information section provides additional information that may help the resolver.
- **Security of DNS**
 - DNS is one of the most important systems in the Internet infrastructure; it provides crucial services to Internet users.
 - Applications such as Web access or e-mail are heavily dependent on the proper operation of DNS.
 - DNS can be attacked in several ways including:
 1. The attacker may read the response of a DNS server to find the nature or names of sites the user mostly accesses. This type of information can be used to find the user's profile. To prevent this attack, DNS messages need to be confidential (see Chapters 31 and 32).
 2. The attacker may intercept the response of a DNS server and change it or create a totally new bogus response to direct the user to the site or domain the attacker wishes the user to access. This type of attack can be prevented using message origin authentication and message integrity.
 3. The attacker may flood the DNS server to overwhelm it or eventually crash it. This type of attack can be prevented using the provision against denial-of-service attack.
 - To protect DNS, IETF has devised a technology named DNS Security (DNSSEC) that provides message origin authentication and message integrity using a security service called digital signature.
 - DNSSEC, however, does not provide confidentiality for the DNS messages.
 - There is no specific protection against the denial-of- service attack in the specification of DNSSEC.
 - However, the caching system protects the upper-level servers against this attack to some extent.

15.1.8 SNMP

- Network management plays an important role in the Internet as it becomes larger and larger.
- The failure of a single device may interrupt the communication from one point of the Internet to the other.

- The concept of network management and discusses five general areas of network management: configuration, fault, performance, security, and accounting.
- Configuration management is related to the status of each entity and its relationship to other entities.
- Fault management is the area of network management that handles issues related to interruptions in the system.
- Performance management tries to monitor and control the network to ensure that it is running as efficiently as possible.
- Security management is responsible for controlling access to the network based on predefined policy.
- Accounting management is the controlling of users' access to network resources through charges.
- Simple Network Management Protocol (SNMP) as a framework for managing devices in an internet using the TCP/IP protocol suite.
- It shows how a manager as a host runs an SNMP client and any agents as a router or host runs a server program.
- The section defines the three components of the management protocol in the Internet.
- The section also defines Structure of Management Information (SMI) as the language that specifies how data types and objects in SNMP should be identified.
- Finally, the section introduces Management Information Base (MIB), which designates the objects to be managed in SNMP according to the rules defined in SMI.
- **SNMP**
 - Several network management standards have been devised during the last few decades.
 - The most important one is Simple Network Management Protocol (SNMP), used by the Internet.
 - SNMP is a framework for managing devices in an internet using the TCP/IP protocol suite.
 - It provides a set of fundamental operations for monitoring and maintaining an internet.
 - SNMP uses the concept of manager and agent.
 - That is, a manager, usually a host, controls and monitors a set of agents, usually routers or servers (see Figure 15.22)

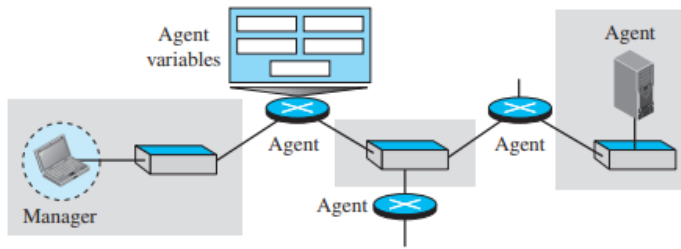


Figure 15.22 SNMP concept

- SNMP is an application-level protocol in which a few manager stations control a set of agents.
- The protocol is designed at the application level so that it can monitor devices made by different manufacturers and installed on different physical networks.
- In other words, SNMP frees management tasks from both the physical characteristics of the managed devices and the underlying networking technology.
- It can be used in a heterogeneous internet made of different LANs and WANs connected by routers made by different manufacturers.
- **Managers and Agents**
 - A management station, called a manager, is a host that runs the SNMP client program.
 - A managed station, called an agent, is a router (or a host) that runs the SNMP server program.
 - Management is achieved through simple interaction between a manager and an agent.
 - The agent keeps performance information in a database. The manager has access to the values in the database.
 - For example, a router can store in appropriate variables the number of packets received and forwarded.
 - The manager can fetch and compare the values of these two variables to see if the router is congested or not.
 - The manager can also make the router perform certain actions.
 - For example, a router periodically checks the value of a reboot counter to see when it should reboot itself.
 - It reboots itself, for example, if the value of the counter is 0.

- The manager can use this feature to reboot the agent remotely at any time.
- It simply sends a packet to force a 0 value in the counter.
- Agents can also contribute to the management process.
- The server program running on the agent can check the environment and, if it notices something unusual, it can send a warning message (called a Trap) to the manager.
- In other words, management with SNMP is based on three basic ideas:
 1. A manager checks an agent by requesting information that reflects the behavior of the agent.
 2. A manager forces an agent to perform a task by resetting values in the agent database.
 3. An agent contributes to the management process by warning the manager of an unusual situation
- **Management Components**
 - To do management tasks, SNMP uses two other protocols:
 - Structure of Management Information (SMI) and Management Information Base (MIB).
 - In other words, management on the Internet is done through the cooperation of three protocols: SNMP, SMI, and MIB, as shown in Figure 15.23

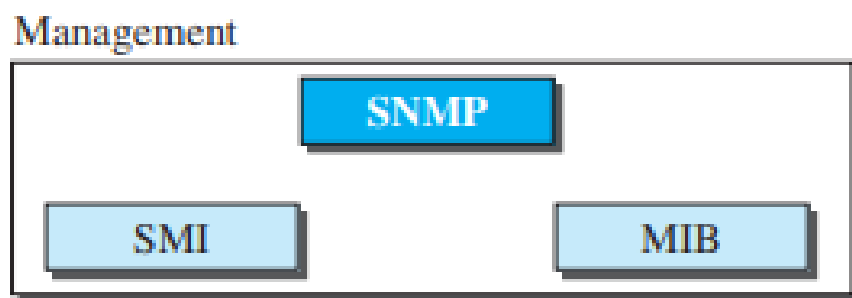


Figure 15.23 Components of network management on the Internet

- **Role of SNMP**
 - SNMP has some very specific roles in network management.
 - It defines the format of the packet to be sent from a manager to an agent and vice versa.

- It also interprets the result and creates statistics (often with the help of other management software).
- The packets exchanged contain the object (variable) names and their status (values).
- SNMP is responsible for reading and changing these values.
- **SNMP**
 - SNMP uses both SMI and MIB in Internet network management.
 - It is an application program that allows:
 - A manager to retrieve the value of an object defined in an agent.
 - A manager to store a value in an object defined in an agent.
 - An agent to send an alarm message about an abnormal situation to the manager.
- **PDUs**
 - SNMPv3 defines eight types of protocol data units (or PDUs): GetRequest, GetNextRequest, GetBulkRequest, SetRequest, Response, Trap, InformRequest, and Report (see Figure 15.24)

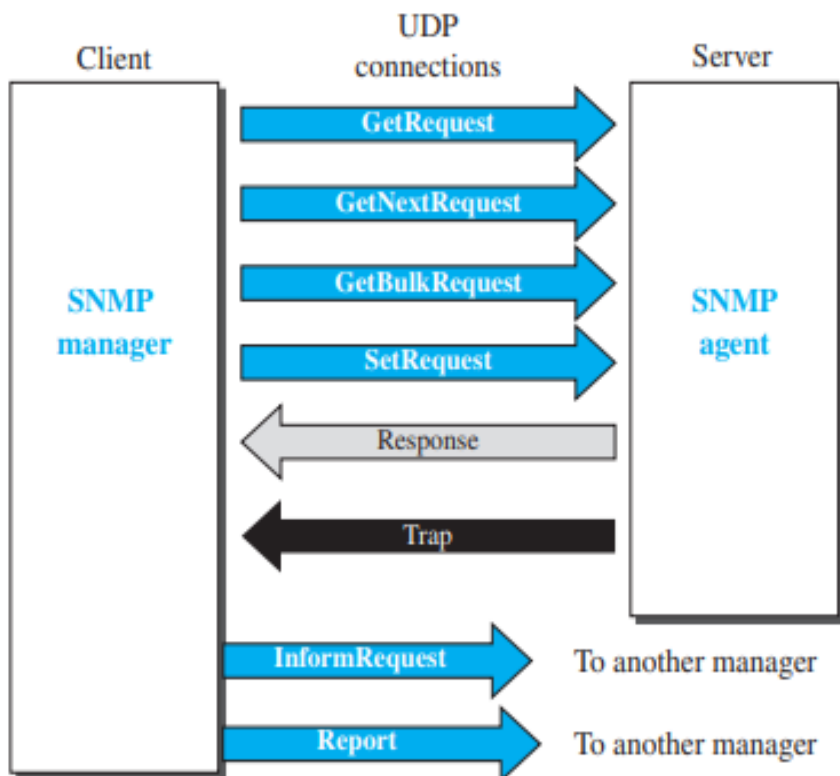


Figure 15.24 SNMP PDUs

15.3 LIST OF REFERENCES

1. <https://ieeexplore.ieee.org/document/756745>
2. <https://ieeexplore.ieee.org/document/678239>

15.2 BIBLIOGRAPHY

1. Data Communications and Networking, Behrouz A. Forouzan, Fifth Edition, TMH, 2018.
2. Computer Network, Bhushan Trivedi, Oxford University Press, 2016

15.3 UNIT END EXERCISES

1. In FTP, can a server retrieve a file from the client site?
2. Write note on DNS.
3. What is the World Wide Web?
4. Explain HTTP.
5. Explain the format of Electronic Mail.
6. What is SSH?
7. Write a note on TELNET.
