

CIT 305: NETWORKING AND COMMUNICATION TECHNOLOGY





NATIONAL OPEN UNIVERSITY OF NIGERIA

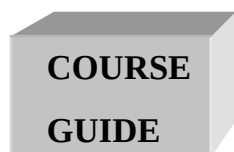
FACULTY OF SCIENCE

DEPARTMENT OF COMPUTER SCIENCE

COURSE CODE: CIT305

COURSE TITLE:

NETWORKING AND COMMUNICATION TECHNOLOGY



CIT305

NETWORKING AND COMMUNICATION TECHNOLOGY

CourseTeam Prof A.S. Sodiya (Developer/Writer)–FUNAAB
Dr. A. A. Afolorunso (Coordinator)-NOUN

CONTENTS

PAGE

Introduction.....	1	
WhatYouWillLearnin ThisCourse.....	1	
Course Aims.....	1	
CourseObjectives.....	1	
WorkingthroughThisCourse.....	2	
CourseMaterials.....	2	2
StudyUnits		
AssignmentFile.....	3	
Tutor-MarkedAssignments(TMAs).....	3	
ExaminationandGrading.....	4	
PresentationSchedule.....	4	
CourseMarkingScheme.....	5	
CourseOverview.....	5	
Howto Getthe MostfromThisCourse.....	6	
Facilitators/TutorsandTutorials.....		7
Summary.....	8	

Introduction

CIT305–Networking and Communication Technology is a three[3] credit unit course of fifteen units. It teaches the various forms of networking, network design, communication technology used by people to accomplish different organisational or individual tasks.

It also gives an insight into various forms of computer networking ranging from LAN to WAN and even go as far as looking into the wireless networks that are in use in today's technology. The course also explains the enterprise network, which forms a branch computer networking technology.

What You Will Learn in This Course

The main purpose of this course is to provide the necessary tools for designing and managing *information systems*. It makes available the steps and tools that will enable you to make proper and accurate decisions on database designs and operations whenever the need arises.

Course Aims

This course sets out to achieve the following aims:

- introduce the concepts associated with information systems development;
- provide necessary tools for analysing, designing, developing a database of any size;
- provide you with the necessary foundation in Database programming
- introduction of web services and their architectural frameworks;

Course Objectives

Certain objectives have been set out to ensure that the course achieves its aims. Apart from the course objectives, every unit of this course has set objectives. In the course of the study, you will need to confirm, at the end of each unit, if you have met the objectives set at the beginning of each unit. By the end of this course you should be able to:

- explain the fundamentals of Networking
- define the term “data links”
- explain network protocols

- discuss what is involved in the building of internet networks using
 - TCP/IP and routers
- explain the network standards (IEEE 802 Standards)
- state the fundamentals of enterprise network
- discuss what is involved in signal transmission and impairment
- explain digital technology
- describe the concept of packet switching

Working through This Course

In order to have a thorough understanding of the study units, you will need to read and understand the contents, practise the steps by designing an information system of your own, and become committed to learning and implementing your knowledge.

This course is designed to cover approximately seventeen weeks, and it will require your devoted attention. You should do the exercises in the Tutor-Marked Assignments (TMAs) and submit to your tutors.

Course Materials

These include the list below:

1. Course guide
2. Study units
3. Recommended texts

4. A file for your assignments and for records to monitor your progress.

Study Units

There are Twentythree study units in this course-as listed below.

Module 1 Networking

- Unit 1 Introduction to Networking
- Unit 2 Data Links
- Unit 3 Deploying Physical Media
- Unit 4 Network Protocols

Module 2 Network Design

- Unit 1 Harnessing Wifi for User Mobility
- Unit 2 Building Internet Networks Using Tcp/Ip and Routers
- Unit 3 Network Standards (Ieee 802 Standards)
- Unit 4 Implementing Security Best Practices

Module3 EnterpriseNetwork

- Unit1 CreatingEnterpriseNetwork
- Unit2 PlanningandSelectionof EnterpriseNetwork
- Unit3 LANandWAN

Module4 CommunicationTechnology

- Unit1 ModemandModulationConcepts
- Unit2 Multiplexers
- Unit3 DigitalTechnologies
- Unit4 SignalTransmissionandImpairment

Module5 NetworkTechnologies

- Unit1 *ISDN*
- Unit2 *DSL*
- Unit3 *SONET*
- Unit4 Packet Switching
- Unit5 Internetand*TCP/IP*

Make use of the course materials, do the exercises to enhance your learning.

Assignments File

These are of two types - the self-assessment exercises and the Tutor-Marked Assignments (TMAs). The self-assessment exercises will enable you to monitor your performance by yourself, while the Tutor-Marked Assignment (TMA) is a supervised assignment. The assignments taken a certain percentage of your total score in this course. The tutor-marked assignments will be assessed by your tutor within a specified period.

The examination at the end of this course will aim at determining the level of mastery of the subject matter. This course includes twelve tutor-marked assignments and each must be done and submitted accordingly. Your best scores however, will be recorded for you. Be sure to send these assignments to your tutor before the deadline to avoid loss of marks.

Tutor-Marked Assignments (TMAS)

There are twelve tutor-marked assignments in this course. You need to submit all the assignments. The total marks for the best four (4) assignments will be 30% of your total course mark.

Assignment questions for the units in this course are contained in the Assignment File. You should be able to complete your assignments from the information and materials contained in your set textbooks, reading and study units. However, you may wish to use other references to broaden your viewpoint and provide a deeper understanding of the subject.

When you have completed each assignment, send it together with form to your tutor. Make sure that each assignment reaches your tutor on or before the deadline given. If, however, you cannot complete your work on time, contact your tutor before the assignment is done to discuss the possibility of an extension.

Examination and Grading

The final examination for the course will carry 70% percentage of the total marks available for this course. The examination will cover every aspect of the course, so you are advised to revise all your corrected assignments before the examination.

This course endows you with the status of a teacher and that of a learner. This means that you teach yourself and that you learn, as your learning capabilities would allow. It also means that you are in a better position to determine and to ascertain the what, the how, and the when of your language learning. No teacher imposes any method of learning on you.

The course units are similarly designed with the introduction following the table of contents, then a set of objectives and then the dialogue and so on.

The objectives guide you as you go through the unit to ascertain your knowledge of the required terms and expressions.

Presentation Schedule

This gives you the important dates for the completion of tutor-marked assignments and tutorials. Remember, you are required to submit all your assignments by the due date. You should guard against lagging behind in your work.

Course Marking Scheme

This table shows the breakdown of the actual marking scheme for the course.

Assessment	Marks
Assignment 1- 4	Four assignments, best three marks of the four count at 20% of course marks
Final Examination	70% of overall course marks
Total	100% of course marks

Course Overview

Unit	Title of Work	Weeks	Assessment (End of Unit)
	Course Guide	Week 1	
Module 1			
1	Introduction to Networking	Week 1	Assignment 1
2	Data Links	Week 1	Assignment 1
3	Deploying Physical Media	Week 2	Assignment 2
4	Network Protocols	Week 3	Assignment 3
Module 2			
1	Harnessing Wi-Fi for user mobility	Week 4	Assignment 4
2	Building Internet networks using TCP/IP and DNS	Week 5	Assignment 5
3	Network Standards (IEEE 802 Standards)	Week 6	Assignment 6
4	Implementing Security best practices	Week 7	Assignment 7
Module 3			
1	Creating Enterprise Network	Week 8	Assignment 8
2	Planning and Selection of Enterprise Network		
3	Advanced WAN and LAN Classes	Week 9	Assignment 9
Module 4			
1	Modem and Modulation Concepts	Week 10	Assignment 10
2	Multiplexers	Week 11	Assignment 11

3	DigitalTechnologies	Week12	Assignment12
Module5			
1	IntegratedServicesDigitalNetwork (ISDN)	Week13	Assignment12
2	Digital SubscriberLine(DSL)	Week14	Assignment13
3	SynchronousOpticalNetwork(SONET)	Week15	Assignment14
4	Packet Switching	Week16	Assignment15
5	InternetandTCP/IP	Week16	Assignment16

	Revision	Week17	
	Examination		
	Total	17Weeks	

How to Get the Most from This Course

In distance learning the study units replace the university lecturer room situation. This is one of the great advantages of distance learning; you can read and work through specially designed study materials at your own pace, and at a time and place that suits you best. Think of it as reading the lecture instead of listening to a lecturer. In the same way that a lecturer might set you some reading to do, the study unit tells you when to read your set books or other material. Just as a lecturer might give you an in-class exercise, your study units provide exercises for you to do at appropriate points.

Each of the study units follows a common format. The first item is an introduction to the subject matter of the unit and how a particular unit is integrated with the other units and the course as a whole. Next is a set of learning objectives. These objectives enable you to know what you should be able to do by the time you have completed the unit. You should use these objectives to guide your study. When you have finished the units you must go back and check whether you have achieved the objectives. If you make a habit of doing this you will significantly improve your chances of passing the course.

Remember that your tutor's job is to assist you. When you need help, don't hesitate to call and ask your tutor to provide it.

1. Read this *Course Guide* thoroughly.
2. Organise a study schedule. Refer to the 'course overview' for more details. Note the time you are expected to spend on each unit and how the assignments relate to the units. Whatever method you choose to use, you should decide on it and write in your own dates for working on each unit.
3. Once you have created your own study schedule, do everything you can to stick to it. The major reason that students fail is that they lag behind in their course work.
4. Assemble the study materials. Information about what you need for a unit is given in the 'overview' at the beginning of each unit. You will almost always need both the study unit you are working on and one of your set of books on your desk at the same time.

5. Work through the unit. The content of the unit itself has been arranged to provide a sequence for you to follow. As you work through the unit you will be instructed to read sections from your set books or other articles. Use the unit to guide your reading.
6. Review the objectives for each study unit to confirm that you have achieved them. If you feel unsure about any of the objectives, review the study material or consult your tutor.
7. When you are confident that you have achieved a unit's objectives, you can then start on the next unit. Proceed unit by unit through the course and try to pace your study so that you keep yourself on schedule.
8. When you have submitted an assignment to your tutor for marking, do not wait for its return before starting on the next unit. Keep to your schedule. When the assignment is returned, pay particular attention to your tutor's comments, both on the tutor-marked assignment form and also written on the assignment. Consult your tutor as soon as possible if you have any questions or problems.
9. After completing the last unit, review the course and prepare yourself for the final examination. Check that you have achieved the unit objectives (listed at the beginning of each unit) and the course objectives (listed in this *Course Guide*).

Facilitators/Tutors and Tutorials

There are 12 hours of tutorials provided in support of this course. You will be notified of the dates, times and location of these tutorials, together with the name and phone number of your tutor, as soon as you are allocated a tutorial group.

Your tutor will mark and comment on your assignments, keep a close watch on your progress and on any difficulties you might encounter and provide assistance to you during the course. You must mail or submit your tutor-marked assignments to your tutor well before the due date (at least two working days are required). They will be marked by your tutor and returned to you as soon as possible.

Do not hesitate to contact your tutor by telephone, or e-mail if you need help. The following may be circumstances in which you would find help necessary. Contact your tutor if:

- you do not understand any part of the study units or the assigned readings,
- you have difficulty with the self-tests or exercises,
- you have a question or problem with an assignment, with your tutor's comments on an assignment or with the grading of an assignment.

You should try your best to attend the tutorials. This is the only chance to have face-to-face contact with your tutor and to ask questions which are answered instantly. You can raise any problem encountered in the course of your study. To gain the maximum benefit from course tutorials, prepare a question list before attending them. You will learn a lot from participating in discussions actively.

Summary

Networking and Communication Technology introduces you to the concepts associated with computer network development which is critical in understanding the various computer technology and data communication technology. The content of the course material has been planned and written to ensure that you acquire the proper knowledge and skills for the appropriate situations. The essence is to help you in acquiring the necessary knowledge and competence by equipping you with the necessary tools to accomplish this.

We hope that by the end of this course you would have acquired the required knowledge to view Computer Network in a new way.

I wish you success with the course and hope that you will find it both interesting and useful.

CourseCode CIT305
CourseTitle NetworkingandCommunicationTechnology

CourseTeam Prof . A.S.Sodiya(Developer/Writer)-
FUNAAB
Afolorunso,A.A.(Coordinator)-NOUN



NATIONAL OPEN UNIVERSITY OF NIGERIA

National Open University of Nigeria

Headquarters

14/16 Ahmadu Bello Way

Victoria Island

Lagos

Abuja Office

5 Dares Salaam Street

Off Aminu Kano Crescent

Wuse II, Abuja

e-mail: centralinfo@nou.edu.ng

URL: www.nou.edu.ng

Published By:

National Open University of Nigeria

First Printed 2012

Reviewed 2020

ISBN: 978-058-910-4

All Rights Reserved

CONTENTS**PAGE****Module1 Networking..... 1**

Unit1 Introductionto Networking..... 1

Unit2 DataLinks.....14

Unit3 DeployingPhysicalMedia.....20

Unit4 NetworkProtocols.....24

Module2 NetworkDesign.....31

Unit1 HarnessingWififorUserMobility..... 31

Unit2 BuildingInternetworksUsingTcp/IpandRouters..... 38

Unit3 NetworkStandards(Ieee802Standards)..... 46

Unit4 ImplementingSecurityBestPractices..... 49

Module3 EnterpriseNetwork..... 53

Unit2 PlanningandSelectionof EnterpriseNetwork..... 57

Unit3 LanandWan..... 63

Module4 CommunicationTechnology..... 67

Unit2 Multiplexers..... 79

Unit3 DigitalTechnologies..... 87

Unit4 SignalTransmissionandImpairment..... 91

Unit1 ISDN.....98

Unit2 DSL.....108

Unit3 SONET.....115

Unit4	Switching.....	122
Unit5	InternetandTCP/IP.....	127

MODULE1 NETWORKING

Unit1	Introductionto Networking
Unit2	Datalinks
Unit3	DeployingPhysicalMedia
Unit4	NetworkProtocols

UNIT1 INTRODUCTIONTONETWORKING

CONTENTS

1.0	Introduction
2.0	Objectives
3.0	MainContent
3.1	Overview
3.2	Definitionof Networking
3.3	Historyof ComputerNetworks
3.4	Purposeof ComputerNetworks
3.5	How Network Affect our lives
3.6	Network Representations
3.7	NetworkClassification
3.8	Typesof NetworkBasedonPhysicalScope
3.9	Introductionto TelecommunicationCircuit
4.0	Conclusion
5.0	Summary
6.0	Tutor-MarkedAssignment

7.0 References/Further Reading

1.0 INTRODUCTION

Having read through the course guide, you will have a general understanding of what this unit is about and how it fits into the course as a whole. This unit will describe the general concept of networking, types and its application areas.

2.0 OBJECTIVES

At the end of this unit, you should be able to:

- explain the term networking system
- identify the various types of networking
- highlight the history of networking
- discuss how network can affect lives
- describe the areas of work of networking.

3.0 MAINCONTENT

3.1 Overview

The concept of a network is pretty simple. A couple of computers have some cable strung between them, and send data back and forth using electrical signaling over the cable. More or less the same as telephones do or, in a very rough sense, like two kids speaking into tin cans connected by a string; but how does the data actually move from computer A to computer B? How does computer A find the physical location of computer B on the network? If they communicate with electrical signaling, so the data is traveling "at the speed of light"; why does it take so long to send a big file across the network?

The internet is not one single entity; it is a massive interconnection of hosts such as your computer, and another half-way across the world. We understand that our Internet Service Providers (ISPs) provide internet access to our host, but it may not necessarily provide service to the server that you are communicating with. In this case, how does information make its way from your host to the other host and vice versa?

The answer lies in the interconnection of many ISPs themselves. ISPs can be categorised into Tier 1, 2 and 3 ISPs. Tier 1 ISPs are major internet service providers that usually sell access to smaller Tier 2 ISPs. Tier 2 ISPs may service entire countries or cities, but not the rest of the world. Tier 3 ISPs are also customers of Tier 2 ISPs, and usually service end users such as you. ISPs peer with each other to allow data from your host to reach the other host. From this, we understand that data passes through multiple ISPs in order to be delivered from one location to another.

3.2 Definition of Networking

A computer network, often simply referred to as a network, is a collection of computers and devices interconnected by communications channels that facilitate communications among users and allow users to share resources.

A network can also be defined as a collection of computers or other hardware devices that are connected together, either physically or logically, using special hardware and software, to allow them to exchange information and cooperate. Networking is the term that describes the processes involved in designing, implementing, upgrading, managing and otherwise working with networks and network technologies.

3.3 History of Computer Networks

Before the advent of computer networks that were based upon some type of telecommunications system, communication between calculation machines and early computers was performed by human users by carrying instructions between them. Many of the social behaviours seen in today's Internet were demonstrably present in the nineteenth century and arguably in even earlier networks using visual signals.

In September 1940, George Stibitz used a teletype machine to send instructions for a problem set from his Model at Dartmouth College to his Complex Number Calculator in New York and received results back by the same means. Linking output systems like teletype to computers was an interest at the Advanced Research Projects Agency (ARPA) when, in 1962, J.C.R. Licklider was hired and developed a working group he called the "Intergalactic Network", a precursor to the ARPANET.

In 1964, researchers at Dartmouth developed the Dartmouth Time Sharing System for distributed users of large computer systems. The same year, at Massachusetts Institute of Technology, a research group supported by General Electric and Bell Labs used a computer to route and manage telephone connections.

Throughout the 1960s, Leonard Kleinrock, Paul Baran and Donald Davies independently conceptualised and developed network systems which used packets that could be used in a network between computer systems. In 1965, Thomas Merrill and Lawrence G. Roberts created the first wide area network (WAN).

The first widely used telephone switch that used true computer control was introduced by Western Electric in 1965. In 1969, the University of California at Los Angeles, the Stanford Research Institute, University of California at Santa Barbara, and the University of Utah were connected as the beginning of the ARPANET network using 50 kbit/s circuits. Commercial services using X.25 were deployed in 1972, and later used as an underlying infrastructure for expanding TCP/IP networks.

Today, computer networks are the core of modern communication. All modern aspects of the Public Switched Telephone Network (PSTN) are computer-controlled, and telephony increasingly runs over the Internet Protocol, although not necessarily the public Internet. The scope of communication has increased significantly in the past decade, and this boom in communications would not have been possible without the progressively advancing computer network. Computer networks and the technologies needed to connect and communicate through and between

them, continue to drive computer hardware, software, and peripherals industries. This expansion is mirrored by growth in the numbers and types of users of networks from the research to the home user.

3.4 Purpose of Computer Networks

Computer networks can be used for a variety of purposes; these are listed below.

i. Facilitating communications

Using a network, people can communicate efficiently and easily via email, instant messaging, chatrooms, telephone, videotelephone calls, and videoconferencing.

ii. Sharing hardware

In a networked environment, each computer on a network may access and use hardware resources on the network, such as printing a document on a shared network printer.

iii. Sharing files, data, and information

In a network environment, authorized users may access data and information stored on other computers on the network. The capability of providing access to data and information on shared storage devices is an important feature of many networks.

iv. Sharing software

Users connected to a network may run application programs on remote computers.

3.5 How Network Affect our lives

The followings are ways in which network can affect our lives:

1. Network connect us

Communication is almost as important to us as our reliance on air, water, food, and shelter. In today's world, through the use of networks, we are connected like never before.

2. Education

Networks Support the Way We Learn. The internet and personal computing have transformed the modern world. You do not have to be in school anymore to take a class, you do not have to be in a classroom to have a teacher.

3. Business

Network technology also is an essential part in business. It affects the way companies communicate and establish relations with their clients. In a fast moving and business environment, it is vital for them to interact with clients regularly and quickly to gain their trust and to obtain customer loyalty. With the use of Internet and online social networks, firms interact with consumers and answer all their queries about the product. Establishing effective communication with customers not only creates rapport with them, but it also creates strong public image

4. Television and Radio

Rapid growth of network technology affects the innovation of television and radio broadcasting causing the rapid increase in numbers cable, satellite television, fiber lines, satellite radio, video-on-demand and pay-per-view

5. Communication

Communication is an important part of the human's life. Nowadays, technology evolves rapidly and communication becomes easier than before. The effects of modern technology on how we communicate with one another are vast and varied, with both negative and positive impacts for our business and personal lives.

6. Internet

The Internet changed our life enormously; there is no doubt about that. Many advantages of the Internet show you the importance of this new medium. Internet

changed our life in many positive ways. One of the advantages of the internet is that we can create new social contact all around the world which we could not do it easily without the Internet. Such communities can also help people who cannot go out to find friends in the real life because they are disabled.

3.6 Network Components

A network can be as simple as a single cable connecting two computers or as complex as a collection of networks that span the globe.

Network infrastructure contains three broad categories of network components:

- i. Devices
 1. End Devices
 - An end device is where a message originates from or where it is received.
 - Data originates with an end device, flows through the network, and arrives at an end device
 - ii. Intermediary Network Devices

An intermediary device interconnects end devices in a network. Examples include: switches, wireless access points, routers, and firewalls.

- The management of data as it flows through a network is also the role of an intermediary device including:
 - Regenerate and retransmit data signals.
 - Maintain information about what pathways exist through the network and internetwork.
 - Notify other devices of errors and communication failures.
- 2. Network Media

Communication across a network is carried through a medium which allows a message to travel from source to destination.

- Networks typically use three types of media:
 - Metallic wires within cables, such as copper
 - Glass, such as fiber optic cables
 - Wireless transmission

3.6 Network Representations

Network diagrams, often called topology diagrams, use symbols to represent devices within the network. In addition to the device representations on the right, it is important to remember and understand the following terms:

- Network Interface Card (NIC)
- Physical Port
- Interface

3.7 Network Classification

The following list presents categories used for classifying networks.

a. Connection method

Computer networks can be classified according to the hardware and software technology that is used to interconnect the individual devices in the network, such as optical fiber, Ethernet, wireless LAN, HomePNA, powerline communication or G.hn.

Ethernet, as it is defined by IEEE 802, utilises various standards and mediums that enable communication between devices. Frequently

deployed devices include hubs, switches, bridges, or routers. Wireless LAN technology is designed to connect devices without wiring. These devices use radio waves or infrared signals as a transmission medium. *ITU-T G.hn* technology uses existing home wiring (coaxial cable, phone lines and power lines) to create a high-speed (up to 1 Gigabit/s) local area network.

b. Wired technologies

- i. **Twisted pair** wire is the most widely used medium for telecommunication. Twisted-pair cabling consists of copper wires that are twisted into pairs. Ordinary telephone wires consist of two insulated copper wires twisted into pairs. Computer networking cabling consists of 4 pairs of copper cabling that can be utilised for both voice and data transmission. The use of two wires twisted together helps to reduce crosstalk and electromagnetic induction. The transmission speed ranges from 2 million bits per second to 100 million bits per second. Twisted pair cabling comes in two forms which are Unshielded Twisted Pair (UTP) and Shielded twisted-pair (STP) which are rated in categories which are manufactured in different increments for various scenarios.
- ii. **Coaxial cable** is widely used for cable television systems, office buildings, and other work-sites for local area networks. The cables consist of copper or aluminum wire wrapped with an insulating layer typically of a flexible material with a high dielectric constant, all of which are surrounded by a conductive layer. The layers of insulation help to minimise interference and distortion. Transmission speeds range from 200 million to more than 500 million bits per second.

- iii. **Optical fiber cable** consists of one or more filaments of glass fiber wrapped in protective layers. It transmits light which can travel over extended distances. Fiber-optic cables are not affected by electromagnetic radiation. Transmission speed may reach trillion of bits per second. The transmission speed of fiber optics is hundred of times faster than for coaxial cables and thousand of times faster than a twisted-pair wire.

SELF-ASSESSMENT EXERCISE 1

Identify the key problems with wired media.

c. **Wireless technologies**

Takenoteofthefollowing.

- i. **Terrestrial microwave** –terrestrialmicrowavesuseEarth-based transmitterandreceiver.Theequipment lookssimilartosatellitedishes. Terrestrialmicrowavesuselow-gigahertz range,whichlimitsall communicationsto line-of-sight.Path between relay stations spaced approx,30milesapart.Microwaveantennasareusuallyplacedontopof buildings,towers,hills,andmountainpeaks.
- ii. **Communicationssatellites**– thesatellitesusemicrowaveradioastheir telecommunicationsmediumwhich are not deflectedby the Earth's atmosphere. Thesatellitesarestationed inspace,typically22,000miles (forgeosynchronous satellites)abovetheequator.TheseEarth-orbiting systems are capable of receivingand relaying voice, data, and TV signals.
- iii. **Cellular andPCSystems**–useseveralradiocommunications technologies.Thesystemsaredividedtodifferentgeographicareas. Eachareahasalow-power transmitter orradiorelayantennadeviceto relaycallsfromoneareato thenextarea.
- iv. **WirelessLANs**–wirelesslocalareanetworkuseahigh-frequency radiotechnologysimilartodigitalcellularandalow-frequency radio technology.WirelessLANsusespreadspectrum technologytoenable communicationbetweenmultipledevicesinalimitedarea.Anexample of open-standardswirelessradio-wavetechnologyisIEEE.

- v. **Infrared communication**- this can transmit signals between devices within small distances not more than 10 meters peer to peer or (face to face) without anybody in the line of transmitting.

3.8 Types of Networks Based on Physical Scope

Common types of computer networks may be identified by their scale.

a. Local area network

A local area network (LAN) is a network that connects computers and devices in a limited geographical area such as home, school, computer laboratory, office building, or closely positioned group of buildings.

Each computer or device on the network is a node. Current wired LANs are most likely to be based on Ethernet technology, although new standards like *ITU-TG.hn* also provide a way to create a wired LAN using existing home wires (coaxial cables, phone lines and power lines).

Typically a library network, in a branching tree topology and controlled access to resources. All interconnected devices must understand the network layer (layer 3), because they are handling multiple subnets (the different colors).

Those inside the library, which have only 10/100 Mbit/s Ethernet connections to the user device and a Gigabit Ethernet connection to the central router, could be called "layer 3 switches" because they only have Ethernet interfaces and must understand IP. It would be more correct to call them access routers, where the router at the top is a distribution router that connects to the Internet and academic networks' customer access routers. The defining characteristics of LANs, in contrast to

WANs (Wide Area Networks), include their higher data transfer rates, smaller geographic range, and no need for leased telecommunication lines. Current Ethernet or other IEEE 802.3 LAN technologies operate at speeds up to 10 Gbit/s. This is the data transfer rate. IEEE has projects investigating the standardization of 40 and 100 Gbit/s.

b. Personal area network

A Personal Area Network (PAN) is a computer network used for communication among computer and different information technological devices close to one person. Some examples of devices that are used in a PAN are personal computers, printers, fax machines, telephones, PDAs, scanners, and even video game consoles. A PAN may include wired and wireless devices. The reach of a PAN typically extends to 10 meters. A wired PAN is usually constructed with USB and Firewire connections while technologies such as Bluetooth and infrared communication typically form a wireless PAN.

c. Home area network

A Home Area Network (HAN) is a residential LAN which is used for communication between digital devices typically deployed in the home, usually a small number of personal computers and accessories, such as printers and mobile computing devices. An important function is the sharing of Internet access, often a broadband service through a CATV or Digital Subscriber Line (DSL) provider. It can also be referred to as an office area network (OAN).

d. Wide area network

A Wide Area Network (WAN) is a computer network that covers a large geographic area such as a city, country, or spans even intercontinental distances, using a communications channel that combines many types of media such as telephonelines, cables, and airwaves. A WAN often uses

transmission facilities provided by common carriers, such as telephone companies. WAN technologies generally function at the lower three layers of the OSI reference model: the physical layer, the data link layer, and the network layer.

e. Campus network

A campus network is a computer network made up of an interconnection of local area networks (LAN's) within a limited geographical area. The networking equipments (switches, routers) and transmission media (optical fiber, copper plant, Cat5 cabling etc.) are almost entirely owned (by the campus tenant/owner: an enterprise, university, government etc.).

In the case of a university campus-based campus network, the network is likely to link a variety of campus buildings including; academic departments, the university library and student residence halls.

f. Metropolitan area network

A Metropolitan area network is a large computer network that usually spans a city or a large campus.

g. Enterprise private network

An enterprise private network is a network built by an enterprise to interconnect various company sites, e.g., production sites, head offices, remote offices, shops, in order to share computer resources.

h. Virtual private network

A Virtual Private Network (VPN) is a computer network in which some of the links between nodes are carried by open connections or virtual circuits in some larger network (e.g., the Internet) instead of by physical wires. The data link layer protocols of the virtual network are said to be tunneled through the larger network when this is the case. One common application is secured communications through the public Internet; but a VPN need not have explicit security features, such as authentication or content encryption. VPNs, for example, can be used to separate the traffic of different user communities over an underlying network with strong security features.

VPN may have best-effort performance, or may have a defined service level agreement (SLA) between the VPN customer and the VPN service provider. Generally, a VPN has a topology more complex than point-to-point.

i. Internet network

An internet network is the connection of two or more private computer networks via a common routing technology (OSI Layer 3) using routers. The Internet is an aggregation of many internet networks; hence its name has been shortened to Internet.

j. Backbone network

A Back-Bone Network (BBN)-or network backbone, is part of a computer network infrastructure that interconnects various pieces of network, providing a path for the exchange of information between different LANs or subnetworks. A backbone can tie together diverse networks in the same building, in different buildings in a campus environment,

over wide areas. Normally, the backbone's capacity is greater than the networks connected to it.

A large corporation that has many locations may have a backbone network that ties all of the locations together, for example, if a server cluster needs to be accessed by different departments of a company that are located at different geographical locations. The pieces of the network connections (for example: Ethernet, wireless) that bring these departments together is often mentioned as network backbone. Network congestion is often taken into consideration while designing backbones. Backbone networks should not be confused with the internet backbone.

k. Global area network

A Global Area Network (GAN) is a network used for supporting mobile communications across an arbitrary number of wireless LANs, satellite coverage areas, etc. The key challenge in mobile communications is handing off the user communications from one local coverage area to the next. In IEEE Project 802, this involves a succession of terrestrial wireless LANs.

l. Internet

The internet is a global system of interconnected governmental, academic, corporate, public, and private computer networks. It is based on the networking technologies of the Internet Protocol Suite. It is the successor of the Advanced Research Projects Agency Network (ARPANET) developed by DARPA of the United States Department of Defense. The Internet is also the communications backbone underlying the World Wide Web (WWW).

Participants in the Internet use a diverse array of methods of several hundred documented, and often standardised, protocols compatible with the Internet Protocol Suite and an addressing system (IP addresses) administered by the Internet Assigned Numbers Authority and address registries. Service providers and large enterprises exchange information about their reachability of their address space through the Border Gateway Protocol (BGP), forming a redundant world-wide mesh of transmission paths.

m. Intranets and extranets

Intranets and extranets are parts or extensions of a computer network, usually a local area network.

An intranet is a set of networks, using the Internet Protocol and IP-based tools, such as web browsers and file transfer applications, which are under the control of a single administrative entity. That administrative entity closes the intranet to all but specific, authorised users. Most commonly, an intranet is the internal network of an organisation. A large intranet will typically have at least one web server to provide users with organisational information.

An extranet is a network that is limited in scope to a single organisation or entity and also has limited connections to the other networks. A company's customers may be given access to some part of its intranet while at the same time the customers may not be considered trusted from a security standpoint. Technically, an extranet may also be categorised as a *CAN*, *MAN*, *WAN*, or other type of network, although an extranet cannot consist of a single *LAN*; it must have at least one connection with an external network.

n. Overlay network

An overlay network is a virtual computer network that is built on top of another network. Nodes in the overlay are connected by virtual or logical links, each of which corresponds to a path, perhaps through many physical links, in the underlying network.

A sample overlay network: *IP over SONET over Optical*

For example, many peer-to-peer networks are overlay networks because they are organised as nodes of a virtual system of links run on top of the Internet. The Internet was initially built as an overlay on the telephone network. Overlay networks have been around since the invention of networking when computer systems were connected over telephone lines using modems, before any data network existed.

Nowadays the Internet is the basis for many overlaid networks that can be constructed to permit routing of messages to destinations specified by an IP address. For example, distributed hash tables can be used to route messages to a node having a specific logical address, whose IP address is known in advance.

Overlay networks have also been proposed as a way to improve Internet routing, such as through quality of service guarantees to achieve higher-quality streaming media. Previous proposals such as *IntServ*, *DiffServ*, and *IP Multicast* have not seen wide acceptance largely because they require modification of all routers in the network. On the other hand, an overlay network can be incrementally deployed on end-hosts running the overlay protocol software, without cooperation from Internet service providers. The overlay has no control over how packets are routed in the

underlying network between two overlay nodes, but it can control, for example, the sequence of overlay nodes a message traverses before reaching its destination.

SELF-ASSESSMENT EXERCISE 2

What is the relevance of *ISP* in internet connections?

3.9 Introduction to Telecommunication Circuit

A **telecommunication circuit** is any line, conductor, or other conduit by which information is transmitted. A **dedicated circuit, private circuit, or leased line** is a line that is dedicated to only one use. Originally, this was an analog, and was often used by radio stations as a studio/transmitter link (STL) or remote pickup unit (RPU) for their audio, sometimes as a backup to other means. Later lines were digital, and used for private corporate data networks.

A **telecommunication circuit** may be defined as follows.

- The complete path between two terminals over which one-way or two-way communications may be provided. See communications protocol.
- An electronic path between two or more points, capable of providing a number of channels.
- A number of conductors connected together for the purpose of carrying an electric current.
- An electronic closed-loop path among two or more points used for signal transfer.

- A number of electrical components, such as resistors, inductances, capacitors, transistors, and power sources connected together in one or more closed loops.

4.0 CONCLUSION

In this unit, you have been introduced to the fundamentals of Networking. You have also learnt the different types of Networking and its classification. You are advised to go over this unit again since it is the basis for understanding the course.

5.0 SUMMARY

In this unit, you have learnt that:

- computer network, often simply referred to as a network, is a collection of computers and devices interconnected by communications channels that facilitate communications among users and allow users to share resources.
- a telecommunication circuit is any line, conductor, or other conduit by which information is transmitted.
- types of telecommunication circuit are dedicated circuit, private circuit, or leased line

6.0 TUTOR-MARKED ASSIGNMENT

1. Briefly describe the following:- i.

Internet network

ii. Backbone Network

iii. Intranet

2. State and explain all the types of networks based on their physical scope.

7.0 REFERENCES/FURTHER READING

Hafner, Katie. (1998). *Where Wizards Stay Up Late: The Origins of the Internet*.

CCNA Routing and Switching. "Introduction to Networks" www.netacad.com

["How Does the Internet Work?"](http://tldp.org/HOWTO/Unix-and-Internet-Fundamentals-HOWTO/internet.html) <http://tldp.org/HOWTO/Unix-and-Internet-Fundamentals-HOWTO/internet.html>. Retrieved 15 June 2009.

Johna, Till, Johnson. "Net was born of economic necessity, not fear".
<http://www.networkworld.com/columnists/2004/0607johnson.html>. Retrieved June 15, 2009. URL is sufficient attribution.

Leonard, Kleinrock (2005). *The History of the Internet*.
http://www.lk.cs.ucla.edu/personal_history.html. Retrieved 2009-05-28.

UNIT2 DATALINKS

CONTENTS

- 1.0 Introduction
- 2.0 Objectives
- 3.0 MainContent
 - 3.1 Introductionto OpenSystemInterconnection(OSI)Model
 - 3.2 DataLinkLayer
 - 3.3 CategoriesofDataLink
 - 3.4 TypesofDataLink
- 4.0 Conclusion
- 5.0 Summary
- 6.0 Tutor-MarkedAssignment
- 7.0 References/FurtherReading

1.0 INTRODUCTION

In the previous unit, you have been introduced to the basic concept of networking. In this unit, you will learn about another concept of networking- which is data link. In telecommunication, a **data link** is the means of connecting one location to another for the purpose of transmitting and receiving digital information. It can also refer to a set of electronics assembly, consisting of a transmitter and a receiver [two Data Terminal Equipments (DTEs)] and the interconnecting data telecommunication circuit. These are governed by a link protocol enabling digital data to be transferred from a data source to a data sink.

2.0 OBJECTIVES

At the end of this unit, you should be able to:

- discuss OSI model
- explain the term “datalink”
- identify the various types of datalink
- state the relevance of datalink to networking.

3.0 MAIN CONTENT

3.1 Introduction to Open System Interconnection (OSI) Model

OSI is not a physical model, though, it is a set of guidelines that application developers can use to create and implement applications that run on a network. It also provides a framework for creating and

implementing networking standards, devices and internetworking schemes.

OSI has seven different layers, divided into two groups. The top three layers define how the applications within the end stations will communicate with each other and with other users. The bottom four layers define how data is transmitted end to end.

Table 2.1 below shows the layers of the OSI reference model

Table 2.1: OSI Reference Model

	Layers	Summary of Functions
7	Application layer	Interacts with the operating systems and applications
6	Presentation layer	Transformation of data. Responsible for functions like data compression, encryption, etc.
5	Session layer	Defines a connection from a user to a network server or from a peer on a network to another peer
4	Transport layer	Manages the flow of information from one network node to another. Ensures that packets are decoded in the proper sequence and
3	Network layer	Defines how data move from one point to another on a network. It determines what goes into each packet and defines different packet
2	Data link layer	Defines standards that assign meaning to bits carried by the physical layer
1	Physical layer	Defines the properties of the physical medium used to make network connection

The OSI model is fully discussed in other courses in computer network. However, data link layer is expected to be fully discussed in this course.

3.2 DataLinkLayer

The datalink layer is the second layer in the Open System Interconnection (OSI) reference model. The datalink layer defines the rules for accessing and using the physical layer. The datalink layer provides the physical transmission of the data and handles error notification, network topology and flow control. This means that the datalink layer will ensure that messages are delivered to the proper device on a LAN using hardware addresses and will translate messages from the Network layer into bits for the physical layer to transmit.

The datalink layer in the model combines packets and bytes into frames, it also provides access to media using MAC address. The datalink layer does perform the error detection but it should be noted that it does not perform the error correction process.

The data layer formats the message into pieces, each called a data frame, and adds a customised header containing the hardware destination and source address.

3.3 Categories of Data Link

There are at least three categories of basic data-link configurations that can be conceived of and used:

- simplex communications - referring to all communications in one direction only e.g. radio transmission.
- half-duplex communications - referring to communications in both directions, but not both ways simultaneously, e.g. walkie-talkie.

- duplex communications, communications in both directions simultaneously,e.g. telephoneconversation.

3.4 Types of Data Links

Let us consider these one after the other a.

Industrial Ethernet

Industrial ethernet is used to provide data link solutions for the communications and automation industry. Traditional office grade ethernet cannot meet the reliability demanded by industrial applications. A brief loss of service in an office environment may not be such a big issue, but in an industrial environment it may represent significant loss on your capital investment.

Industrial ethernet is specifically designed to operate in harsh environments such as factory floor automation, process control, HVAC, medical, manufacturing. Typically, industrial automation devices include, rugged case, din rail attachment, wide temperature specification, broad power source input, these features give us a reliable ethernet connection in demanding environments.

b. Radiomodems

Radiomodems are radio frequency transceivers for serial data communications. They connect to serial ports RS232, RS422/485 and transmit and receive signals from other matching radio (point to point) or radio (multidrop) network.

Wireless radiomodems are designed to be transparent to the systems within which they operate. All communication appears to your system as if communicating across directly connected cables; no special preparation of your data is needed. MaxStream units provide you true plug-and-communicate wireless capability operating in the internationally recognised 2.4GHz license free band.

c. Ethernet to RS232, RS485 serial devices servers

An Ethernet to RS232 or RS485 device server allows your network to enable virtually any serial RS232/422/485 port device. They provide the ability to remotely monitor, control or diagnose your equipment over your LAN or even WAN (internet/web) link. Allowing you to maintain the existing investment you have made in serial interface plant and machine equipment.

d. Wireless RS232 link

When creating an RS-232 wireless link, you can replace conventional expensive RS232 serial cable runs, allowing for an easy to use, invisible connection. Handywave bluetooth is the cable replacement solution for RS-232. Simply plug one unit into your RS-232 device and the other into your PC for an instant wireless link with minimal setup and also gives the added flexibility and mobility not available with traditional wired RS232 links.

e. **GSM and GPRS**

A GSM modem is a wireless modem that works with a GSM wireless network. A wireless modem behaves like a dial-up modem. The main difference between them is that a dial-up modem sends and receives data through a fixed telephone line while a wireless modem sends and receives data through radio waves. GPRS modem is a GSM modem which, additionally, supports the GPRS technology for data transmission. GPRS stands for General Packet Radio Service. It is a packet-switched technology that is an extension of GSM. (GSM is a circuit-switched technology.) A key advantage of GPRS over GSM is that GPRS has a higher data transmission speed. This technology is ideal for M2M (machine-to-machine communications) applications such as meter reading, remote maintenance, traffic control systems, vending machines and building management systems HVAC.

f. **Power Over Ethernet (POE)**

Power Over Ethernet or POE technology describes any system to transmit electrical power, along with data, to remote devices over standard twisted-pair cable in an Ethernet network. The standard is IEEE 802.3af which calls for 48 Volts DC over two pairs of a four-pair cable at a maximum current of 350 AMP for a maximum load power of 16 Watts.

g. **Outdoor Ethernet switches**

An outdoor ethernet switch is specifically designed for the toughest industrial environments. An outdoor switch is constructed from a rugged weather-tight aluminum case and the design usually carries an IP rating which provides a waterproof and dust-tight connection. An outdoor Ethernet switch can be easily adopted in almost all kinds of industrial applications and provides the most reliable solutions for your network in outdoor environments, typical applications include: railway, moving vehicles, factory automation, and marine (DNV Approval).

3.5 Function of Data-link Layer

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

i. 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 end, data link layer picks up signals from hardware and assembles them into frames.

ii. 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.

iii. 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.

iv. 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.

v. 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.

3.6 Error Detection

Besides framing, datalink layers also include mechanisms to detect and sometimes even recover from transmission error. To allow a receiver to detect transmission errors, a sender must add some redundant information as an error detection code to the frame sent. This error detection code is computed by the sender on the frame that it transmits.

When the receiver receives a frame with an error detection code, it recomputes it and verifies whether the received error detection code matches the computer error detection code. If they match, the frame is considered to be valid.

Many error detection schemes exist and entire books have been written on the subject. A detailed discussion of these techniques is outside the scope of this book, and we will only discuss some examples to illustrate the key principles.

To understand error detection codes, let us consider two devices that exchange bit strings containing N bits. To allow the receiver to detect a transmission error, the sender converts each string of N bits into a string of $N+r$ bits. Usually, the r redundant bits are added at the beginning or the end of the transmitted bit string, but some techniques interleave redundant bits with the original bits. An error detection code can be defined as a function that computes the r redundant bits corresponding to each string of N bits. The simplest error detection code is the parity bit. There are two types of parity schemes : even and odd parity. With the even (resp. odd) parity scheme,

the redundant bit is chosen so that an even (resp. odd) number of bits are set to 1 in the transmitted bit string of $N+r$ bits. The receiver can easily recompute the parity of each received bit string and discard the strings with an invalid parity. The parity scheme is often used when 7-bit characters are exchanged. In this case, the eighth bit is

often a parity bit. The table 3.6a shows the parity bits that are computed for bit strings containing three bits.

Table 3.6a: Odd and Even Parity

3 bits string	Odd parity	Even parity
000	1	0
001	0	1
010	0	1
100	0	1
111	0	1
110	1	0
101	1	0
011	1	0

The parity bit allows a receiver to detect transmission errors that have affected a single bit among the transmitted $N+r$ bits. If there are two or more bits in error, the receiver may not necessarily be able to detect the transmission error. More powerful error detection schemes have been defined. The Cyclical Redundancy Checks (CRC) are widely used in datalink layer protocols. An N -bits CRC can detect all transmission errors affecting a burst of less than N bits in the transmitted frame and all transmission errors that affect an odd number of bits.

It is also possible to design a code that allows the receiver to correct transmission errors. The simplest error correction code is the triple modular redundancy (TMR). To transmit a bit set to 1 (resp. 0), the sender transmits 111 (resp. 000). When there are no transmission errors, the receiver can decode 111 as 1. If transmission errors have affected a single bit, the receiver performs majority voting as shown in the table below. This scheme allows the receiver to correct all transmission errors that affect a single bit.

It is also possible to design a code that allows the receiver to correct transmission errors. The simplest error correction code is the triple modular redundancy (TMR). To transmit a bit set to 1 (resp. 0), the sender transmits 111 (resp. 000). When there are no transmission errors, the receiver can decode 111 as 1. If transmission errors have affected a single bit, the receiver performs majority voting as shown in the table below. This scheme allows the receiver to correct all transmission errors that affect a single bit.

Table 3.6b: Error Correction

3 bits string	Odd parity
000	0
001	0
010	0
100	0
111	1
110	1
101	1
011	1

Other more powerful error correction codes have been proposed and are used in some applications. The Hamming Code is a clever combination of parity bits that provides error detection and correction capabilities.

4.0 CONCLUSION

In this unit, you have been introduced to what is referred to as data link. You have also learnt the different types of data link in use and as well as the relevance of data link to networking.

5.0 SUMMARY

In this unit, you have learnt about:

- OSI model - which has 7 layers: application, presentation, session, transport, network, data link and physical.
- data link layer which defines rules for accessing and using physical layer
- 7 types of data link which include industrial Ethernet, radio modems, Ethernet to RS232, wireless RS232 link, GSM and GPRS, Power Over Ethernet (POE) and outdoor Ethernet
- Error -detection

6.0 TUTOR-MARKED ASSIGNMENT

1. What do you understand by the term “datalink”?
2. How relevant is datalink to networking of computers?
3. State some protocols that are defined by the datalink layer.

7.0 REFERENCES/FURTHER READING

Bates, R., & Gregory, D. (2006). Voice & Data Communications Handbook, Fifth Edition: McGraw-hill.

Mike Meyers (2007). (6th Ed.). All in One CompTIA A+ Certification Exam Guide. McGraw Hill.

Olivier Bonaventure (2011), Computer Networking: Principles, Protocols and Practice Release 0.25. Saylor URL: <http://www.saylor.org/courses/cs402/>

<http://www.apan.net/meetings/busan03/cs-history.htm>. Retrieved December 25, 2005.

["APRICOT webpage"](http://www.apricot.net/). Apricot.net. 2009-05-04. <http://www.apricot.net/>. Retrieved 2009-05-28.

"A Brief History of the Internet in China". *China Celebrates 10 Years of Being Connected to the Internet*.

<http://www.pcworld.idg.com.au/index.php/id;854351844;pp;2;fp;2;fpid;>

1. Retrieved December 25, 2005.

UNIT3 DEPLOYING PHYSICAL MEDIA

CONTENTS

- 1.0 Introduction
- 2.0 Objectives
- 3.0 Main Content
 - 3.1 Introduction to Physical Media
 - 3.2 Physical Media Comparison
- 4.0 Conclusion
- 5.0 Summary
- 6.0 Tutor-Marked Assignment
- 7.0 References/Further Reading

1.0 INTRODUCTION

In the previous units, you have been introduced to network and data link layer of the OSI reference model. This unit will introduce you to physical media used in the networking of computer and it will, as well, teach you how to deploy the physical media.

2.0 OBJECTIVES

At the end of this unit, you should be able to:

- explain the term physical media
- identify the various types of physical media
- identify which physical media is good for each of the network types.

3.0 MAIN CONTENT

3.1 Introduction to Physical Media

In the OSI reference model, any physical means for transmitting data is referred to as the physical media. The bottom of the OSI model's physical layer provides an interface to such media. Specifications for the physical media themselves are not part of the OSI model.

3.2 Types of Physical Media

Let us look at these one after another.

- a. **Twisted pair**-wire twisted to avoid crosstalk interference. It may be shielded or unshielded.

- UTP-Unshielded Twisted Pair. Normally UTP contains 8 wires or 4 pair; 100 meter maximum length; 4-100 Mbps speed.
- STP-Shielded Twisted Pair. 100 meter maximum length; 16-155 Mbps speed; lower electrical interference than UTP.

b. Coaxial-two conductors separated by insulation such as TV75 ohm cable; maximum length of 185 to 500 meters; It is of two types, namely:

thinnet-thinnet uses a British Naval Connector (BNC) on each end.

Thinnet is part of the RG-58 family of cable. Maximum cable length is 185 meters. Transmission speed is 10 Mbps. *Thinnet* cables should have 50 ohm impedance; and its terminator has 50 ohm impedance; barrel connector will have no impedance. Maximum *thinnet* nodes are 30 on a segment. One end of each cable is grounded.

thicknet-half inch rigid cable; maximum cable length is 500 meters. Transmission speed is 10 Mbps; it is expensive and is not commonly used- (RG-11 or RG-8). A vampire tap or piercing tap is used with a transceiver attached to connect computers to the cable. 100 connections may be made. The computer has an Attachment Unit Interface (AUI) on its network card which is a 15 pin DB-15 connector. The computer is connected to the transceiver at the cable from its AUI on its network card using a drop cable. Maximum *thicknet* nodes are 100 on a segment. One end of each cable is grounded.

The RG value for cable types refers to its size. Coax cable types are listed below.

- RG-58/U - 50 ohm, with a solid copper wire core for thin ethernet.
- RG-58 A/U-50ohm, with a stranded wire core.
- RG-58C/U-Military version of RG-58 A/U.
- RG-59-75ohm, for broadband transmissions such as cable TV.
- RG-62-93ohm, primarily used for ArcNet.
- RG-6-used for satellite cable (if you want to run a cable to a satellite).
- RG-8-50ohm thick ethernet.
- RG-11-75ohm thick ethernet.

c. Fiber-optic-data is transmitted using light rather than electrons.

Usually there are two fibers, one for each direction. Cable length of 2 Kilometers; speed from 100Mbps to 2Gbps. This is the most expensive and most difficult to install, but is not subject to interference. There are two types of cables:

- single mode cables for use with lasers have greater bandwidth and cost more. Injection laser diodes (ILD) work with single mode cable.
- multimode cables for use with Light Emitting Diode (LED) drivers; all signals appear to arrive at the same time. P-N diodes or photodiodes are used to convert light to electric signals when using multimode.

Table 3.1: Types of Fiber Cable

Fiber thickness (microns)	Cladding thickness (microns)	Mode
8.3	125	single
62.5	125	multi
50	125	multi
100	140	multi

SELF-ASSESSMENT EXERCISE 1

How is the fiberoptic medium better than coaxial cable?

3.2 Physical Media Comparisons

Let us look at the features of different types of physical media in the table below.

Table 3.2: Comparing Physical Media

Media	Distance(meters)	Speed	Approx.cost/station
UTP	100	4-100Mbps	\$90
STP	100	16-155Mbps	\$125
Thinnet	185	10Mbps	\$25
Thicknet	500	10Mbps	\$50
Fiber	2000	100Mbps-2Gbps	\$250(multimode)

4.0 CONCLUSION

This unit has taken you through understanding the physical media and its various types. In this unit, we also compare the media using the distance they can cover; their effective speed of transferring data as our basis of comparison. The next unit will introduce you to network protocols. Protocol can be viewed as a language that has to be understood by all the computers on a network.

5.0 SUMMARY

In this unit, you have learnt that:

- a physical medium is any physical means for transmitting data
- the 3 types of physical media are twisted pair, coaxial cable and fiber optic.

6.0 TUTOR-MARKED ASSIGNMENT

1. Discuss the various types of physical media as regards the distance they can cover.
2. State the features of 8 types of coaxial cable.
3. Explain the two types of fiber optic cable.

7.0 REFERENCES/FURTHER READING

David Roessner, Barry Bozeman, Irwin Feller, Christopher Hill, Nils Newman (1997). *The Role of NSF's Support of Engineering in Enabling Technological Innovation*.
<http://www.sri.com/policy/csted/reports/techin/inter2.html>.
 Retrieved 2009-05-28.

Hauben, Ronda (2004). "The Internet: On its International Origins and Collaborative

Vision". *Amateur Computerist* **12**(2). <http://www.ais.org/~jrh/acn/A>
Cn12-2.a03.txt. Retrieved 2009-05-
29.

"RFC675-Specification of internet transmission control program".
Tools.ietf.org. <http://tools.ietf.org/html/rfc675>. Retrieved 2009-
05-28.

Tanenbaum, Andrew S. (1996). *Computer Networks*. Prentice Hall.

UNIT4 NETWORK PROTOCOLS

CONTENTS

- 1.0 Introduction
- 2.0 Objectives
- 3.0 Main Content
 - 3.1 Introduction to Network Protocols
 - 3.2 Network Protocol Overview
 - 3.3 Types of Network Protocols
- 4.0 Conclusion
- 5.0 Summary
- 6.0 Tutor-Marked Assignment
- 7.0 References/Further Reading

1.0 INTRODUCTION

In this unit, you will be taken through Network Protocol and its general concepts.

2.0 OBJECTIVES

At the end of this unit, you should be able to:

- explain the meaning of network protocols

□ statethetypesof networkprotocols.

3.0 MAINCONTENT

3.1 IntroductiontoNetworkProtocols

Aprotocolisasetofrules thatgovernsthecommunications between computersonanetwork. Theserulesincludeguidelinesthatregulatethe following characteristics of anetwork:access method,allowed physical topologies,types of cabling,and speed of data transfer.

Packets can be transmitted across networks or over telepho nelines. In fact, network protocols and several communications protocols use packet switching to establish a connection and route information. The format of a packet depends on the protocol that creates the packet the network computer must have a network protocol driver loaded. This program may be referred to as the transport protocol, or just as protocol. It operates between the adapter and the initial layer of network software to package and unpack data for the LAN.

3.2 Network Protocol Overview

The *OSI* model, and any other network communication model, provides only a conceptual framework for communication between computers, but the model itself does not provide specific methods of communication. Actual communication is defined by various communication protocols. In the context of data communication, a protocol is a formal set of rules, conventions and data structure that governs how computers and other network devices exchange information over a network. In other words, a protocol is a standard procedure and format that two data communication devices must understand, accept and use to be able to talk to each other.

In modern protocol design, protocols are "layered" according to the *OSI* 7 layer model or a similar layered model. Layering is a design principle which divides the protocol design into a number of smaller parts, each part accomplishing a particular sub-task and interacting with the other parts of the protocol only in a small number of well-defined ways. Layering allows the parts of a protocol to be designed and tested without a combinatorial explosion of cases, keeping each design relatively simple. Layering also permits familiar protocols to be adapted to unusual circumstances.

The header and/or trailer at each layer reflect the structure of the protocol. Detailed rules and procedures of a protocol or protocol group are often defined by a lengthy document. For example, *IETF* uses *RFCs*

(Request For Comments) to define protocols and updates to the protocols.

A wide variety of communication protocols exist. These protocols were defined by many different standard organisations throughout the world and by technology vendors over years of technology evolution and development. One of the most popular protocol suites is *TCP/IP*, which is the heart of Internet networking communications. The IP- Internet Protocol is responsible for exchanging information between routers, so that the routers can select the proper path for network traffic; while *TCP* is responsible for ensuring that data packets are transmitted across the network, reliably, and error free. *LAN* and *WAN* protocols are also critical protocols in network communications. The *LAN* protocol suite is for the physical and data link layers of communications over various *LAN* media, such as ethernet wires and wireless radiowaves. The *WAN* protocol suite is for the lowest three layers and defines communication over various wide-area media, such as fiber optic and copper cables.

Network communication has slowly evolved. Today's new technologies are based on the accumulation of technologies over the years- which

may be either still existing or obsolete. As a result of this, the protocols which define network communication are highly inter-related. Many protocols rely on others for operation. For example, many routing protocols use the network protocol to exchange information between routers.

In addition to standards for individual protocols in transmission, there are now also interface standards for different layers to talk to the ones above or below (usually operating system specific). For example, Winsock and Berkeley socket between layers 4 and 5; *NDIS* and *ODI* between layers 2 and 3.

The protocols for data communication cover all areas as defined in the *OSI* model. However, the *OSI* model is only loosely defined. A protocol may perform the functions of one or more of the *OSI* layers, which introduces complexity to understanding protocols relevant to the *OSI* 7 layer model. In real-world protocols, there is some argument as to where the distinctions between layers are drawn; there is no one black and white answer.

To develop a complete technology that is useful for the industry, very often a group of protocols is required in the same layer or across many different layers. Different protocols often describe different aspects of a single communication; taken together, these form a protocol suite. For example, Voice over IP (VOIP), a group of protocols developed by many vendors and standard organisations, has many protocols across the

4 top layers in the *OSI* model.

Protocols can be implemented either in hardware or software or a mixture of both. Typically, the lower layers are implemented in hardware, with the higher layers being implemented in software.

Protocols could be grouped into suites (or families, or stacks) by their technical functions, or origin of the protocol introduction, or both. A protocol may belong to one or multiple protocol suites, depending on how you categorise it. For example, the Gigabit Ethernet protocol *IEEE 802.3* is a LAN (Local Area Network) protocol and it can also be used in MAN (Metropolitan Area Network) communications.

Most recent protocols are designed by the *IETF* for internetworking communications and by the *IEEE* for local area networking and metropolitan area networking. The *ITU-T* contributes mostly to Wide Area Networking (WAN) and telecommunications protocols. *ISO* has its own suite of protocols for internetworking communications, which is mainly deployed in European countries.

3.4 Types of Network Protocols

The most common network protocols are listed below.

- TCP/IP
- FTP
- AppleTalk
- HTTP
- SHTTP
- DHCP

i. TCP/IP Protocols

The TCP/IP protocol suite establishes the technical foundation of the Internet. It is the conceptual model and set of communications protocols used in the Internet and similar computer networks.

The TCP/IP suite's core functions are addressing and routing (IP/IPv6 in the networking layer) and transportation control (TCP, UDP in the transport layer).

IP - Internet Protocol

The Internet Protocol is the principal communications protocol in the Internet protocol suite for relaying datagrams across network boundaries.

On a LAN, messages are sent between machines by supplying the six bytes unique identifier (the "MAC" address). In an SNA network, every machine has Logical Units with their own network addresses. DECNET, AppleTalk, and Novell IPX all have a scheme for assigning numbers to each local network and to each workstation attached to the network.

TCP - Transmission Control Protocol

TCP provides a reliable stream delivery and virtual connection service to applications through the use of sequenced acknowledgment with retransmission of packets when necessary. TCP provides stream data transfer, transportation reliability, efficient flow control, full-duplex operation, and multiplexing. Check the TCP section for more details

ii. FTP: File Transfer Protocol

File Transfer Protocol (FTP) enables file sharing between hosts.

FTP uses TCP to create a virtual connection for control information and then creates a separate TCP connection for data transfers. The control connection uses an image of the TELNET protocol to exchange commands and messages between hosts.

The key functions of FTP are:

1. to promote sharing of files (computer programs and/or data);
2. to encourage indirect or implicit (via programs) use of remote computers;
3. to shield a user from variations in file storage systems among hosts; and
4. to transfer data reliably and efficiently.

FTP, though usable directly by a user at a terminal, is designed mainly for use by programs. FTP control frames are TELNET exchanges and can contain TELNET commands and option negotiation. However, most FTP control frames are simple ASCII text and can be classified as FTP commands or FTP messages. FTP messages are responses to FTP commands and consist of a response code followed by explanatory text.

iii. HTTP: Hypertext Transfer

The Hypertext Transfer Protocol (HTTP) is an application-level protocol with the lightness and speed necessary for distributed, collaborative, hypermedia information systems. HTTP has been in use by the World-Wide Web global information initiative since 1990.

HTTP allows an open-ended set of methods to be used to indicate the purpose of a request. It builds on the discipline of reference provided by the Uniform Resource Identifier (URI), as a location (URL) or name (URN), for indicating the resource on

which a method is to be applied. Messages are passed in a format similar to that used by Internet Mail and the Multipurpose Internet Mail Extensions (MIME).

HTTP is also used as a generic protocol for communication between user agents and proxies/gateways to other Internet protocols, such as SMTP, NNTP, FTP, Gopher and WAIS, allowing basic hypermedia access to resources available from diverse applications and simplifying the implementation of user agents.

The HTTP protocol is a request/response protocol. A client sends a request to the server in the form of a request method, URI, and protocol version, followed by a MIME-like message containing request modifiers, client information, and possible body content over a connection with a server. The server responds with a status line, including the message's protocol version and a success or error code, followed by a MIME-like message containing server information, entity meta information, and possible entity body content.

iv. DHCP: Dynamic Host Configuration Protocol

Dynamic Host Configuration Protocol (DHCP) is a communications protocol enabling network administrators manage centrally and to automate the assignment of IP addresses in a network.

In an IP network, each device connecting to the Internet needs a unique IP address. DHCP lets a network administrator supervise and distribute IP addresses from a central point and automatically sends a new IP address when a computer is plugged into a different place in the network.

DHCP uses the concept of a "lease" or amount of time that a given IP address will be valid for a computer. The lease time can vary depending on how long a user is likely to require the Internet connection at a particular location. It's especially useful in education and other environments where users change frequently. Using very short leases, DHCP can dynamically reconfigure networks in which there are more computers than there are available IP addresses.

DHCP supports static addresses for computers containing Web servers that need a permanent IP address.

v. **S-HTTP: Secure Hypertext**

Secure HTTP (S-HTTP) is a secure message-oriented communications protocol designed for use in conjunction with HTTP.

S-HTTP is designed to coexist with HTTP's messaging model and to be easily integrated with HTTP applications.

Secure HTTP provides a variety of security mechanisms to HTTP clients and servers, providing the security service options appropriate to the wide range of potential end uses possible for the World-Wide Web (WWW). S-HTTP provides symmetric capabilities to both client and server (in that equal treatment is

given to both requests and replies, as well as for the preferences of both parties) while preserving the transaction model and implementation characteristics of HTTP.

Several cryptographic message format standards may be incorporated into S-HTTP clients and servers. S-HTTP supports interoperation among a variety of implementations and is compatible with HTTP. S-HTTP aware clients can communicate with S-HTTP oblivious servers and vice-versa, although such transactions obviously would not use S-HTTP security features.

vi. **AppleTalk**

AppleTalk is a routable protocol that is used primarily in Macintosh environments to connect multiple systems in a network environment.

- Phase 1 This was designed for small work group environments and therefore supports a much smaller number of nodes on the network. Phase 1 supports non-extended networks. Each network segment can be assigned only a single network number, and only one zone is allowed in non-extended network.
- Phase 2 This was designed for larger networks and supports more than 200 hosts on the network. Phase 2 supports extended networks,

thereby allowing one network segment to be assigned multiple network numbers and allowing for multiple zones on that network segment.

4.0 CONCLUSION

This unit has introduced basic rudiments of network protocols, under which you have learnt different types of network protocols which include *TCP/IP*, *HTTP*, *FTP* etc.

5.0 SUMMARY

In this unit, you have learnt about:

- network protocol, which is a set of rules that governs communication between computers on a network
- types of network protocols, namely-*FTP*, *DHCP*, *HTTP* and *TCP/IP*
- Comparison between different network protocols.

6.0 TUTOR-MARKED ASSIGNMENT

1. What is the significance of network protocol?
2. Differentiate between 5 types of network protocols.
3. State the main features of *HTTP* and *SMTP*

7.0 REFERENCES/FURTHER READING

OGC-00-33R Department of Commerce: Relationship with the Internet Corporation
for Assigned Names and Numbers. Government Accountability Office. 7
July 2000. p.5. <http://www.gao.gov/new.items/og00033r.pdf>.

"DDNNIC". IAB Recommended Policy on Distributing Internet Identifier
Assignment. <http://www.rfc-editor.org/rfc/rfc1174.txt>.
Retrieved December 26, 2005.

Olivier Bonaventure (2011), Computer Networking: Principles, Protocols and
Practice Release 0.25. Saylor URL: <http://www.saylor.org/courses/cs402/>

MODULE2 NETWORKDESIGN

Unit1	HarnessingWi-Fi for User Mobility
Unit2	Building Internet Networks Using <i>TCP/IP</i> and Routers
Unit3	Network Standards (<i>IEEE 802</i> Standards)
Unit4	Implementing Security Best Practices

UNIT1 HARNESSINGWI-FIFORUSERMOBILITY

CONTENTS

1.0	Introduction
2.0	Objectives
3.0	Main Content
3.1	Overview of <i>Wi-Fi</i> Technology
3.2	Introduction to <i>Wi-Fi</i> Technology
3.3	<i>WLAN</i> Performance Metrics
3.4	Testing Methodology for Wireless Networks
3.5	Test Setup and Sample Results
4.0	Conclusion
5.0	Summary
6.0	Tutor-Marked Assignment
7.0	References/Further Reading

1.0 INTRODUCTION

This unit will expose you to the concept of Wi-Fi technology and its features.

2.0 OBJECTIVES

At the end of this unit, you should be able to:

- define the concept of Wi-Fi technology
- explain the details of what brought about Wi-Fi technology.

3.0 MAIN CONTENT

3.1 Overview of Wi-Fi Technology

Testing techniques developed for wired devices and networks fall short when applied to the WLAN market. The inherent instability of the unwired medium—air—in which the wireless world operates and the constant mobility of the WLAN user make the 802.11 protocol an order of magnitude more complex than equivalent wired protocols. As a

result, the metrics used to benchmark wired protocols are only a starting point for the WLAN industry.

Differences between wired and wireless networks require metrics and methodology for performance benchmarking that address the intricacies of the 802.11 protocol. This unit addresses wireless-specific functions central to business-critical applications on a wi-fi infrastructure and includes typical results and interpretations.

3.2 Introduction to Wi-Fi Technology

The IEEE 802.11b standard is popularly known as “Wireless Fidelity” (or Wi-Fi for short). It has become widely popular for wireless LANs in office environments. Proponents of this technology consider it great competition to third-generation wireless networks, which also provide high data rate, mobile internet access. Wi-Fi can be used to provide broadband wireless internet access.

Access Points (APs) can be installed at various locations in the city. The APs are also called “hotspot”. All the APs in a city can be interconnected through an ATM-based backbone network. As the wireless device moves from one location to another, the mobile device is connected to the nearest AP.

3.3 WLAN Performance Metrics

Wi-fi protocols address differences between wired and wireless networks, and the implementation of the more advanced wireless protocol demands performance validation. Algorithms used in a network's clients and APs (and the capacity of these devices to process the algorithms) limit the network's performance. The objective of the validation process and test metrics is to identify critical test parameters and find the correct method of testing them.

Testing Ethernet network performance is essentially a measure of packet forwarding rate. In addition to packet forwarding measurements, WLANs must undergo tests related to the unstable physical layer and end-user mobility; including automatic data rate adaptation, roaming, verification of security, QoS and overlapping BSSs, as well as behavioral tests that measure performance under abnormal network conditions. The primary focus of the testing effort should be parameters that eventually affect network efficiency and operation.

Data rate adaptation-wired LANs support fixed data rates: 10/100/1000 Mb/s. Wireless networks support multiple data rates: 11/5.5/2/1 Mb/s for 802.11b; 54/48/36/24/18/12/9/6 Mb/s for 802.11a and 802.11g. The

critical difference is that WLANs support dynamic rate adaptation and can operate at multiple data rates automatically determined by the end point (AP or client).

In addition, because the 802.11 standard does not specify exact criteria for data rate adaptation, the algorithms can vary from device to device. The rate adaptation algorithm should be based on optimising throughput; that is, when the number of errors at a specific data rate increases to the point where throughput is severely affected, the device should drop to a lower data rate to recover the best throughput at that distance from the AP.

The challenge is how to measure this repeatedly, while creating a metric that can be set as the *golden standard*. To measure throughput at fixed points, many vendors often use an interference-free environment with a long, direct line-of-sight area where they can simulate data rate adaptation by wheeling a client up and down on a cart. This method, however, lacks accurate rate adaptation data and is less efficient than newer devices that offer controlled RF environments and accurate signal attenuation through test setup automation. While characterising *range* vs. *data* rate, the test should simultaneously characterise *range* vs. *throughput* and *range* vs. *packet error rate*.

Roaming—as a client moves out of range of one AP, it dissociates from the AP and must associate and authenticate with another. If the client predicts this roam will occur by noticing the drop in signal and

searching for an alternate AP before it is actually disassociated from the first, it can optimise the roam time and network disruption caused by the roam.

The client device makes the decision to roam based on its position relative to different APs and their signal strengths. The client might periodically analyse signal strength of the APs that surround it and decide which one to associate with if it needs to roam. Load-sharing protocols used by some WLAN network vendors depart from the traditional client-based decision process, orchestrating client devices to associate with specific APs and spreading the load evenly among APs and optimising the entire network throughput. In addition, the IEEE is advancing its work on roaming through better RF measurement (802.11k) and fast roaming processes (802.11r). As the roaming process increases in complexity, it is critical to have standard roaming metrics for testing WLAN networks and equipment.

Roaming is critical because it takes time, which can cause data loss that can ultimately disrupt a communication session. Data loss is particularly important for time-sensitive enterprise applications, such as VoWLAN,

that are especially susceptible to packet delay caused by roaming. Roaming metrics include *roaming time*, *packet loss* and *session continuity*. Roaming time can be broken down into the following stages - scanning, associating to the new AP, authentication with the new AP and data flow. Analysing the time of each phase of this process will help ensure the most efficient roam time.

Packet forwarding - forwarding rate is a function of a device: in wired networks, the ethernet switch; in WLANs, the AP. Packet forwarding rate testing is always done at the highest signal strength and at the highest data rate because this puts the most demand on the device and measures its packet processing power in the most extreme case.

Like wired throughput tests, a wireless packet forwarding test varies the packet size to ensure the ability of the device to work with diverse traffic; but unlike wired devices, there are other factors to consider. The most critical is *security*, because wireless network devices must encrypt each packet. This additional overhead must be added to evaluate its effect on the packet forwarding rate. Another important factor is client capacity. Running the test with a large number of users stresses the AP's ability to handle a large number of users, each sending a portion of the bandwidth. This also affects how the AP functions under such conditions.

Security-this is a critical consideration for enterprise networks; because they are susceptible to intruders, wireless networks have more stringent security requirements than their wired counterparts. Wireless security protocols (802.11i) rely heavily on authentication and encryption, which depend on the processing power of the AP and client, and cryptography accelerators for data encryption. The efficiency with which the devices handle key management and encryption will have an effect on performance measurements, such as forwarding rate and roaming.

When a client initially accesses the network or roams between APs, authentication occurs using protocols such as *EAP-TLS*, *EAP-TTLS* and *LEAP*. Complex key derivation algorithms can overload APs if multiple simultaneous authentication requests are made. Authentication of wireless networks is tested by measuring how efficiently and quickly an AP manages simultaneous authentication requests.

Encryption protocols used in Wi-Fi, such as *WEP*, *TKIP* and *AES/CCMP*, can also impact throughput performance. This security metric is performed by making a series of comparative throughput measurements using different encryption methods.

Quality of Service (QoS)-because 802.11 is a shared media protocol

without QoS, WLANs cannot prioritise real-time applications such as voice and video, over data applications. QoS protocols for WLANs must account for jitter, delay and packet loss, which have required minimums for real-time applications including VOIP and multimedia streaming. Jitter, or inter-packet delay, is particularly critical in packetised voice.

3.4 Testing Methodology for Wireless Networks

Traditionally, wireless system designers have had a variety of testing options. Most are home-grown or custom-built, and include isolated screen rooms for RF control, large open spaces for testing mobility and expensive off-the-shelf meters that focus on point-to-point tests of the physical layer. Another approach is emerging for integrated chassis-based Wi-Fi testing. The older methods are costly, and in many cases do not provide the systems level configuration needed to accurately and usefully provide relevant information about the metrics discussed above.

Isolated screen rooms for controlling RF interference—WLAN system developers reduce the effects of RF interference by conducting tests in a large screen room that isolates devices under test from extraneous RF interference. It is the wireless equivalent of a "clean room." Although screen rooms can eliminate interference effects, they cannot test real-world network conditions such as mobility and roaming. Screen rooms can be expensive to erect and maintain; and are not portable, which limits their use and effectiveness.

SELF-ASSESSMENT EXERCISE

Describe the performance of Wi-Fi technology.

3.5 Test Setup and Sample Results

The roaming test is an example that shows the efficiency that can be achieved using a chassis-based test platform. Using two test modules—a WLA and an RFM—a phone or a PC client is connected between the two APs through two 80 dB programmable attenuators. The attenuators are programmed to force the client to roam in a controlled way from one AP to another. Data collection is performed on the source and destination channels simultaneously by the integrated dual-channel WLA module. The roaming test is fully automated and can be configured to repeat the measurements for a set period.

One attenuator is initially set to minimum and the other to maximum so that the client receives a strong signal from AP1 and associates with it; while AP2 is out of the client's range. The attenuator between AP1 and the client is then gradually increased, eventually making AP1 invisible to the client while the attenuator between the client and AP2 is gradually

decreased, "moving" AP2 within range and forcing a roam. The ranges and the rate of change of attenuators are configurable within the test script. At the end of the roaming test, the test script tabulates the details of each roam.

This automated roaming test implementation provides accurate time measurements and identifies specific time intervals in a way that emulates real-life roaming to the clients and access points: gradual signal strength decrease and increase. The entire test can be repeated as many times as is needed, and multiple roams can be performed in a short period without human involvement.

4.0 CONCLUSION

This unit introduces you to Wi-Fi technology and also teaches the advantage of the technology over wired technology. This unit has also exposed you to the testing methodology for wireless networks.

5.0 SUMMARY

In this unit, you have learnt that:

- The *IEEE 802.11b* standard is popularly known as 'Wireless Fidelity' (or Wi-Fi).

- Wi-fi protocols address the differences between wired and wireless networks
- Wi-fi can be used to provide broadband wireless internet access.
- *WLAN* must be tested continually; the primary focus of the testing effort should be parameters that eventually affect network efficiency and operation.
- Wi-fi technology has become popular for wireless *LAN*s in office environment

6.0 TUTOR-MARKED ASSIGNMENT

1. What do you understand by the term “wi-fi?”
2. In a tabular form, compare and contrast wi-fi and the *WLAN*.
3. State and explain the metrics for testing *WLAN*

7.0 REFERENCES/FURTHER READING

“NASA Successfully Tests First Deep Space Internet”. NASA media
release 08-298, November 18, 2008 Archived.

Prasad, K. V. (2009). *Principles of Digital Communication Systems and
Computer Networks*. Dreamtech Press.

<http://www.azimuthsystems.com/>

UNIT2 BUILDING INTERNETWORKS USING TRANSPORT CONTROL PROTOCOL (TCP)/INTERNET PROTOCOL (IP) AND ROUTERS

CONTENTS

- 1.0 Introduction
- 2.0 Objectives
- 3.0 Main Content
 - 3.1 TCP Technology
 - 3.1.1 Internet Protocol
 - 3.2 Routing in IP Environments
- 4.0 Conclusion
- 5.0 Summary
- 6.0 Tutor-Marked Assignment
- 7.0 References/Further Reading

1.0 INTRODUCTION

The language of the internet is Transport Control Protocol/Internet Protocol (TCP/IP). No matter what type of computer platform or software is being used, the information must move across the internet in this format. This protocol calls for data to be grouped together, in bundles, called network packets.

2.0 OBJECTIVES

At the end of this unit, you should be able to:

- explain the fundamentals of the *TCP/IP*
- state the relevance of *TCP/IP* in internet networks.

3.0 MAIN CONTENT

3.1 TCP Technology

TCP is a connection-oriented transport protocol that sends data as an unstructured stream of bytes. By using sequence numbers and acknowledgment messages, *TCP* can provide a sending node with delivery information about packets transmitted to a destination node. Where data has been lost in transit from source to destination, *TCP* can retransmit the data until either a timeout condition is reached or until successful delivery has been achieved. *TCP* can also recognise duplicate messages and will discard them appropriately. If the sending computer is

transmitting too fast for the receiving computer, *TCP* can employ flow control mechanisms to slow data transfer. *TCP* can also communicate delivery information to the upper-layer protocols and applications it supports. All these characteristics make *TCP* an end-to-end reliable transport protocol. *TCP* is specified in *RFC793*.

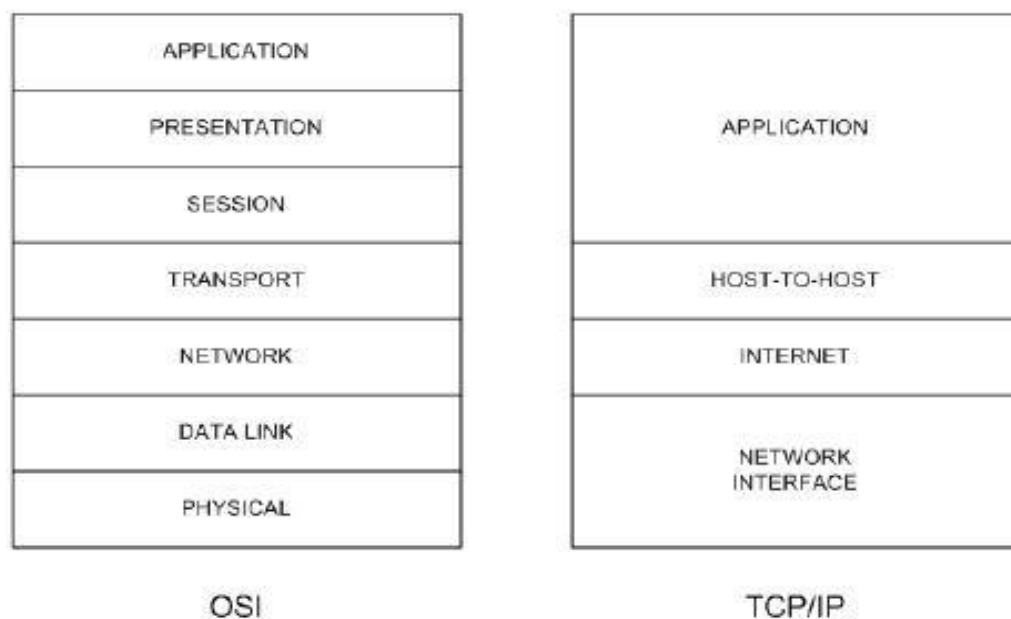


Figure 2.1: TCP/IP Protocol Suite in Relation to the OSI Reference Model

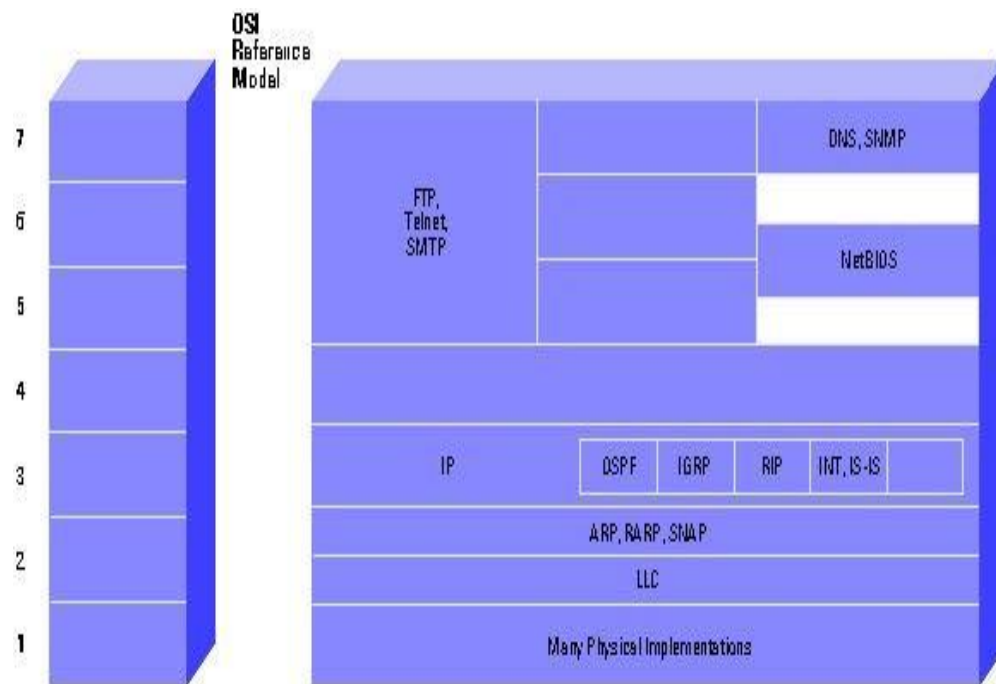


Figure 2.2: Important Internet Protocols in Relation to the OSI Reference Model

3.1.1 Internet Protocol (IP)

IP is the primary layer 3 protocol in the internet suite. In addition to internet network routing, *IP* provides error reporting and fragmentation and re-assembly of information units called datagrams for transmission over networks with different maximum data unit sizes. *IP* represents the heart of the internet protocol suite (**note:** the term *IP* in this section refers to *IPv4* unless otherwise stated explicitly).

IP addresses are globally unique, 32-bit numbers assigned by the network information center. Globally unique addresses permit *IP* networks anywhere in the world to communicate with each other. An *IP* address is divided into two parts. The first part designates the network address while the second part designates the host address. The *IP* address space is divided into different network classes. Class A networks are intended mainly for use with a few very large networks, because they provide only 8 bits for the network address field. Class B networks allocate 16 bits, and Class C networks allocate 24 bits for the network address field. Class C networks only provide 8 bits for the host field; however, so the number of hosts per network may be a limiting factor. In all three cases, the leftmost bit(s) indicate the network class. *IP* addresses are written in dotted decimal format; for example, 34.0.0.1.

Figure 2.3 below shows the address formats for Class A, B, and C *IP* networks.

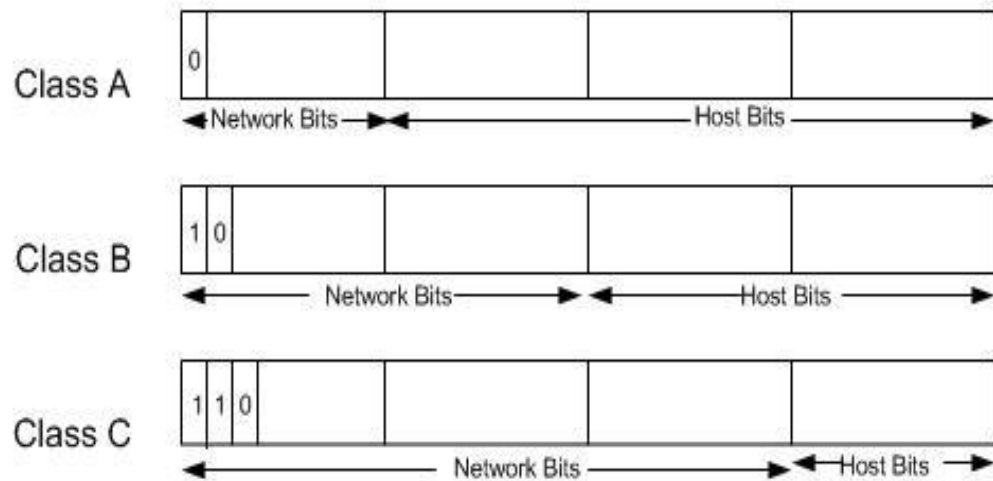


Figure 2.3: Address Formats for Class A, B, and C IP Networks

IP networks also can be divided into smaller units called subnetworks or 'subnets'. Subnets provide extra flexibility for the network administrator. For example, assume that a network has been assigned a Class A address and all the nodes on the network use a Class A address. Further assume that the dotted decimal representation of this network's

address is 34.0.0.0. (All zeros in the host field of an address specify the entire network). The administrator can subdivide the network using subnetting. This is done by "borrowing" bits from the host portion of the address and using them as a subnet field, as depicted in Figure 2.4 below.

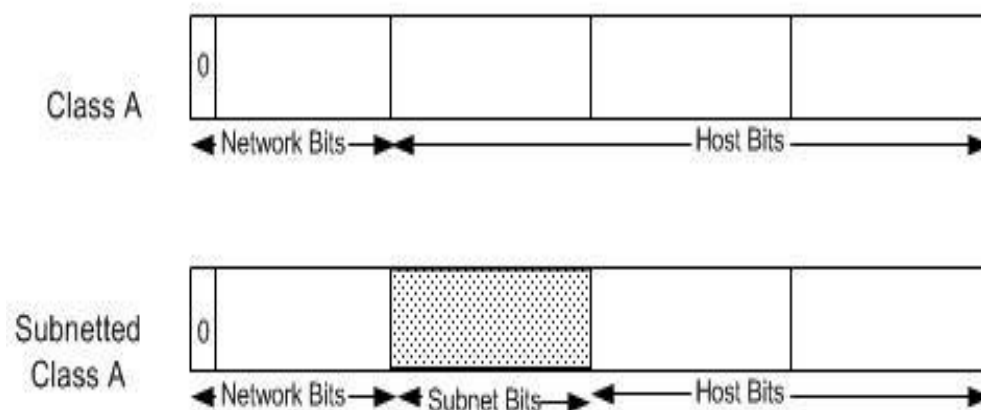


Figure 2.4: "Borrowing" Bits

If the network administrator has chosen to use 8 bits of subnetting, the second octet of a Class A IP address provides the subnet number. In our example, address 34.1.0.0 refers to network 34, subnet 1; address 34.2.0.0 refers to network 34, subnet 2, and so on.

The number of bits that can be borrowed for the subnet address varies.

To specify how many bits are used to represent the network and the subnet portion of the address, IP provides subnet masks. Subnet masks use the same format and representation technique as IP addresses. Subnet masks have ones in all bits except those that specify the host

field. For example, the subnet mask that specifies 8 bits of subnetting for Class A address 34.0.0.0 is 255.255.0.0. The subnet mask that specifies 16 bits of subnetting for Class A address 34.0.0.0 is 255.255.255.0. Both of these subnet masks are pictured in Figure 2.5 below. Subnet masks can be passed through a network on demand so that new nodes can learn how many bits of subnetting are being used on their network.

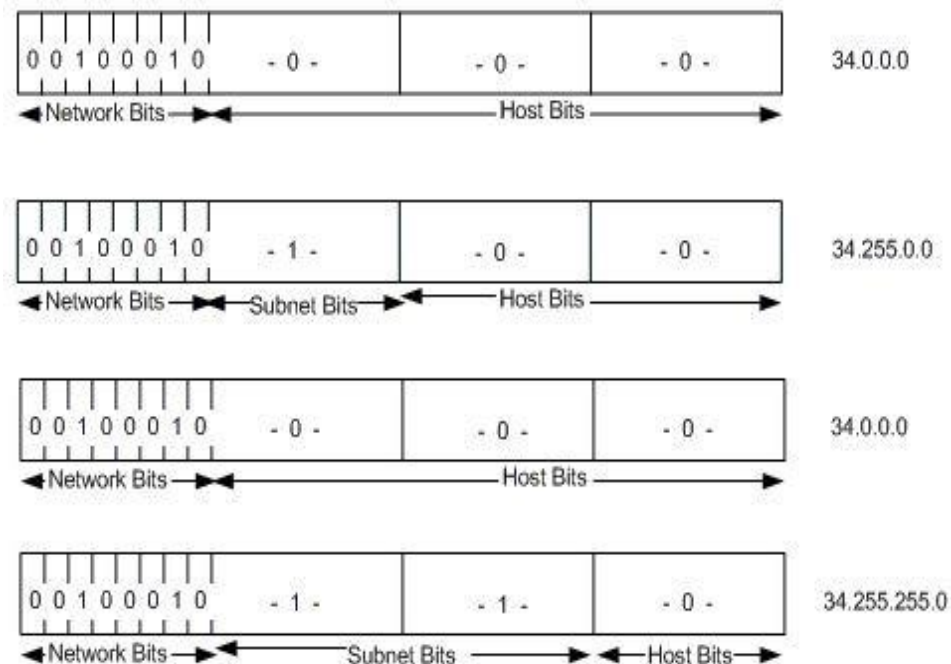


Figure 2.5: Subnet Masks

Traditionally, all subnets of the same network number used the same subnet mask. In other words, a network manager would choose an eight-bit mask for all subnets in the network. This strategy is easy to manage for both network administrators and routing protocols. However, this practice wastes address space in some networks. Some subnets have many hosts and some have only a few, but each consumes an entire subnet number. Serial lines are the most extreme example, because each has only two hosts that can be connected via a serial line subnet. As IP subnets have grown, administrators have looked for ways to use their address space more efficiently.

One of the techniques of subnetting is called Variable Length Subnet Masks (VLSM). With VLSM, a network administrator can use a long mask on networks with few hosts and a short mask on subnets with many hosts. However, this technique is more complex than making them all one size; and addresses must be assigned carefully.

Of course in order to use VLSM, a network administrator must use a routing protocol that supports it. Cisco routers support VLSM with Open Shortest Path First (OSPF), Integrated Intermediate System to Intermediate System (Integrated IS-IS), Enhanced Interior Gateway Routing Protocol (Enhanced IGRP), and static routing. On some media, such as IEEE 802 LANs, IP addresses are dynamically discovered through the use of two other members of the Internet Protocol suite: Address Resolution Protocol (ARP) and Reverse Address Resolution Protocol (RARP). ARP uses broadcast messages to determine the

hardware (MAC layer) address corresponding to a particular network-layer address. ARP is sufficiently generic to allow use of IP with virtually any type of underlying media access mechanism. RARP uses broadcast messages to determine the network-layer address associated with a particular hardware address. RARP is especially important to diskless nodes, for which network-layer addresses usually are unknown at boot time.

3.2 Routing in IP Environments

An internet is a group of interconnected networks. The internet, on the other hand, is the collection of networks that permits communication between most research institutions, universities, and many other organisations around the world. Routers within the internet are organised hierarchically. Some routers are used to move information through one particular group of networks under the same administrative authority and control. (Such an entity is called an autonomous system.) Routers used for information exchange within an autonomous system are called interior routers, and they use a variety of Interior Gateway Protocols (IGPs) to accomplish this end. Routers that move information between autonomous systems are called exterior routers; they use the Exterior Gateway Protocol (EGP) or Border Gateway Protocol (BGP).

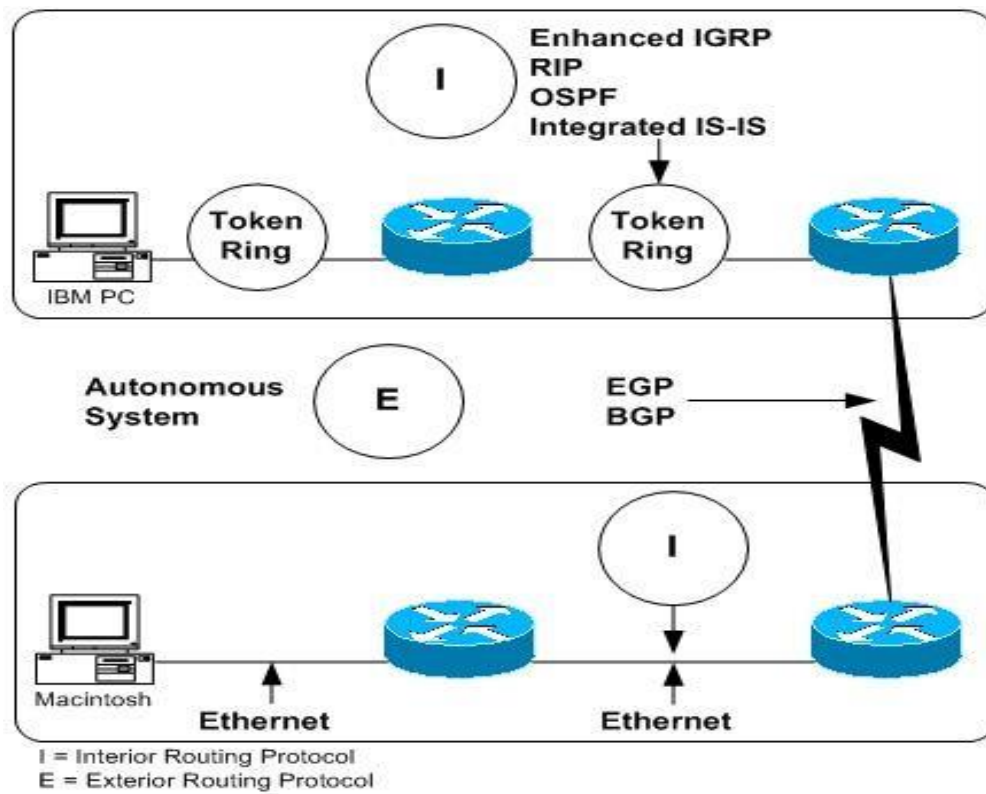


Figure 2.6: Representation of the Internet Architecture

Routing protocols used with *IP* are dynamic in nature. Dynamic routing requires the software in the routing device to calculate routes. Dynamic routing algorithms adapt to changes in the network and automatically select the best routes. In contrast with dynamic routing, static routing calls for routes to be established by the network administrator. Static routes do not change until the network administrator changes them. *IP* routing tables consist of destination address/next hop pairs.

4.0 CONCLUSION

This unit has taken you through the *TCP/IP* networks, its model and implementation procedure.

5.0 SUMMARY

In this unit, you have learnt that:

- *TCP* is a connection-oriented transport protocol that sends data as an unstructured stream of bytes
- *IP* is the primary layer 3 protocol in the internet protocol suite and *IP* represents the heart of the protocol
- one of the techniques of subnetting is called Variable Length Subnet Masks (VLSM); with VLSM, a network administrator can use a long mask on networks with few hosts and a short mask on subnets with many hosts

- *IP* networks also can be divided into smaller units called subnetworks or subnets. Subnets provide extra flexibility for the network administrator
- routing protocols used with *IP* are dynamic in nature.

6.0 TUTOR-MARKED ASSIGNMENT

1. Explain the *TCP/IP* model of networking.
2. With the aid of diagram, explain internet architecture.
3. What are the differences between the Class A, Class B and Class C *IP* networks?

7.0 REFERENCES/FURTHER READING

"Events in British Telecoms History". *Events in British Telecoms History*. Archived from the original on 2003-04-05.

http://web.archive.org/web/20030405153523/http://www.sigtel.com/tel_hist_brief.html. Retrieved November 25, 2005.

A Brief History of Internet.
<http://www.isoc.org/internet/history/brief.shtml>. Retrieved 2009-05-28.

Prasad, K. V. (2009). *Principles of Digital Communication Systems and Computer Networks*. Dreamtech Press.

UNIT3 NETWORK STANDARDS(IEEE802 STANDARDS)

CONTENTS

- 1.0 Introduction
- 2.0 Objectives
- 3.0 Main Content
 - 3.1 Overview of IEEE802
- 4.0 Conclusion
- 5.0 Summary
- 6.0 Tutor-Marked Assignment
- 7.0 References/Further Reading

1.0 INTRODUCTION

IEEE802 refers to a family of IEEE standards dealing with local area networks and metropolitan area networks.

2.0 OBJECTIVES

At the end of this unit, you should be able to:

- state what the *IEEE* 802 standard stands for
- explain the various *IEEE* 802 standards and their descriptions.

3.0 MAIN CONTENT

3.1 Overview of *IEEE* 802

Specifically, the *IEEE* 802 standards are restricted to networks carrying variable-size packets. (By contrast, in cell relay networks, data are transmitted in short, uniformly sized units called cells. Isochronous networks, where data is transmitted as a steady stream of octets, or groups of octets, at regular time intervals, are also out of the scope of this standard). The number 802 was simply the next free number *IEEE* could assign, though “802” is sometimes associated with the date the first meeting was held—February 1980.

The services and protocols specified in *IEEE* 802 map to the lower two layers (data link and physical) of the seven-layer *OSI* networking reference model. In fact, *IEEE* 802 splits the *OSI* data link layer into two sub-layers named Logical Link Control (LLC) and Media Access Control (MAC), so that the layers can be listed like this:

- datalinklayer
- physicallayer

The *IEEE802* family of standards is maintained by the *IEEE802* LAN/MAN Standards Committee (LMSC). The most widely used standards are for the ethernet family, token ring, wireless LAN, bridging and virtual bridged LANs. An individual working group provides the focus for each area. Table 3.1 below presents different *IEEE* network standards and their descriptions.

Table 3.1: *IEEE* Network Standards and their Descriptions

Name Description Note

IEEE802.1	Bridging (networking) and Network Management	
IEEE802.2	LLC	Inactive
IEEE802.3	Ethernet	
IEEE802.4	Tokenbus	Disbanded
IEEE802.5	Defines the MAC layer for a Token Ring	Inactive
IEEE802.6	MANs	Disbanded
IEEE802.7	Broadband LAN using Coaxial Cable	Disbanded
IEEE802.8	Fiber Optic TAG	Disbanded
IEEE802.9	Integrated Services LAN	Disbanded
IEEE802.10	Interoperable LAN Security	Disbanded

IEEE 802.11 a/b/g/n	WirelessLAN(WLAN)&Mesh(Wi- Fi certification)	
IEEE802.12	100BaseVG	Disbanded
IEEE802.13	Unused	
IEEE802.14	Cablemodems	Disbanded
IEEE802.15	WirelessPAN	
IEEE802.15.1	Bluetoothcertification	
IEEE802.15.2	IEEE 802.15 and IEEE 802.11 coexistence	
IEEE802.15.3	High-RatewirelessPAN	
IEEE802.15.4	Low-RatewirelessPAN(e.g. ZigBee)	
IEEE802.15.5	MeshnetworkingforWPAN	
IEEE802.16	BroadbandWirelessAccess(WiMAX certification)	

IEEE802.16.1	Local Multipoint Distribution Service	
IEEE802.17	Resilient packet ring	
IEEE802.18	Radio Regulatory TAG	
IEEE802.19	Coexistence TAG	
IEEE802.20	Mobile Broadband Wireless Access	
IEEE802.21	Media Independent Handoff	
IEEE802.22	Wireless Regional Area Network	
IEEE802.23	Emergency Services Working Group	New (March, 2010)

4.0 CONCLUSION

In this unit, you learnt that the *IEEE802* standards are restricted to networks carrying variable-size packets. You also learnt about the services and protocols specified in *IEEE802*.

5.0 SUMMARY

In this unit, you have learnt that:

- the services and protocols specified in *IEEE802* map to the lower two layers (data link and physical) of the seven-layer *OSI* networking reference model
- the *IEEE802* family of standards is maintained by the *IEEE802* LAN/MAN Standards Committee (LMSC).

- *IEEE802 standards are widely used for the ethernet family, token ring, wireless LAN, bridging and virtual bridged LANs.*
- *there are variations of IEEE802 standards.*

6.0 TUTOR-MARKED ASSIGNMENT

1. Describe the importance of *IEEE802* standards.
2. State all the working group of the *IEEE802* standard and their description.

7.0 REFERENCES/FURTHER READING

The IEEE standard Sourcebook.

"The First Network Email". *The First Network Email.*

<http://openmap.bbn.com/~tomlinso/ray/firstemailframe.html>. Retrieved December 23, 2005.

UNIT4 IMPLEMENTINGSECURITYBEST PRACTICES

CONTENTS

- 1.0 Introduction
- 2.0 Objectives
- 3.0 MainContent
 - 3.1 Overviewof BestNetworkSecurity
 - 3.2 StrategiesforBestNetworkSecurity
- 4.0 Conclusion
- 5.0 Summary
- 6.0 Tutor-MarkedAssignment
- 7.0 References/FurtherReading

1.0 INTRODUCTION

Networksecurity involves maintainingConfidentiality,Integrity and Availability(CIA)ofnetwork.Confidentiality involvespreventing unauthoriseddisclosureof resourcesorinformationon thenetwork. Integritymeanspreventionofunauthorisedmodificationofresourcesor information on the network. Integrity involves prevention of unauthoriseddenialof resourcesor informationonthecomputersystem.

2.0 OBJECTIVES

At the end of this unit, you should be able to:

- explain what network security is all about
- state the best network security
- highlight the features of best security practices.

3.0 MAIN CONTENT

3.1 Overview of Best Network Security

Efficient network security is needed for corporations, universities, schools, public libraries, internet cafes and other applications where administrator has to secure and maintain a lot of network workstations located in different places. Administrator does not need to physically visit workstation to change security settings or install patches. An efficient network security programme is needed for securing, protecting, and maintaining PC workstations within a corporatenetwork. There exist network-based password-protected security software that allows you to completely secure workstation so very your network as well as

maintain them by uploading and installing any executable patches remotely. They support tons of security restrictions, options and tweaks to control access to every bit of windows.

You can deny access to each individual component of several control panel applets, including display, network, passwords, printers, system and internet options; disable the boot keys, context menus, Disk Operating System (DOS) windows, registry editing, internet and network access. You can also hide desktop icons, individual drives, start menu items, and taskbar; apply password protection to windows and restrict users to running specific applications only, control internet usage and much more. In total, best network security support over 600 different security restrictions, options and tweaks that allow you to restrict access to almost every corner of windows.

After installing the remote client service application on your workstations, the maintenance becomes absolutely hassle-free. You just connect your administrator's computer to the net from any place and remotely change security settings, upload and execute patches as well as schedule reboots, shutdowns, and windows explorer restarts just with a click of the mouse. The best solution for corporations, universities, schools, public libraries, internet cafes etc., is to upload and install any executable patches remotely, apply security restrictions, options and tweaks. It supports over 600 different security restrictions.

There are two main types of network security that must be addressed:

Network infrastructure security

- Physical security of network devices

- Preventing unauthorized access to the management software on those devices

Information Security

- Protection of the information or data transmitted over the network

Goals of network security

There are three main goals of network security:

- Confidentiality – only intended recipients can read the data
- Integrity – assurance that the data has not be altered with during transmission
- Availability – assurance of timely and reliable access to data for authorized users

3.2 Strategies for Best Network Security

Then (10) key strategies for ensuring efficient network security are as listed below.

i.

Create security training and awareness programme within the organisation. Let the personnel in your organisation be aware of:

- emerging threats
- possible consequences of a single security breach.
- security requirement

There must be regular security training. Collaborations with other cyber security organisations are also necessary.

ii. Have an efficient access control system

In order to control insider threats, there must be clear definition of roles and responsibilities in organisations. Efficient authorisation system should be implemented in order to limit the activities of personnel.

iii. Implement trusted and efficient intrusion prevention and detection systems

To effectively present cyberattacks, a trusted and adaptive intrusion prevention and detection system must be implemented. Nowadays, it is also necessary to implement an intrusion forecasting system.

iv. Identify and document all resources on the network

All categories of resources on the network must be well documented and their status monitored and reported.

v.

All connections to the network must be known, documented and monitored

Identify and monitor all connections (cable, wireless, local or internet) and ensure that these connections are well protected. All unwanted or unnecessary connections should be disconnected.

vi. Have an effective programme for the risk analysis and incident building

A periodic and thorough analysis of potential risks to network resources and vulnerability assessment is an essential requirement for effective cybersecurity management.

vii. Prevent unknown applications

Implement strategies that disallow any unknown application to run on your system.

viii. Use security applications from reliable vendors

Trusted security tools from known vendors should be used. The vendors must implement security features in form of patches and upgrades.

ix. Establish effective configuration management processes

Since changes in hardware and software configurations can introduce vulnerabilities, an assessment of these security implications of all changes should be carried out.

x. Ensure periodic evaluation of all security apparatus

Organisations should ensure periodic testing and evaluation of all security apparatus.

SELF-ASSESSMENT EXERCISE

What do you understand by cybersecurity?

4.0 CONCLUSION

This unit discusses the security issues in a network; and you have also learnt about the best security practices.

5.0 SUMMARY

In this unit, you have learnt that:

- network security involves maintaining confidentiality, integrity and availability (CIA) of network
- efficient network security is needed for corporations, universities, schools, public libraries, internet cafes and other applications where an administrator has to secure and maintain a lot of network workstations located in different places.
- there are certain key strategies required for ensuring best network security practices

6.0 TUTOR-MARKED ASSIGNMENT

1. Why do we need to protect a network?
2. Describe how the network within your organisation is protected.

7.0 REFERENCES/FURTHER READING

<http://www.azimuthsystems.com/>

NASA Successfully Tests First Deep Space Internet. NASA media release 08-298, November 18, 2008 Archived.

MODULE3 ENTERPRISENETWORK

Unit1	CreatingEnterpriseNetwork
Unit2	PlanningandSelectionof EnterpriseNetwork
Unit3	LANandWAN

UNIT1 CREATINGENTERPRISENETWORK

CONTENTS

1.0	Introduction
2.0	Objectives
3.0	MainContent
3.1	Overviewof EnterpriseNetwork
3.2	PredefinedEnterpriseNetwork
3.2.1	ResidualNetworks
3.4	Threatto EnterpriseNetwork
3.4.1	TypesofThreat
4.0	Conclusion
5.0	Summary
6.0	Tutor-MarkedAssignment
7.0	References/FurtherReading

1.0 INTRODUCTION

This unit introduces you to enterprise network. An enterprise network is used to connect computers and other digital equipment of an organisation together.

2.0 OBJECTIVES

At the end of this unit, you should be able to:

- explain enterprise network
- describe predefined enterprise network.

3.0 MAIN CONTENT

3.1 Overview of Enterprise Network

An enterprise private network is a network built by an enterprise to interconnect various company sites, e.g., production sites, head offices, remote offices, shops, in order to share computer resources. Enterprise networks are used for configuring access rules in an enterprise policy.

that can be applied to any array in the enterprise and for configuring enterprise network rules that apply to all arrays in the enterprise. Only enterprise networks can be used to create enterprise-level rules. Enterprise networks can also be used for defining array-level access and publishing rules and for defining array-level network rules. However, you cannot use array-level networks when creating enterprise-level rules.

Any number of user-defined enterprise networks can also be included in a network defined in an array by including references to them. An enterprise network which is a set of IP address ranges corresponds exactly to the IP addresses included in a protected array-level network, such as the Internal network, defined in one array can be used to reference that network in all the other arrays of the enterprise.

When you configure enterprise networks, you specify only the IP address ranges and do not specify any of the other properties that you would define for array-level networks. In particular, you cannot configure Network Load Balancing or Cache Array Routing Protocol (CARP) for an enterprise network.

The IP addresses that are included in an enterprise network are excluded from the default external network in each array in the enterprise, even if the enterprise network is not included in any network defined in the array.

3.2 Predefined Enterprise Networks

The following predefined enterprise networks are recreated upon installation:

- External
- LocalHost
- QuarantinedVPNClients
- VPNClients

These predefined enterprise networks implicitly define the same *IP* address sets as their array-level counterparts. They can be used for defining rules in an enterprise policy and for defining enterprise network rules. When an enterprise policy is assigned to an array, each predefined enterprise network in a rule will be interpreted as the array-level network of the same name. For example, you can create an enterprise access rule that applies to requests sent to the local host enterprise network. When a request is handled in an array to which the enterprise policy containing this access rule is assigned, the rule will apply to the

IP addresses in the local host network on the array member handling the request.

3.2.1 Residual Networks

IP addresses that belong to a configurable enterprise network, but do not belong to any configurable array-level network are considered to be part of a residual network.

3.4 Threat to Enterprise Network

Today, there is an ever-growing dependency on computer networks for business transactions. With the free flow of information and the high availability of many resources, managers of enterprise networks have to understand all the possible threats to their networks. These threats take many forms, but all result in loss of privacy to some degree and possibly malicious destruction of information or resources that can lead to large monetary losses.

3.4.1 Types of Threat

Listed below are some common types of threats—especially when the enterprise network is connected to the Internet.

a. Insider threats

An insider attack involves someone from the inside, such as a disgruntled employee, attacking the network. Insider attacks can be malicious or not malicious. Malicious insiders intentionally eavesdrop, steal, or damage information; use information in a fraudulent manner; or deny access to other authorised users. This can be controlled using efficient authentication and authorisation systems and IDS.

b. Malware

One of the biggest threats to computer users on the internet today is malware. Malware are malicious codes that can cause distortion on your browser and eventually cripple your system. They can hijack your browser, redirect your search attempts, serve up nasty pop-ups, track what websites you visit, and prevent you from performing certain functions. Examples are worm, viruses, Trojan horse, spyware, adware, dialers, hijackers, etc. They are controlled using frequently updated and highly rated antivirus, anti-spyware, etc. Efficient management of patches, SPAM filters and two-way authentication are also used for their control.

c. Socially engineered attacks

Most traditional social engineering attacks capitalise on this vulnerability. Examples are phishing, *pharming*, masquerading. General technical controls such as a firewall, internet content filtering, antivirus software, anti-spam software, security awareness and patch management can help reduce or eliminate many phishing attacks.

d. Authentication attacks

An attacker tries to crack the passwords stored in a network account database or a password-protected file. Examples include dictionary, a brute-force, shoulder surfing, keylogging and hybrid attacks. They are controlled using robust and multifactor authentication and authorisation systems.

SELF-ASSESSMENT EXERCISE

Mention other types of threats that are not listed here

4.0 CONCLUSION

This unit has taken you through the rudiments of enterprise network. It has also introduced you to residual networks.

5.0 SUMMARY

In this unit, you have learnt that:

- an enterprise private network is a network built by an enterprise to interconnect various company sites, e.g., production sites, head offices, remote offices, shops, in order to share computer resources.
- common types of threats to enterprise network are insider attacks, socially engineered, malware and authentication attacks.
- predefined enterprise network are used for defining rules in an enterprise policy and for defining enterprise network rules.

6.0 TUTOR-MARKED ASSIGNMENT

1. Discuss the three types of threat to enterprise network.
2. What do you understand by enterprise network?

7.0 REFERENCES/FURTHER READING

Abbate, Janet (1999). *Inventing the Internet*. Cambridge: MIT.

[Bemer, Bob, "A History of Source Concepts for the Internet/Web".](#)

UNIT2 PLANNINGANDSELECTIONOF ENTERPRISENETWORK

CONTENTS

- 1.0 Introduction
- 2.0 Objectives
- 3.0 MainContent
 - 3.1 Overview
 - 3.2 Stepsto EffectiveNetworkPlanningDesign
- 4.0 Conclusion
- 5.0 Summary
- 6.0 Tutor-MarkedAssignment
- 7.0 References/FurtherReading

1.0 INTRODUCTION

Thisunitgivesyouanoverviewofnetworkandthestepsinvolvedin planningandselectingenterprisenetwork.

2.0 OBJECTIVES

At theendof thisunit,youshouldbeableto:

- highlighthowtomakeproperplanningforeveryphaseofan enterprisenetwork

□ explain the design phases in enterprise networks.

3.0 MAIN CONTENT

3.1 Overview

Network planning and design methodology describes a process with 9 specific steps and a sequence for those activities. It is an engineering life cycle that supports technical initiatives such as Windows migration, IP telephony and wireless design, to name a few examples. The methodology begins with examining company business requirements. It is absolutely essential that you understand the company business model, business drivers and how they are growing from a business perspective. That will build the foundation for a design proposal that serves the business, technical and operational requirements of the company.

3.2 Steps to Effective Network Planning and Design

Here, let us look at the following steps.

Step 1- Business requirements

Any design project starts with an understanding of what the company does and what they need to accomplish from a business perspective.

This begins with an understanding of their business model, which really describes how their company works from an operational and business perspective to generate revenues and reduce costs. Many vendors today have conducted their own return on investment (ROI) studies for new implementations such as unified communications and telephony. It is an effective sales tool that illustrates the cost benefits compared with investment over a specified period of time.

Some typical business drivers are listed below.

- Reduce operating costs
- Generate revenue
- Client satisfaction
- Employee productivity

This is a list of some typical project business requirements.

- Budget constraints
- Office consolidations
- Company mergers and acquisitions
- Business partner connectivity
- Telecommuter remote access
- New offices and employees
- New data center applications
- Reduce network outage costs
- Cost effective network management
- Vendor contracts

Step2- Design requirements

Now that you have learnt the basic business requirements of the company, you can determine the standard and specific design requirements. The design requirements process is focused on defining requirements from a technical perspective. Those requirements along with the business requirements will build the framework that is used to define infrastructure, security and management. Design requirements are defined as standard and miscellaneous. The standard design requirements are generic and represent those considered with many

design projects. Miscellaneous requirements are those that are not defined with any of the standard requirements.

Standard design requirements are:

- performance
- availability
- scalability
- standards compatibility
- rapid deployment

Step3- Network assessment

A network assessment is conducted after we have finished the business and design requirements of the company. A network assessment provides a quick snapshot of the current network with an examination of the infrastructure, performance, availability, management and security. That information is utilised for making effective strategy recommendations and design proposals to the client concerning specific information systems modifications. The network assessment model has three sequential activities, which are- assessment, analysis and recommendations. The current network is examined using five primary surveys, namely- infrastructure, performance, availability, management and security. When the surveys are completed, the information collected is then reviewed for trends, problems and issues that are negatively affecting the network.

Step4- Infrastructureselection

Afterdoinganetworkassessmentwearereadyto startselectingspecific infrastructure components for the networkdesign. This phase starts building theinfrastructure withaspecificsequencethatpromotes effectiveequipment selectionanddesign.Itisimportantthatyou considerbusinessrequirements,designrequirements andthenetwork assessmentwhenbuildingyourinfrastructure.

The following numbered list describes the specific infrastructure componentsandtheirparticularsequence.

- a. EnterpriseWANTopology
- b. Campustopology
- c. Traffic model
- d. Equipmentsselection
- e. Circuits
- f. Routingprotocoldesign g. Addressing

- h. Naming conventions
- i. IOS Services
- j. Domain name services
- k. DHCP services

Step 5: Security strategy

We must now define a security strategy for securing the infrastructure. The need for enterprise network security should not be ignored with the proliferation of the Internet. Companies are continuing to leverage the public infrastructure for connecting national and international offices, business partners and new company acquisitions. These security requirements and network assessment recommendations should drive the selection of security equipment, protocols and processes. It identifies what assets must be protected, what users are allowed access and how those assets will be secured.

Step 6- Network management strategy

This section will define a network management strategy for managing all equipment defined from infrastructure and security. It is necessary to define how the equipment is going to be monitored and determine if the current management strategy is adequate or if new applications, equipment, protocols and processes must be identified. Management components are then integrated with infrastructure and security to finish building the proposed design. These primary elements comprise any well-

defined management strategy and should be considered when developing your strategy.

- Management groups
- SNMP applications
- Monitored devices and events

Step 7- Proof of concept

All infrastructure, security and management components must now be tested with a proof of concept plan. It is important to test the current design, configuration and IOS versions in a non-production environment or on the production network with limited disruption. Implementation of new network modules at a router, for instance, could require that you change the current IOS version that is implemented. Making those changes can affect WAN or campus modules already installed at production routers. That is the real value of doing a proof of concept and certifying that the new equipment and IOS versions integrate with each device as well as the network. The following list describes the advantages of doing a proof of concept with your network design. The

proof of concept test results should be examined and used to modify current infrastructure, security and management specifications before generating a design proposal. The proof of concept model suggested herein involves prototype design, equipment provisioning, defining tests, building equipment scripts and examining test results.

1. Prototype design
2. Provision of equipment
3. Define tests
4. Build equipment scripts
5. Review test results

Step 8- Design proposal/review

With the proof of concept finished, you are now ready to build a design proposal for the design review meeting. Your intended audience could be the Director, CIO, CTO, Senior network engineer, Consultant or anyone that is approving a budget for the project. It is important to present your ideas with clarity and professionalism. If a presentation is required, Powerpoint slides work well and could be used to support concepts from the design proposal document. The focus is on what comprises a standard design proposal and the sequence for presenting that information.

The working design proposal is presented to the client after addressing any concerns from proof of concept assurance testing. The design review is an opportunity for you to present your design proposal to the client and discuss any issues. It is an opportunity for the client to identify concerns they have and for the design engineer to clarify issues. The focus is to agree on any modifications, if required, and make changes to the infrastructure, security and management before implementation starts. Business and design requirements can change from when the project started which sometimes will necessitate changes to infrastructure, security and management specifications. Any changes should then go through proof of concept testing again before final changes to the design proposal.

Step 9- Implementation

The final step entails defining an implementation process for the specified design. This describes a suggested implementation methodology of the proposed design, which should have minimal disruption to the production network. As well, it should be efficient and as cost effective as possible. As with previous methodologies, there is a sequence that should be utilised as well.

Once the implementation is finished, the network has to be monitored to checkmate problems which may arise. Design and configuration modifications are then made to address any problems or concerns.

SELF-ASSESSMENT EXERCISE

A manufacturing company is interested in setting up a network. Describe how you will carry out network planning and design for the organisation.

4.0 CONCLUSION

In this unit, you have learnt the necessary steps to consider in implementing an enterprise network.

5.0 SUMMARY

In this unit, you have learnt that:

- network planning and design methodology is an important activity towards having an efficient network
- network planning and design involve in specific steps, namely- identification of business requirements, design requirements,

network assessment, infrastructure selection, security strategy, network management strategy, proof of concept and implementation.

6.0 TUTOR-MARKED ASSIGNMENT

1. What are the possible consequences of lack of effective network planning?
2. State and explain the phases involved in creating an enterprise network

7.0 REFERENCES/FURTHER READING

Krol, Ed. (1992). *Whole Internet User's Guide and Catalog*. O'Reilly & Associates, 1992.

Scientific American Special Issue on Communications, Computers, and Networks, September, 1991.

<http://www.eBookmall.com>

Prasad, K. V. (2009). "Principles of Digital Communication Systems and Computer Networks", Dreamtech Press.

UNIT3 LOCAL AREA NETWORK(LAN) AND WIDE AREA NETWORK(WAN)

CONTENTS

- 1.0 Introduction
- 2.0 Objectives
- 3.0 Main Content
 - 3.1 Understanding WAN and LAN
 - 3.2 Difference between LAN and WAN
- 4.0 Conclusion
- 5.0 Summary
- 6.0 Tutor-Marked Assignment
- 7.0 References/Further Reading

1.0 INTRODUCTION

Computers networked together in a self-contained group form a Local Area Network or LAN. A LAN typically is contained within a single building or a group of neighbouring buildings. Two computers linked together at home are the simplest form of a LAN. Several hundred computers cabled together across several buildings at school form a more complex LAN. LANs are usually connected with coaxial or CAT5 cable.

On the other hand, a WAN is geographically large. It is often formed by the joining together of LANs in distant places. A national banking

organisation, for example, may use a WAN to connect all of its branches across the country. The difference between LANs and WANs is getting blurry as fibre optic cables have allowed LAN technologies to connect devices many kilometers apart. WANs are usually connected using the internet, ISDN landlines or satellite.

2.0 OBJECTIVES

At the end of this unit, you should be able to:

- discuss the background of LAN and WAN
- explain the differences between LAN and WAN
- state the relationship between LAN and WAN.

3.0 MAIN CONTENT

3.1 Understanding WAN and LAN

A Local Area Network (LAN) exists in a house or a university campus, while a Wide Area Network (WAN) exists across many office buildings separated by a vast distance. The office buildings in a WAN may be in different countries or even continents. For example, the headquarters building may be in USA, the regional office building may be in UK, and the branch office building may be in India. The workers in the three buildings can use WAN to collaborate with one another. The internet can also be considered as a WAN.

3.2 Differences between LAN and WAN

One of the differences between LAN and WAN, is the speed of the network. The maximum speed of a LAN can be 10 megabits per second, while the speed of a WAN can go up to 150 megabits per second. This means the speed of a WAN, is one-tenth of the speed of a LAN. A WAN is usually slower because it has lower bandwidth.

Computers in a LAN can share a printer, if they are all in the same LAN. On the other hand, a WAN cannot share a printer, so a computer in one country cannot use a printer in another country. A LAN does not need a dedicated computer to direct traffic to and from the internet, unlike a

WAN that needs a special-purpose computer, whose only purpose is to send and receive data from the internet.

Another basis for comparing LAN and WAN is the cost of the network. A WAN is more expensive than a LAN. It is easier to expand a LAN than a WAN. The equipment needed for a LAN is a Network Interface Card (NIC), a switch and a hub. On the other hand, the equipment needed to connect a WAN to the internet is a modem and a router. The modem may be a cable modem or a DSL modem that is connected to a wall jack, while the routers should be configured so that it can handle the packets traveling between the WAN and the internet.

Also, between LAN and WAN, there is a difference in the networking standard used. A LAN uses the ethernet standard, while a WAN uses the T1 standard. Before ethernet, the protocols used for LAN were Attached Resource Computer Network (ARCNET) and token ring. The protocols used for WAN are frame relay and Asynchronous Transfer Mode (ATM). Another protocol for WAN is Packet Over SONET/SDH (POS), where SONET stands for Synchronous Optical Networking and SDH stands for Synchronous Digital Hierarchy. The first WAN protocol was X.25, while an advanced WAN protocol is multiprotocol label switching.

(mpls). The hardware in a LAN is connected with 10base-t cable connectors, while a WAN is connected via leased lines or satellites.

Furthermore, a LAN is easy to set up, as you need to slip the NIC into the PCI slot (for desktop computers) or PCMCIA slot (for laptop computers). You also need to install the driver for the NIC. The NIC can be connected to the network using the RJ45 port.

On the other hand, a WAN is very difficult to set up. There is often an appliance to optimise the WAN. There is also a device to cache WAN data, so workers in the branch office can quickly access documents. The router also has quality of service (QoS) built in, so that it gives priority to certain kinds of traffic.

There are various topologies available in LAN and WAN networking.

The most common topologies in LAN and WAN networks are ring and star. The ring topology is a network in which every node (every computer) is connected to exactly two other nodes. The star topology is a network in which all the nodes (called leaf nodes or peripheral nodes) are connected to a central node.

Table 3.2: LAN versus WAN comparison chart

	LAN	WAN
Stands For	Local Area Network	Wide Area Network
Covers	Local areas only (e.g., homes,	Large geographic areas (e.g., cities,

	LAN	WAN
	offices, schools)	states, nations)
Definition	LAN (Local Area Network) is a computer network covering a small geographic area, like a home, office, school, or group of buildings.	WAN (Wide Area Network) is a computer network that covers a broad area (e.g., any network whose communications links cross metropolitan, regional, or national boundaries over a long distance).
Speed	High speed (1000 mbps)	Less speed (150 mbps)
Data transfer rates	LANs have a high data transfer rate.	WANs have a lower data transfer rate compared to LANs.
Example	The network in an office building can be a LAN	The Internet is a good example of a WAN
Technology	Tend to use certain connectivity technologies, primarily Ethernet and Token Ring	WANs tend to use technologies like MPLS, ATM, Frame Relay and X.25 for connectivity over longer distances
Connection	One LAN can be connected to other LANs over any distance via telephone lines and radio waves.	Computers connected to a wide-area network are often connected through public networks, such as the telephone system. They can also be connected through leased lines or satellites.
Components	Layer 2 devices like switches and bridges. Layer 1 devices like hubs and repeaters.	Layers 3 devices Routers, Multi-layer Switches and Technology specific devices like ATM or Frame-relay Switches etc.

	LAN	WAN
Fault Tolerance	LANs tend to have fewer problems associated with them, as there are smaller number of systems to deal with.	WANs tend to be less fault tolerant as they consist of large number of systems.
Data Transmission Error	Experiences fewer data transmission errors	Experiences more data transmission errors as compared to LAN
Ownership	Typically owned, controlled, and managed by a single person or organization.	WANs (like the Internet) are not owned by any one organization but rather exist under collective or distributed ownership and management over long distances.
Set-up costs	If there is a need to set-up a couple of extra devices on the network, it is not very expensive to do that.	For WANs since networks in remote areas have to be connected the set-up costs are higher. However WANs using public networks can be setup very cheaply using just software (VPN etc).
Geographical Spread	Have a small geographical range and do not need any leased telecommunication lines	Have a large geographical range generally spreading across boundaries and need leased telecommunication lines
Maintenance costs	Because it covers a relatively small geographical area, LAN is easier to maintain at relatively low costs.	Maintaining WAN is difficult because of its wider geographical coverage and higher maintenance costs.

	LAN	WAN
Bandwidth	High bandwidth is available for transmission.	Low bandwidth is available for transmission.
Congestion	Less congestion	More congestion

SELF-ASSESSMENT EXERCISE

Get valid diagrams to represent the structure of LAN and WAN

4.0 CONCLUSION

This unit has taken you through the nitty-gritty of Local Area Network and Wide Area Network classes. It also examined the relationship between LAN and WAN and also the differences between the two types of network.

5.0 SUMMARY

In this unit, you have learnt that:

- a local area network (LAN) exists in a house or a university campus, while a wide area network (WAN) exists across many office buildings separated by a vast distance.
- one of the differences between LAN and WAN, is the speed of the network.
- between LAN and WAN, as well, there is a difference in the networking standard used.

6.0 TUTOR-MARKED ASSIGNMENT

1. Discuss the differences between the LAN and the WAN classes.
2. Discuss other types of networks apart from LAN and WAN

7.0 REFERENCES/FURTHER READING

["http://en.wikipedia.org/wiki/History_of_the_Internet"](http://en.wikipedia.org/wiki/History_of_the_Internet)

Prasad, K. V. (2009). *Principles of Digital Communication Systems and Computer Networks*. Dreamtech Press.

MODULE4 COMMUNICATIONTECHNOLOGY

Unit1	ModemandModulationConcepts
Unit2	Multiplexers
Unit3	DigitalTechnologies
Unit4	SignalTransmissionandImpairment

UNIT1 MODEMANDMODULATIONCONCEPTS

CONTENTS

1.0	Introduction
2.0	Objectives
3.0	MainContent
3.1	History
3.2	Significanceof DigitalModulation
3.3	AnalogModulationMethods
3.4	DigitalModulationMethods
3.5	CommonDigitalModulationTechniques
3.6	FundamentalDigitalModulationMethods
3.7	ModulatorandDetectorPrinciplesof Operation
3.8	ListofknownDigitalModulationTechniques
4.0	Conclusion
5.0	Summary
6.0	Tutor-MarkedAssignment
7.0	References/FurtherReading

1.0 INTRODUCTION

A **modem** (**m**odulator-**d**emodulator) is a device that modulates an analog carrier signal to encode digital information, and also demodulates such a carrier signal to decode the transmitted information. The goal is to produce a signal that can be transmitted easily and decoded to reproduce the original digital data. Modems can be used over any means of transmitting analog signals, from dial-up to radio. The most familiar example is a voice-band modem that turns the digital data of a personal computer into modulated electrical signals in the voice frequency range of a telephone channel. These signals can be transmitted over telephone lines and demodulated by another modem at the receiver's side to recover the digital data.

Modulation is the process of facilitating the transfer of information over a medium. Sound transmission in air has limited range for the amount of power your lungs can generate. To extend the range your voice can reach, we need to transmit it through a medium other than air, such as a

phoneline or radio. The process of converting information (voice in this case) so that it can be successfully sent through a medium (wire or radio waves) is called modulation.

2.0 OBJECTIVES

At the end of this unit, you should be able to:

- define a modem
- explain modulation
- state the types of modulation
- describe common digital modulation techniques.

3.0 MAIN CONTENT

3.1 History

Newswire services in 1920s used multiplex equipment that met the definition, but the modem function was incidental to the multiplexing function, so they are not commonly included in the history of modems. Modems grew out of the need to connect teletype machines over ordinary phone lines instead of more expensive leased lines which had previously been used for current loop-based teleprinters and automated telegraphs. George Stibitz connected a New Hampshire teletype to a computer in New York City by a subscriber telephone line in 1940.

In 1943, IBM adapted this technology to their unit record equipment and was able to transmit punched cards at 25 bits/second. Mass-produced modems in the United States began as part of the SAGE air-defense system in 1958; connecting terminals at various airbases, radar sites, and command-and-control centers to the SAGE director centers scattered around the U.S. and Canada. SAGE modems were described by AT&T's Bell Labs as conforming to their newly published Bell 101 dataset standard. While they ran on dedicated telephonelines, the devices at each end were no different from commercial acoustically coupled Bell 101, 110 baud modems.

In the summer of 1960, the name Data-phonewas introduced to replace the earlier term digital subset. The 202 data-phone was a half-duplex asynchronous service that was marketed extensively in late 1960. In 1962, the 201A and 201B data-phones were introduced. They were synchronous modems using two-bit-per-baud phase-shift keying (PSK). The 201A operated half-duplex at 2,000 bit/s over normal phonelines, while the 201B provided full duplex 2,400 bit/s service on four-wire leased lines, the send and receive channels running on their own set of two wires each.

The famous Bell 103A dataset standard was also introduced by Bell Labs in 1962. It provided full-duplex service at 300 baud over normal phone lines. Frequency-shift keying was used with the call originator transmitting at 1,070 or 1,270 Hz and the answering modem transmitting at 2,025 or 2,225 Hz. The readily available 103A2 gave an important boost to the use of remote low-speed terminals such as the KSR33, the ASR33, and the IBM 2741. AT&T reduced modem costs by introducing the originate-only 113D and the answer-only 113B/C modems.

3.2 Significance of Digital Modulation

The aim of **digital modulation** is to transfer a digital bit stream over an analog bandpass channel; for example over the public switched telephone network (where a bandpass filter limits the frequency range to between 300 and 3400 Hz), or over a limited radio frequency band.

The aim of **analog modulation** is to transfer an analog baseband (or low pass) signal, for example an audio signal or TV signal, over an analog bandpass channel, for example a limited radio frequency band or a cable TV network channel.

Analog and digital modulation facilitate Frequency Division Multiplexing (FDM), where several low pass information signals are transferred simultaneously over the same shared physical medium, using separate passband channels.

The aim of **digital baseband modulation** methods, also known as line coding, is to transfer a digital bit stream over a baseband channel, typically a non-filtered copper wire such as a serial bus or a wired local area network.

The aim of **pulse modulation** methods is to transfer a narrowband analog signal, for example a phone call over a wideband baseband channel or, in some of these schemes, a bit stream over another digital transmission system.

In music synthesizers, modulation may be used to synthesise waveforms with a desired overtone spectrum. In this case the carrier frequency is typically in the same order or much lower than the modulating waveform. See for example frequency modulation synthesis or ring modulation.

3.3 Analog Modulation Methods

In analog modulation, the modulation is applied continuously in response to the analog information signal.

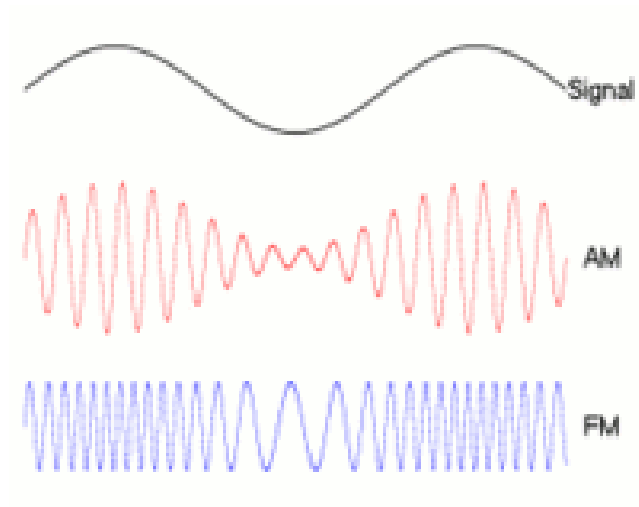


Figure 1.1: Analog Modulation Signal

A low-frequency message signal (top) may be carried by an AM or FM radiowave.

Common analog modulation techniques are listed below.

- Amplitude Modulation (AM) - the amplitude of the carrier signal is varied in accordance to the instantaneous amplitude of the modulating signal
- Double-Sideband Modulation (DSB)
 - Double-sideband modulation with carrier (DSB-WC) is used on the AM radio broadcasting band
 - Double-sideband suppressed-carrier transmission (DSB-SC)

- Double-sideband reduced carrier transmission (DSB-RC)
- Single-sideband modulation (SSB, or SSB-AM),
- SSB with carrier (SSB-WC)
- SSB suppressed carrier modulation (SSB-SC)
- Vestigial sideband modulation (VSB, or VSB-AM)
- Quadrature Amplitude Modulation (QAM)
- Angle modulation
- Frequency Modulation (FM) - here the frequency of the carrier signal is varied in accordance to the instantaneous amplitude of the modulating signal.

- Phase Modulation (PM)-here the phase shift of the carrier signal is varied in accordance to the instantaneous amplitude of the modulating signal

The accompanying figure shows the result of (amplitude-)modulating a signal onto a carrier (both of which are sine waves). At any point along the y -axis, the amplitude of the modulated signal is equal to the sum of the carrier signal.

3.4 Digital Modulation Methods

In digital modulation, an analog carrier signal is modulated by a digital bitstream. Digital modulation methods can be considered as digital-to-analog conversion, and the corresponding demodulation or detection as analog-to-digital conversion. The changes in the carrier signal are chosen from a finite number of M alternative symbols (the modulation alphabet).

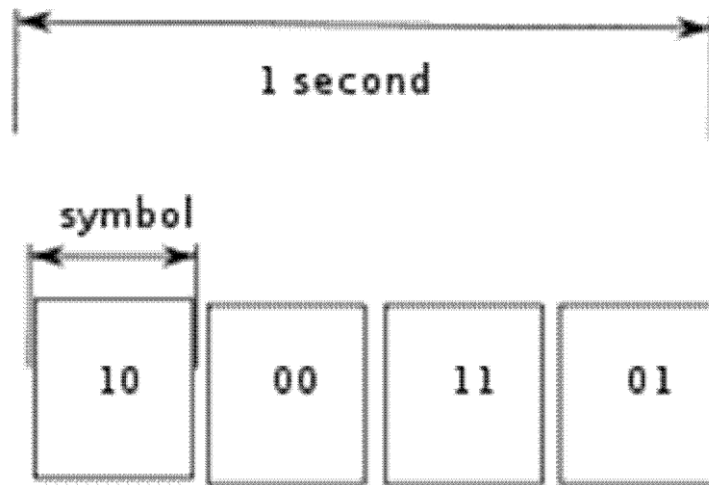


Figure 1.2: Schema of 4Baud (8 bps) Data Link

A simple example-a telephoned line is designed for transferring audible sounds-for example tones, and not digital bits (zeros and ones). Computers may however communicate over a telephoned line by means of modems, which are representing the digital bits by tones, called symbols. If there are four alternative symbols (corresponding to a musical instrument that can generate four different tones, one at a time), the first symbol may represent the bit sequence 00, the second 01, the third 10 and the fourth 11. If the modem plays a melody consisting of 1000 tones per second, the symbol rate is 1000 symbols/second, or baud. Since each tone (i.e., symbol) represents a message consisting of two digital bits in this example, the bit rate is twice the symbol rate, i.e. 2000

bits per second. This is similar to the technique used by dialup modems as opposed to DSL modems.

3.5 Common Digital Modulation Techniques

These are listed below:

- Amplitude-Shift Keying (ASK)
- Frequency-Shift Keying (FSK)
- Phase-Shift Keying (PSK)

i. Amplitude-Shift Keying (ASK)

In ASK, the amplitude of the carrier is changed in response to information, and all else is kept fixed. Bit 1 is transmitted by a carrier of one particular amplitude. To transmit 0, we change the amplitude keeping the frequency constant. On-Off Keying (OOK) is a special form of ASK.

Advantages

Scheme is simple, so it is easy to implement transmitter and receiver with several components.

Low bandwidth requirements

Disadvantages

ASK is heavily effected by noise and interference and can be easily demodulated.

ii. Frequency-Shift Keying (FSK)

In FSK, we change the frequency in response to information, one particular frequency for 1 and another frequency for 0 as shown below for the same bit sequence. In the example below, frequency f_1 for bit 1 is higher than f_2 used for the 0 bit.

$$FSK(t) = \begin{cases} \sin(2\pi f_1 t) & \text{for bit 1} \\ \sin(2\pi f_2 t) & \text{for bit 0} \end{cases}$$

Advantages

FSK is insensitive to channel fluctuations and not easily effected by noise.

Resilient to signal strength variations

Does not require linear amplifiers in the transmitter

Disadvantages

FSK is a low performance type of digital modulation.

iii. Phase-Shift Keying (PSK)

In *PSK*, we change the phase of the sinusoidal carrier to indicate information. Phase, in this context, is the starting angle at which the sinusoid starts. To transmit 0, we shift the phase of the sinusoid by 180°. Phase shift represents the change in the state of the information in this case.

$$PSK(t) = \begin{cases} \sin(2\pi f t) & \text{for bit 1} \\ \sin(2\pi f t + \pi) & \text{for bit 0} \end{cases}$$

Advantages

PSK, phase shift keying enables data to be carried on a radio communication signal in a more efficient manner than Frequency Shift Keying (FSK), and some other forms of modulation.

Disadvantages

Implementation is complex and expensive.

3.6 Fundamental Digital Modulation Methods

The most fundamental digital modulation techniques are based on keying:

- in the case of *PSK* (phase-shift keying), a finite number of phases are used.
- in the case of *FSK* (frequency-shift keying), a finite number of frequencies are used.
- in the case of *ASK* (amplitude-shift keying), a finite number of amplitudes are used.
- in the case of *QAM* (quadrature amplitude modulation), a finite number of at least two phases, and at least two amplitudes are used.

In *QAM*, an in-phase signal (the *I* signal, for example a cosine waveform) and a quadrature phase signal (the *Q* signal, for example a sine wave) are amplitude modulated with a finite number of amplitudes, and summed. It can be seen as a two-channel system, each channel using *ASK*. The resulting signal is equivalent to a combination of *PSK* and *ASK*.

In all of the above methods, each of these phases, frequencies or amplitudes are assigned a unique pattern of binary bits. Usually, each phase, frequency or amplitude encodes an equal number of bits. This number of bits comprises the symbol that is represented by the particular phase, frequency or amplitude.

If the alphabet consists of $M = 2^N$ alternative symbols, each symbol represents a message consisting of N bits. If the symbol rate (also known as the baud rate) is f_s symbols/second (or baud), the data rate is Nf_s bit/second.

For example, with an alphabet consisting of 16 alternative symbols, each symbol represents 4 bits. Thus, the data rate is four times the baud rate.

In the case of PSK, ASK or QAM, where the carrier frequency of the modulated signal is constant, the modulation alphabet is often conveniently represented on a constellation diagram, showing the amplitude of the I signal at the x -axis, and the amplitude of the Q signal at the y -axis, for each symbol.

3.7 Modulator and Detector Principles of Operation

PSK and ASK, and sometimes also FSK, are often generated and detected using the principle of QAM. The I and Q signals can be combined into a complex-valued signal $I + jQ$ (where j is the imaginary unit). The resulting so-called equivalent low-pass signal or equivalent baseband signal is a complex-valued representation of the real-valued modulated physical signal (the so-called passband signal or RF signal). These are the general steps used by the modulator to transmit data.

1. Group the incoming data bits into codewords, one for each symbol that will be transmitted.
2. Map the codewords to attributes, for example amplitudes of the I and Q signals (the equivalent lowpass signal), or frequency or phase values.
3. Adapt pulse shaping or some other filtering to limit the bandwidth and form the spectrum of the equivalent lowpass signal, typically using digital signal processing.
4. Perform Digital-to-Analog Conversion (DAC) of the I and Q signals (since, today, all of the above is normally achieved using Digital Signal Processing DSP).
5. Generate a high-frequency sine wave carrier waveform, and perhaps also a cosine quadrature component. Carry out the modulation, for example by multiplying the sine and cosine wave form with the I and Q signals, resulting in that the equivalent lowpass signal is frequency shifted into a modulated passband signal or RF signal. Sometimes this is achieved using DSP technology, for example direct digital synthesis using a waveform table, instead of analog signal processing. In that case, the above DAC step should be done after this step.
6. Amplification and analog bandpass filtering to avoid harmonic distortion and periodic spectrum

At the receiver side, the demodulator typically performs the following functions.

1. Bandpass filtering.

2. Automatic Gain Control-AGC (to compensate for attenuation, for example fading).
3. Shifting the *RF* signal to the equivalent baseband *I* and *Q* signals, or to an Intermediate Frequency (IF) signal, by multiplying the *RF* signal with a local oscillator sine wave and cosine wave frequency (see the superheterodyne receiver principle).
4. Sampling and Analog-to-Digital Conversion (ADC) Sometimes before or instead of the above point, for example, by means of undersampling.
5. Equalisation filtering; for example a matched filter, compensation for multipath propagation, time spreading, phase distortion and frequency selective fading, to avoid symbol interference and symbol distortion.
6. Detection of the amplitudes of the *I* and *Q* signals, or the frequency or phase of the *IF* signals.
7. Quantisation of the amplitudes, frequencies or phases to the nearest allowed symbol values.
8. Mapping of the quantised amplitudes, frequencies or phases to codewords (bit groups).

9. Parallel-to-serial conversion of the codewords into a bit stream.
10. Pass the resultant bit stream on for further processing such as removal of any error-correcting codes.

As is common to all digital communication systems, the design of both the modulator and demodulator must be done simultaneously. Digital modulation schemes are possible because the transmitter-receiver pair have prior knowledge of how data is encoded and represented in the communication system. In all digital communication systems, both the modulator at the transmitter and the demodulator at the receiver are structured so that they perform inverse operations.

Non-coherent modulation methods do not require a receiver reference clock signal that is phase synchronized with the sender carrier wave. In this case, modulation symbols (rather than bits, characters, or data packets) are asynchronously transferred. The opposite is coherent modulation.

3.8 List of Known Digital Modulation Techniques

The most common digital modulation techniques are listed below.

- Phase-Shift Keying (PSK)-
 - Binary PSK (BPSK), using $M=2$ symbols
 - Quadrature PSK (QPSK), using $M=4$ symbols

- 8PSK, using $M=8$ symbols
 - 16PSK, using $M=16$ symbols
 - Differential PSK (DPSK)
 - Differential QPSK (DQPSK)
 - Offset QPSK (OQPSK)
 - $\pi/4$ -QPSK
- Frequency-Shift Keying (FSK)-
- Audio Frequency-Shift Keying (AFSK)
 - Multi-Frequency Shift Keying (M-ary FSK or MFSK)
 - Dual-Tone Multi-Frequency (DTMF)
 - Continuous-Phase Frequency-Shift Keying (CPFSK)
- Amplitude-Shift Keying (ASK)
- On-Off Keying (OOK), the most common ASK form
- M-ary vestigial sideband modulation, for example 8VSB

- Quadrature amplitude modulation (QAM) - a combination of PSK and ASK-
- Polar modulation like QAM a combination of PSK and ASK.
- Continuous Phase Modulation (CPM) methods-
- Minimum-Shift Keying (MSK)
- Gaussian Minimum-Shift Keying (GMSK)
- Orthogonal Frequency-Division Multiplexing (OFDM) modulation-
- Discrete Multi-Tone (DMT)-including adaptive modulation and bit-loading.
- Wavelet modulation
- Trellis Coded Modulation (TCM), also known as trellis modulation
- Spread-spectrum techniques
- Direct-Sequence Spread Spectrum (DSSS)
- Chirp Spread Spectrum (CSS) according to IEEE 802.15.4a CSS uses pseudo-stochastic coding
- Frequency-Hopping Spread Spectrum (FHSS) applies a special scheme for channel release

SELF-ASSESSMENT EXERCISE

Differentiate between Amplitude-Shift Keying (ASK) and Frequency-Shift Keying (FSK).

4.0 CONCLUSION

In this unit, you have been introduced to the fundamental concepts of modems and modulation. You also learnt about types of modulation; you were also introduced to various methods of modulation. This knowledge will boost your understanding of *ICT* concepts.

5.0 SUMMARY

In this unit, you have learnt that:

- a modem (modulator-demodulator) is a device that modulates an analog carrier signal to encode digital information, and also demodulates such a carrier signal to decode the transmitted information.

- the aim of digital modulation is to transfer a digital bit stream over an analog channel
- common digital modulation techniques are Amplitude-Shift Keying (ASK), Frequency-Shift Keying (FSK) and Phase-Shift Keying (PSK).
-

6.0 TUTOR-MARKED ASSIGNMENT

1. How does the modem perform its function?
2. Explain what you understand by the term “modulation”.
3. Describe the different types of modulation.

7.0 REFERENCES/FURTHER READING

Barry, J. R. E.; Lee, A. & Messerschmidt, D.G. (2004). *Digital Communication*. Kluwer Academic Publishers.

UNIT2 MULTIPLEXERS

CONTENT

- 1.0 Introduction
- 2.0 Objectives
- 3.0 MainContent
 - 3.1 Typesof Multiplexing
 - 3.1.1 SpaceDivisionMultiplexing
 - 3.1.2 Frequency-DivisionMultiplexing
 - 3.1.3 Time-DivisionMultiplexing
 - 3.1.4 Code-DivisionMultiplexing
 - 3.1.5 Wavelength-Division Multiplexing
 - 3.1.6 Synchronous Time-Division Multiplexing
 - 3.1.7 Statistical Time-Division Multiplexing
 - 3.2 Relationto MultipleAccess
 - 3.3 ApplicationAreas
- 4.0 Conclusion
- 5.0 Summary
- 6.0 Tutor-MarkedAssignment
- 7.0 References/FurtherReading

1.0 INTRODUCTION

Inelectronics, a **multiplexer** or **mux** is a device that selects one of several analog or digital input signals and forwards the selected input into a single line. A multiplexer of 2^n input has n select lines, which are

used to select which input line to send to the output. An electronic multiplexer makes it possible for several signals to share one device or resource—for example, one A/D converter or one communication line, instead of having one device per input signal. On the other end, a demultiplexer (or **demux**) is a device taking a single input signal and selecting one of many data-output-lines, which is connected to the single input. A multiplexer is often used with a complementary demultiplexer on the receiving end. An electronic multiplexer can be considered as a multiple-input, single-output switch, and a demultiplexer as a single-input, multiple-output switch.

2.0 OBJECTIVES

At the end of this unit, you should be able to:

- explain the meaning of a multiplexer
- describe the types of multiplexers
- state the areas of application of multiplexers.

3.0 MAINCONTENT

3.1 TypesofMultiplexing

Multiplexing technologies may be divided into several types, all of which have significant variations, namely-Space-Division Multiplexing (SDM), Frequency-Division Multiplexing (FDM), Time-Division Multiplexing (TDM), and Code Division Multiplexing (CDM). Variable bit rate digital bit streams may be transferred efficiently over a fixed bandwidth channel by means of statistical multiplexing, for example packet mode communication. Packet mode communication is an asynchronous mode time-domain multiplexing which resembles time-division multiplexing.

Digital bit streams can be transferred over an analog channel by means of code-division multiplexing (CDM) techniques such as frequency-hopping spread spectrum (FHSS) and direct-sequence spread spectrum (DSSS).

In wireless communications, multiplexing can also be accomplished through alternating polarization (horizontal/vertical or clockwise/counterclockwise) on each adjacent channel and satellite, or through phased multi-antenna array combined with a Multiple-input multiple-output communications (MIMO) scheme.

3.1.1 Space-Division Multiplexing

In wired communication, space-division multiplexing simply implies different point-to-point wires for different channels. Examples include an analogue stereo audio cable, with one pair of wires for the left channel and another for the right channel, and a multi-pair telephone cable. Another example is a switched star network such as the analogue telephone access network (although inside the telephone exchange or between the exchanges, other multiplexing techniques are typically employed) or a switched ethernet network. A third example is a mesh network. Wired space-division multiplexing is typically not considered as multiplexing.

In wireless communication, space-division multiplexing is achieved by multiple antenna elements forming a phased array antenna. Examples are Multiple-Input and Multiple-Output (MIMO), Single-Input and Multiple-Output (SIMO) and Multiple-Input and Single-Output (MISO) multiplexing. For example, a *IEEE 802.11n* wireless router with N antennas makes it possible to communicate with N multiplexed channels, each with a peak bitrate of 54 Mbit/s , thus increasing the total peak bitrate with a factor N . Different antennas will give different

multi-path propagation (echo) signatures, making it possible for digital signal processing techniques to separate different signals from each other. These techniques may also be utilised for space diversity (improved robustness to fading) or beam-forming (improved selectivity) rather than multiplexing.

3.1.2 Frequency-Division Multiplexing

Frequency-division multiplexing (FDM): The spectrum of each input signal is shifted in several distinct frequency ranges.

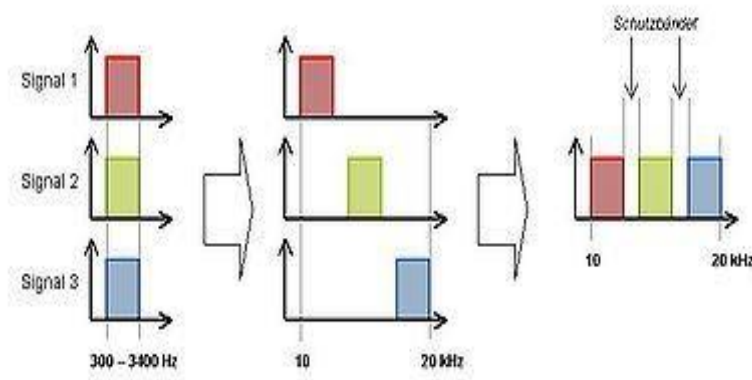


Figure 3.1: Frequency-Division Multiplexing

Frequency-Division Multiplexing (FDM) is inherently an analog technology. FDM achieves the combining of several digital signals into

one medium by sending signals in several distinct frequency ranges over that medium.

One of FDM's most common applications is cable television. Only one cable reaches a customer's home but the service provider can send multiple television channels or signals, simultaneously, over that cable to all subscribers. Receivers must tune to the appropriate frequency (channel) to access the desired signal.

A variant technology, called Wavelength-Division Multiplexing (WDM) is used in optical communications.

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 3 is a conceptual illustration of demultiplexing process.

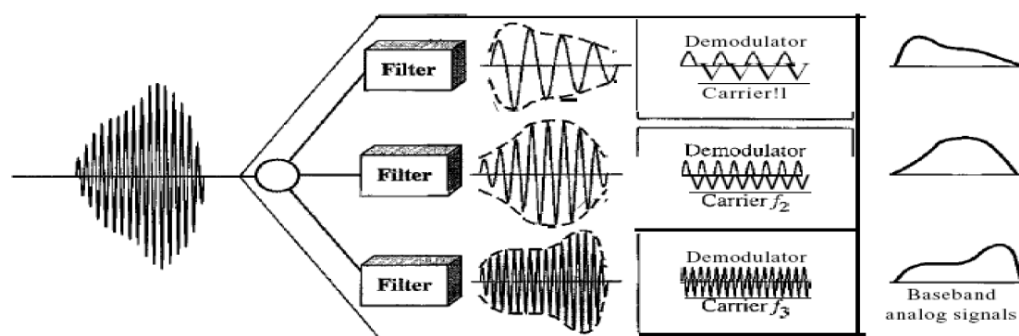


Figure 3.2 FDM de-multiplexing example

3.1.3 Time-Division Multiplexing

Time-Division Multiplexing (TDM) is a digital technology. TDM involves sequencing groups of a few bits or bytes from each individual input stream, one after the other, and in such a way that they can be associated with the appropriate receiver. If done sufficiently and

quickly, the receiving devices will not detect that some of the circuit time was used to serve another logical communication path.

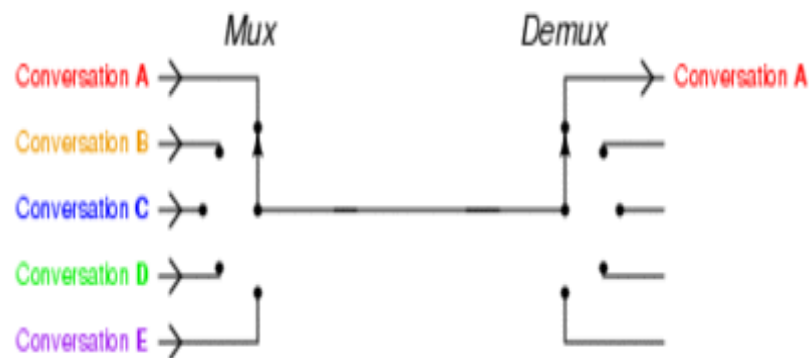


Figure 3.2: Time-Division Multiplexing (TDM)

Consider an application requiring four terminals at an airport to reach a central computer; each terminal communicates at 2400bps , so rather than acquire four individual circuits to carry such a low-speed transmission, the airline has installed a pair of multiplexers. A pair of 9600bps modems and one dedicated analog communications circuit from the airport ticket desk back to the airline data center are also installed.

3.1.4 Code-Division Multiplexing

Code division multiplexing is a technique in which each channel transmits its bits as a coded channel-specific sequence of pulses. This coded transmission typically is accomplished by transmitting a unique time-dependent series of short pulses, which are replaced with a chip

times within the larger bit time. All channels, each with a different code, can be transmitted on the same fiber and asynchronously demultiplexed. Other widely used multiple access techniques are Time Division Multiple Access (TDMA) and Frequency Division Multiple Access (FDMA).

Code Division Multiplexing (CDM) techniques are used as an access technology, namely Code Division Multiple Access (CDMA), in the Universal Mobile Telecommunications System (UMTS) standard for the third generation (3G) mobile communication identified by the ITU. Another important application of the CDMA is the Global Positioning System (GPS).

However, the term *Code Division Multiple Access* is also widely used to refer to a group of specific implementations of CDMA defined by Qualcomm for use in digital cellular telephony, which include *IS-95* and *IS-2000*. The two different uses of this term can be confusing. Actually, CDMA (the Qualcomm standard) and UMTS have been competing for adoption in many markets.

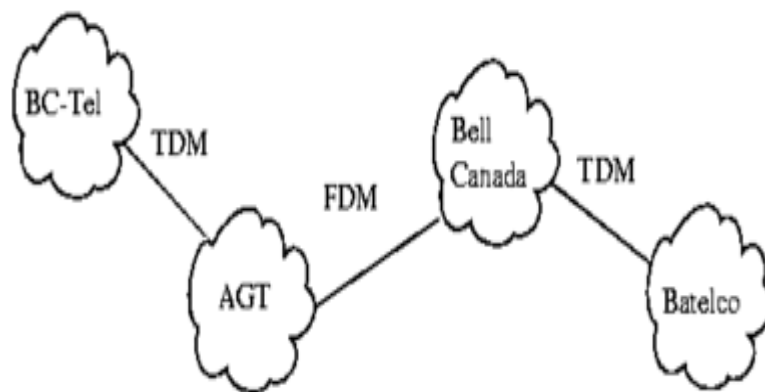


Figure 3.3: Telecommunication Multiplexing

3.1.5 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 3.4 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 3.4 Wavelength-division multiplexing

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 3.5 shows the concept.

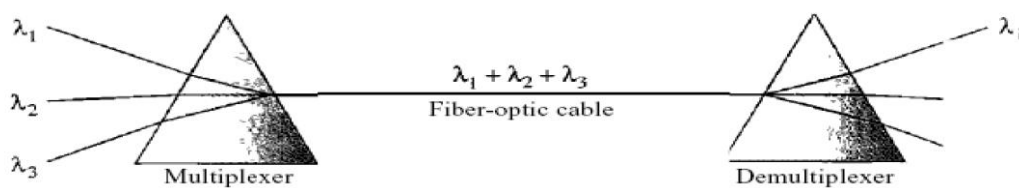


Figure 3.5 Prisms in wavelength-division multiplexing and demultiplexing

3.1.6 Synchronous Time-Division Multiplexing

Time-division multiplexing (TDM) is a digital process that allows several connections to share the high bandwidth of a line. Instead of sharing a portion of the bandwidth as in FDM, time is shared. Each connection occupies a portion of time in the link. Figure 3.6 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.

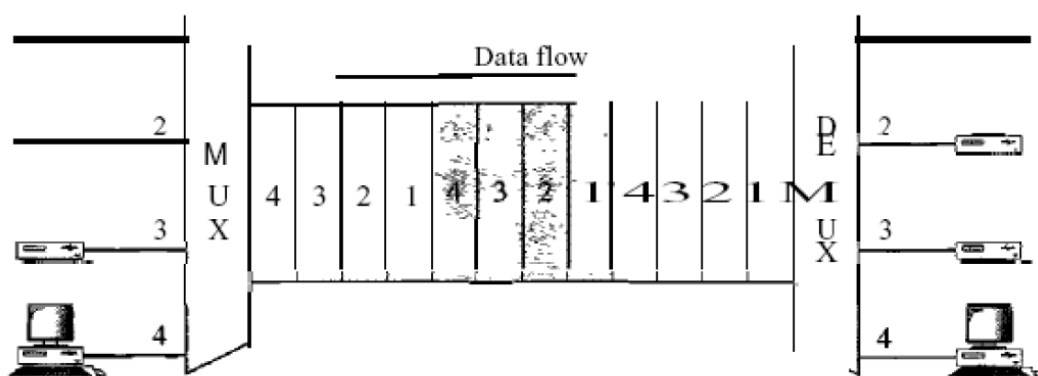


Figure 3.6 TDM

Note that in Figure 3.6 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.

We can divide TDM into two different schemes: synchronous and statistical.

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

Time Slots and Frames

In synchronous TDM, the data flow of each input connection is divided into units, where each input occupies one input time slot. A unit can be 1 bit, one character, or one block of data. Each input unit becomes one output unit and occupies one output time slot. However, the duration of an output time slot is n times shorter than the duration of an input time slot. If an input time slot is T s, the output time slot is T/n s, where n is the number of connections. In other words, a unit in the output connection has a shorter duration; it travels faster. Figure 3.7 shows an example of synchronous TDM where n is 3.

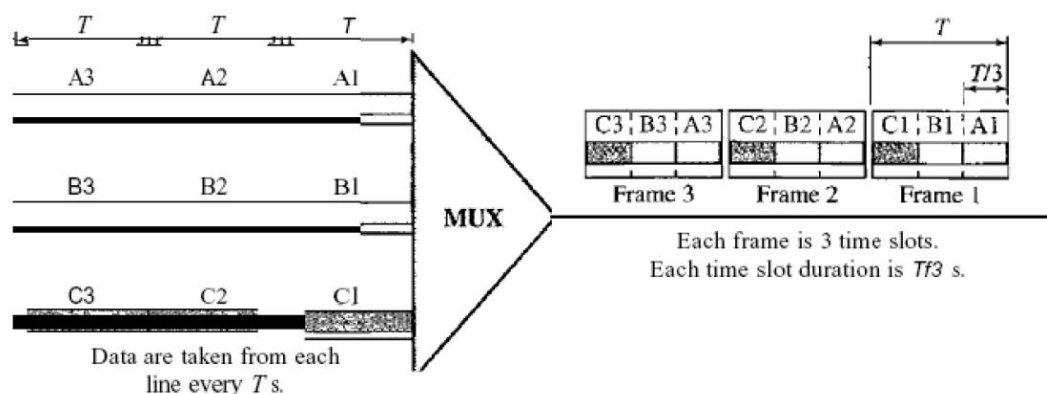


Figure 3.7 Synchronous time-division multiplexing

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 3.7, the data rate of the link is 3 times the data rate of a connection; likewise, the duration of a unit on a connection is 3 times that of the time slot (duration of a unit on the link). In the figure we represent the data prior to multiplexing as 3 times the size of the data after multiplexing. This is just to convey the idea that each unit is 3 times longer in duration before multiplexing than after. Time slots are grouped into frames. A frame consists of one complete cycle of time slots, with one slot dedicated to each sending device. In a system with n input lines, each frame has n slots, with each slot allocated to carrying data from a specific input line.

Interleaving

TDM can be visualized as two fast-rotating switches, one on the 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 has the opportunity to 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 has the opportunity to receive a unit from the path.

Figure 3.8 shows the interleaving process for the connection shown in Figure 3.7.

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

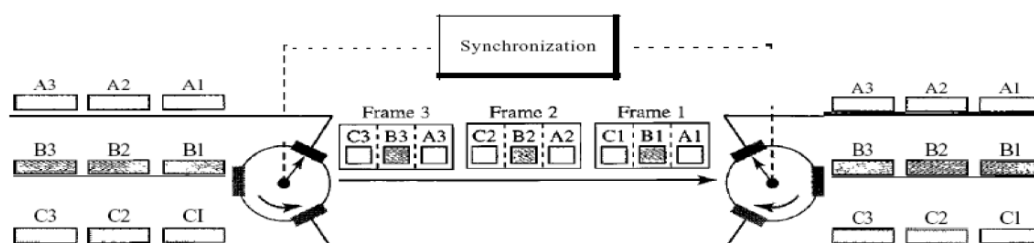


Figure 3.8 Interleaving

3.1.7 Statistical Time-Division Multiplexing

Addressing

Figure 3.9a 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 input 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.

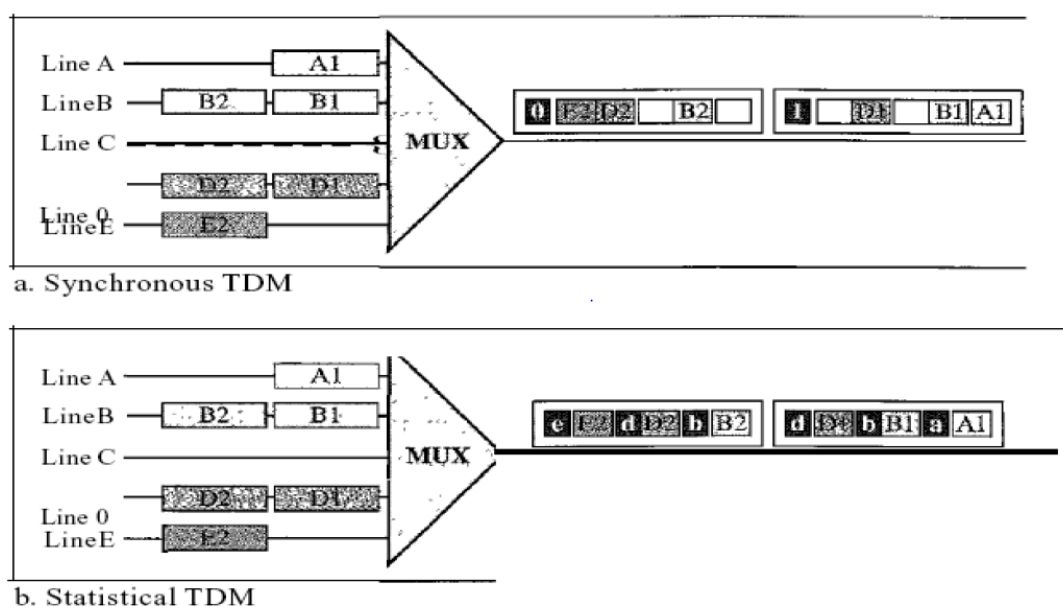


Figure 3.9 a TDM slot comparison

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.

SELF-ASSESSMENT EXERCISE

Find out about spread spectrum multiplexing.

3.2 Relation to multiple access

A multiplexing technique may be further extended into a multiple access method or channel access method, for example *TDM* into Time-Division Multiple Access (TDMA) and statistical multiplexing into Carrier Sense Multiple Access (CSMA). A multiple access method makes it possible for several transmitters connected to the same physical medium to share its capacity.

Multiplexing is provided by the physical layer of the *OSI* model, while multiple access also involves a media access control protocol, which is part of the data link layer. The transport layer in the *OSI* model as well as *TCP/IP* model provides statistical multiplexing of several application layer data flows to/from the same computer.

3.3 Application Areas

Let us consider the following areas.

a. Telegraphy

The earliest communication technology—using electrical wires, and therefore sharing an interest in the economies afforded by multiplexing, was the electric telegraph. Early experiments allowed two separate

messages to travel in opposite directions simultaneously, first using an electric battery at both ends, then at only one end.

- Émile Baudot developed a time-multiplexing system of multiple Hughes machines in the 1870s.
- In 1874, the quadruplex telegraph developed by Thomas Edison transmitted two messages in each direction simultaneously, for a total of four messages transiting the same wire at the same time.
- Several workers were investigating acoustic telegraphy, a frequency-division multiplexing technique, which led to the invention of the telephone.

b. Telephony

In telephony, a customer's telephone line now typically ends at the remote concentrator box down the street, where it is multiplexed along with other telephone lines for that neighborhood or other similar area. The multiplexed signal is then carried to the central switching office on significantly fewer wires and for much further distances than a customer's line can practically go. This is likewise also true for Digital Subscriber Lines (DSL).

Fiber In The Loop (FITL) is a common method of multiplexing, which uses optical fiber as the backbone. It not only connects POTS phone lines with the rest of the PSTN, but also replaces DSL by connecting directly to the Ethernet wire into the home. Asynchronous Transfer Mode is often the communications protocol used.

Since all of the phone (and data) lines have been clumped together, none of them can be accessed except through a demultiplexer. This provides for more-secured communications, though they are not typically encrypted.

The concept is also now used in cable TV, which is increasingly offering the same services as telephone companies. IPTV also depends on multiplexing.

c. Video processing

In video editing and processing systems, multiplexing refers to the process of interleaving audio and video into one coherent *MPEG* transport stream (time-division multiplexing).

In digital video, such a transport stream is normally a feature of a container format which may include metadata and other information, such as subtitles. The audio and video streams may have variable bit

rate. Software that produces such a transport stream and/or container is commonly called a statistical multiplexor or **muxer**. A **demuxer** is a software that extracts or, otherwise, makes available for separate processing the components of such a stream or container.

d. Digital broadcasting

In digital television and digital radio systems, several variable bit-rate data streams are multiplexed together to a fixed bit rate transport stream by means of statistical multiplexing. This makes it possible to transfer several video and audio channels simultaneously over the same frequency channel, together with various services.

In the digital television systems, this may involve several Standard Definition Television (SDTV) programmes (particularly on DVB-T, DVB-S2, ISDB and ATSC-C), or one HDTV, possibly with a single SDTV companion channel over 6 to 8 MHz-wide TV channel. The device that accomplishes this is called a statistical multiplexer. In several of these systems, the multiplexing results in an MPEG transport stream.

The newer DVB standards DVB-S2 and DVB-T2 have the capacity to carry several HDTV channels in one multiplex. Even the original DVB standards can carry more HDTV channels in a multiplex if the most advanced MPEG-4 compression hardware is used.

On communication satellites which carry broadcast television networks and radio networks, this is known as **Multiple Channels Per Carrier** or

MCPC. Where multiplexing is not practical (such as where there are different sources using a single transponder), single channel per carrier mode is used.

Signal multiplexing of satellite TV and radio channels is typically carried out in a central signal play out and uplink centre, such as ASTRA Platform Services in Germany, which provides play out, digital archiving, encryption, and satellite uplinks, as well as multiplexing, for hundreds of digital TV and radio channels.

In digital radio, both the *Eureka 147* system of digital audio broadcasting and the in-band on-channel *HD* Radio, *FMeXtra*, and digital radio *mondiale* systems can multiplex channels. This is essentially required with *DAB*-type transmissions (where a multiplex is called an **ensemble**), but is entirely optional with *IBOC* systems.

e. Analog broadcasting

In *FM* broadcasting and other analog radio media, multiplexing is a term commonly given to the process of adding subcarriers to the audio signal

before it enters the transmitter, where modulation occurs. Multiplexing in this sense is sometimes known as **MPX**, which in turn is also an old term for stereophonic **FM**, seen on stereo systems since the 1960s.

4.0 CONCLUSION

In this unit, you have been introduced to the fundamental concepts of Multiplexers. You also learnt about the types of multiplexers. You were also introduced to the various application areas of multiplexers. All these will go a long way to enrich your understanding of this course as a whole.

5.0 SUMMARY

In this unit, you have learnt that:

- a multiplexer is a device that selects one of several analog or digital input signals and forwards the selected input into a single line
- types of multiplexing are Space-Division Multiplexing (SDM), Frequency-Division Multiplexing (FDM), Time-Division Multiplexing (TDM), and Code-Division Multiplexing (CDM), Wavelength-Division Multiplexing (WDM), Synchronous Time-Division Multiplexing (STDM), Statistical Time-Division Multiplexing (SDTM). Some application areas of multiplexing are telegraphy, telephony, video processing, digital broadcasting, analog broadcasting.

6.0 TUTOR-MARKED ASSIGNMENT

1. How is space-division multiplexing achieved in wireless communication?
2. Describe the implementation of FDM in cable television.
3. How is multiplexing related to OSI model?

7.0 REFERENCES/FURTHER READING

"Federal Standard 1037C: Glossary of Telecommunications Terms".

Institute for Telecommunication Services.

<http://www.its.bldrdoc.gov/fs-1037/fs-1037c.htm>. Retrieved 2009-10-19.

Morris, Mano M. & Charles, Kime, R. (2008). *Logic and Computer Design Fundamentals* (4 Ed.). Prentice Hall.

Voice and Data Communications.

<http://www.perlfect.com/articles/select.shtml>.

Vishnu Universal

Learning <https://www.coursehero.com/file/13176989/Documents-CSECDLNotes2013/>

UNIT3 DIGITAL TECHNOLOGIES

- 1.0 Introduction
- 2.0 Objectives
- 3.0 Main Content
 - 3.1 History
 - 3.2 Properties of Digital Information
- 4.0 Conclusion
- 5.0 Summary
- 6.0 Tutor-Marked Assignment
- 7.0 References/Further Reading

1.0 INTRODUCTION

A **digital** system is a data technology that uses discrete (discontinuous) values. By contrast, non-digital (or analog) systems use a continuous range of values to represent information. Although digital representations are discrete, the information represented can be either discrete, such as numbers, letters or icons, or continuous, such as sounds, images, and other measurements of continuous systems. The word *digital* comes from the same source as the word *digit* and *digitus* (the Latin word for *finger*), as fingers are used for discrete counting. It is most commonly used in computing and electronics, especially where real-world information is converted to binary numeric form as in digital audio and digital photography.

2.0 OBJECTIVES

At the end of this unit, you should be able to:

- explain the meaning of a digital system
- state the properties of digital information.

3.0 MAIN CONTENT

3.1 History

Although, digital signals are generally associated with the binary electronic digital systems used in modern electronics and computing, digital systems are actually ancient, and need not be binary or electronic. Below are some examples of digital signals that are neither binary nor electronic.

- Written text in books (due to the limited character set and the use of discrete symbols - the alphabet in most cases)

- *Anabacus* was created sometime between 1000 BC and 500 BC; it later became a form of calculation frequency. Nowadays it can be used as a very advanced, yet, basic digital calculator that uses beads on rows to represent numbers. Beads only have meaning in discrete up and down states, not in analog in-between states.
- *Abea* con is perhaps the simplest non-electronic digital signal, with just two states (on and off). In particular, *smoke signals* are one of the oldest examples of a digital signal, where an analog "carrier" (smoke) is modulated with a blanket to generate a digital signal (puffs) that conveys information.
- Morse code uses six digital states—dot, dash, intra-character gap (between each dot or dash), short gap (between each letter), medium gap (between words), and long gap (between sentences)—to send messages via a variety of potential carriers such as electricity or light, for example using an electrical telegraph or flashing light.
- The Braille system was the first binary format for character encoding, using a six-bit code rendered as dot patterns.
- Flag semaphore uses rods or flags held in particular positions to send messages to the receiver watching them some distance away.
- International maritime signal flags have distinctive markings that represent letters of the alphabet to allow ships to send messages to each other.
- More recently invented, a modem modulates an analog "carrier" signal (such as sound) to encode binary electrical digital information, as a series of binary digital sound pulses. As slightly

earlier, surprisingly reliable version of the same concept was to bundle a sequence of audio digital "signal" and "no signal" information (i.e. "sound" and "silence") on magnetic cassette tape for use with early home computers.

3.2 Properties of Digital Information

All digital information possesses common properties that distinguish it from analog communications methods. Let us consider these one by one.

- **Synchronisation**- since digital information is conveyed by the sequence in which symbols are ordered, all digital schemes have some method for determining the beginning of a sequence. In written or spoken human language synchronisation is typically provided by pauses (spaces), capitalisation, and punctuation. Machine communications typically use special synchronisation sequences.
- **Language**- all digital communications require *real language*, which in this context consists of all the information that the sender and

receiver of the digital communication must both possess, in advance, in order for the communication to be successful. Languages are generally arbitrary and specify the meaning to be assigned to particular symbol sequences, the allowed range of values, methods to be used for synchronisation, etc.

- **Errors-** disturbances (noise) in analog communications invariably introducesome, generally small deviation or error between the intended and actual communication. Disturbances in a digital communication do not result in errors unless the disturbance is so large as to result in a symbol being misinterpreted as another symbol or disturb the sequence of symbols. It is therefore, generally, possible to have an entirely error-free digital communication. Furthermore, techniques such as check codes may be used to detect errors and guarantee error-free communications through redundancy or retransmission. Errors in digital communications can take the form of *substitution errors* in which a symbol is replaced by another symbol, or *insertion/deletion errors* in which an extra incorrect symbol is inserted into or deleted from a digital message. Uncorrected errors in digital communications have unpredictable and generally large impact on the information content of the communication.

- **Copying-** because of the inevitable presence of noise, making many successive copies of an analog communication is not feasible because each generation increases the noise. Since digital communications are generally error-free, copies of copies can be made indefinitely.

- **Granularity-** when a continuously variable analog value is represented in digital form there is always a decision as to the number of symbols to be assigned to that value. The number of symbols determines the precision or resolution of the resulting datum. The difference between the actual analog value and the digital representation is known as *quantisation error*.

SELF-ASSESSMENT EXERCISE

State other properties of digital information apart from those stated.

4.0 CONCLUSION

In this unit, you have been introduced to the fundamental concepts of digital technologies. You have also learnt the history of digital technologies and some of the properties of digital technologies.

5.0 SUMMARY

In this unit, you have learnt that:

- a digital system is a data technology that uses discrete (discontinuous) values; and that non-digital (or analog) systems use a continuous range of values to represent information.
- digital signals are generally associated with the binary electronic digital systems used in modern electronics and computing
- some properties of digital information are synchronisation, language requirement, low errors, high replication capability and granularity.

6.0 TUTOR-MARKED ASSIGNMENT

1. What are the differences between digital and analog signals?
2. What property of the digital signal makes it better than an analog signal?

7.0 REFERENCES/FURTHER READING

Clark, David D. (1988). "The Design Philosophy of the DARPA Internet Protocols". *Computer Communications Review*, 18(4), pp. 106–114.

"Internet History: People". *Internet History People*.
<http://www.unc.edu/depts/jomc/academics/dri/pioneers2d.html>.
Retrieved July 3, 2006.

"The Risks Digest". *Great moments in e-mail history*.
<http://catless.ncl.ac.uk/Risks/20.25.html#subj3>. Retrieved April
27, 2006.

UNIT4 SIGNAL TRANSMISSION AND IMPAIRMENT

CONTENT

- 1.0 Introduction
- 2.0 Objectives
- 3.0 Main Content
 - 3.1 Types of Impairment
 - 3.1.1 Attenuation
 - 3.1.2 Delay Distortion
 - 3.1.3 Noise
 - 3.2 Noise Types
 - 3.3 Bit Error Rate
 - 3.4 Analog and Digital Transmission
 - 3.4.1 Advantages of Digital Transmission over Analogue
 - 3.5 Signal Transmission Analysis
- 4.0 Conclusion
- 5.0 Summary
- 6.0 Tutor-Marked Assignment
- 7.0 References/Further Reading

1.0 INTRODUCTION

Transmission impairment is a property of a transmission medium which causes the signal to be degraded, reduced in amplitude, distorted or contaminated. Impairment can introduce errors into digital signals. Examples of transmission impairments are attenuation, delay distortion,

and several sources of noise including, thermal noise, impulse noise, and inter-modulation noise. It is important to understand transmission impairments for several reasons. Understanding the source of a transmission impairment like attenuation or dispersion will enable the user to partially correct for (equalize the signal) these effects. Understanding the source of transmission impairments (dispersion, attenuation, impulse noise, and thermal noise) can also help the user understand some of the constraints placed on the transmission of data as a result of these effects. Such constraints include the maximum length of network links, the choice of physical transmission media, the choice of encoding methods, and the data rates supported by the medium.

Attenuation is a property of the transmission medium. It measures how much energy is absorbed and/or radiated from the traveling signal due to its interaction with the transmission medium. Attenuation is measured as a function of the distance traveled through the transmission medium. The transmission medium absorbs energy because the signal is influenced by small impurities within it. Such impurities have different

sizes and distributions depending on the type of medium. Impurities of different sizes effect different frequencies in the signal.

The effect of attenuation is, therefore, a function of frequency. The frequency variation of attenuation can be partially corrected, or equalised, by applying corrections based on a physical model. When a signal is attenuated, its amplitude is reduced. The interpretation of a received signal depends on being able to tell the difference between different signal levels. If the amplitude is reduced too much by attenuation it becomes impossible to accurately tell the difference between the different signal levels, and the information in the signal is lost.

2.0 OBJECTIVES

At the end of this unit, you should be able to:

- define transmission impairment
- describe the types of transmission impairment
- state the effects of noise
- explain the bit error rate
- highlight the differences between analog and digital transmission.

3.0 MAIN CONTENT

3.1 Types of Impairment

The different types of impairment which shall be treated in details in this unit are:

- attenuation
- delay distortion
- noise

Let us now consider this one after the other.

3.1.1 Attenuation

Signal amplitude decreases along a transmission medium. This is known as *signal attenuation*. Amplifiers (in case of analog signals) or repeaters (in case of digital signal) are inserted at intervals along the medium to improve the received signal as close as to its original level. Attenuation and amplification are measured in *decibel (db)*, which is expressed as a constant number of decibels per unit distance.

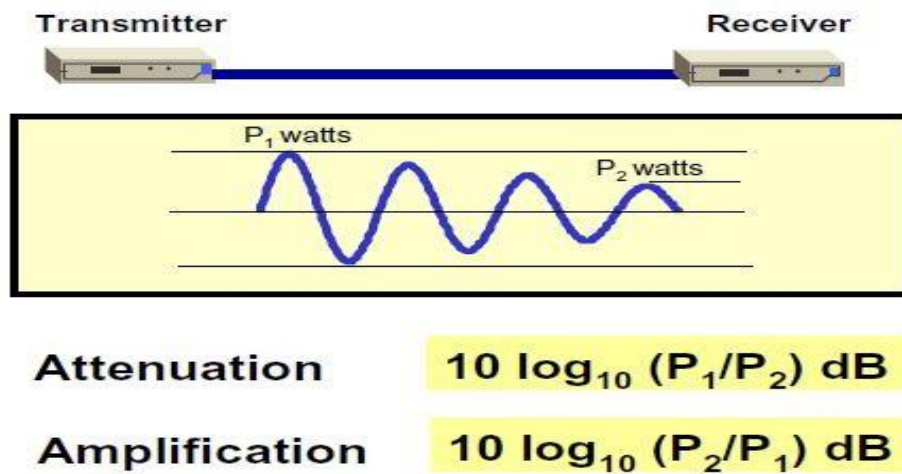


Figure 4.1: Attenuated and Amplified Signal along the Transmission Media

3.1.2 Delay Distortion

The various frequency components in digital signals arrive at the receiver with varying delays, resulting in *delay distortion*. As the bit rate increases, some of the frequency components associated with each bit transition are delayed and start to interfere with frequency components associated with a later bit, causing *inter-symbol interference*, which is a major limitation to maximum bit rate.

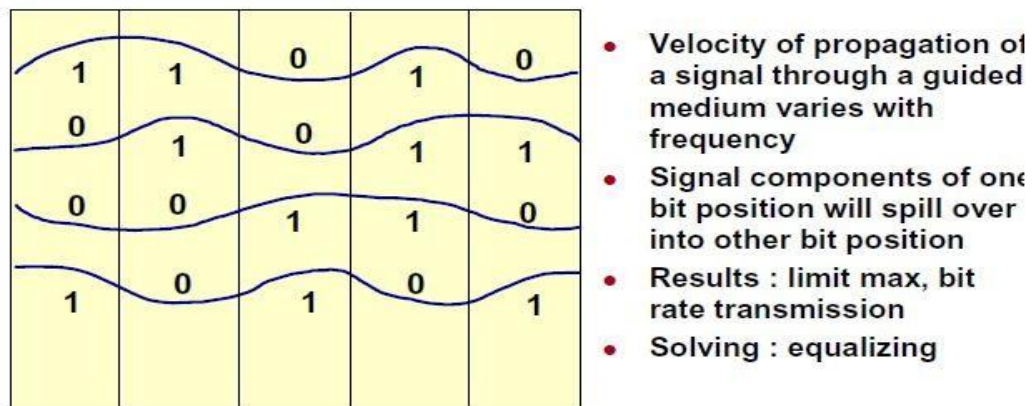


Figure 4.2: Velocity of Propagation of a Signal through Guided Media

3.1.3 Noise

Signal-to-noise ratio (S/N) is a parameter used to quantify how much noise there is in a signal. A high SNR means a high power signal relative to noise level, resulting in a good-quality signal.

S/N is given as

$$S/N_{db} = 10 \log_{10} S/N$$

Where S = average signal power

N = noise power

3.2 Noise Types

Here, we are going to consider some common examples.

Atmospheric noise

- Lightning-static discharge of clouds
- Solar noise-sun's ionised gases
- Cosmic noise-distant stars radiate high frequency signal

Gaussian noise

Thermal noise-generated by random motion of free electrons

Impulse noise-sudden bursts of irregular pulses

Crosstalk-this is interference generated when magnetic fields or current nearby wires interrupt electrical current in a wire. As electrical current travels through a wire, the current generates a magnetic field. Magnetic fields from wires that are closed together can interfere with each other.

Shielding the wire and twisting wire pairs around each other help decrease crosstalk

NEXT (near-end crosstalk)

- interference in a wire at the transmitting end of a signal sent on a different wire NEXT (far-end crosstalk)
- interference in a wire at the receiving end of a signal sent on a different wire

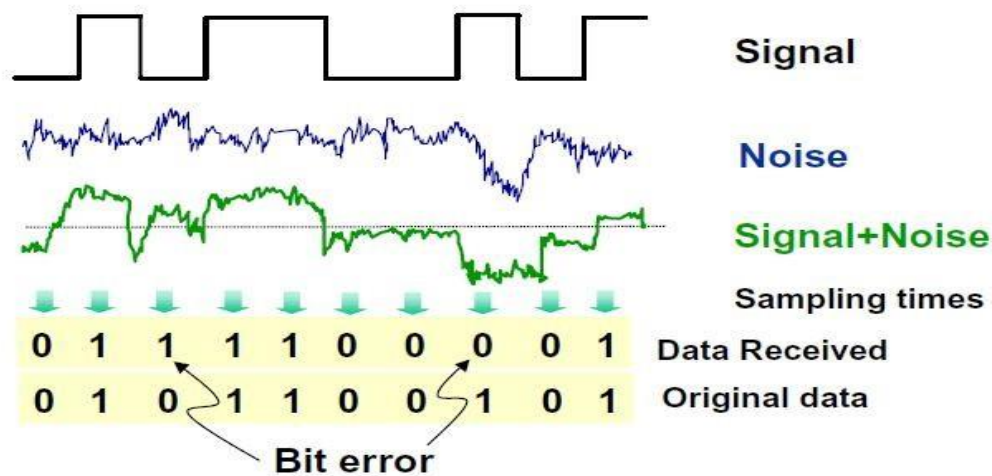


Figure 4.3: Effect of Noise

Impulse noise is the primary source of error for digital data. A sharp spike of energy of 0.01 seconds duration will not destroy any voiced data, but will wash out many bits of digital data.

3.3 Bit Error Rate

- BER (Bit Error Rate) is the probability of a single bit being corrupted in a defined time interval.
- BER of 10^{-5} means on average 1 bit in 10^5 will be corrupted
- A BER of 10^{-5} over voice-grade line is typical.
- BER of less than 10^{-6} over digital communication is common.

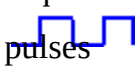
A Bit Error Rate (BER) is a significant measure of system performance in terms of noise. A BER of 10^{-6} , for example, means that one bit of every million may be destroyed during transmission. Several factors that affect BER are listed below.

- Bandwidth
- S/N
- Transmission medium
- Transmission distance
- Environment
- Performance of transmitter and receiver

3.4 Analog and Digital Transmission

Analog signal is characterised by being continuously variable along amplitude and frequency. Digital signal is characterised by series of discrete pulses, representing one and zero bits. Table 4.1 presents the comparison between analog and digital transmission.

Table 4.1: Analog and Digital Transmissions

	Analog	Digital
Data	continuous (e.g., voice)	discrete (e.g., text)
Signal	continuous electromagnetic waves 	sequence of voltage pulses 
transmission	Used mainly for transmitting Transmission of analog signals without regards to their content (the data may be analog or binary). The signals become weaker (attenuated) with the distance. Amplifiers may be used to strengthen the signals, but as side effect they also boost	Used mainly internally Transmission that is concerned with the content of the signal. Repeaters are used to overcome attenuation. A repeater recovers the digital pattern from the signal it gets, and

3.4.1 Advantages of Digital Transmission over Analog

- **Technology**-there is a drop in cost due to *LSI* and *VLSI*
- **Data integrity**-repeater allows longer distances over lines of lesser quality.
- **Capacity utilisation**-digital techniques can be more easily and cheaply utilised, through multiplexing, available transmission links of high bandwidth.
- **Security and privacy**-encryption techniques are more readily applied to digital data
- **Integration**-simplified, if digitised data is used everywhere.

SELF-ASSESSMENT EXERCISE

State the advantages of analog transmission over digital.

3.5 Signal Transmission Analysis

The analysis of electrical signals is a fundamental problem for many engineers and scientists. Even, if the immediate problem is not electrical, the basic parameters of interest are often changed into electrical signals by means of transducers. Common transducers include accelerometers and load cells in mechanical work, EEG electrodes and

blood pressure probes in biology and medicine, and pH and conductivity probes in chemistry. The rewards for transforming physical parameters to electrical signals are great, as many instruments are available for the analysis of electrical signals in the time, frequency and modal domains. The powerful measurement and analysis capabilities of these instruments can lead to rapid understanding of the system under study.

4.0 CONCLUSION

In this unit, you have been introduced to the fundamental concepts of transmission impairment. You also learnt the types of transmission impairments. You were also introduced to bit error rate and some of the major differences between analog and digital transmissions.

5.0 SUMMARY

In this unit, you have learnt that:

- Transmission impairment is a property of a transmission medium which causes the signal to be degraded, reduced in amplitude, distorted or contaminated.
- Examples of transmission impairments are attenuation, delay distortion, and noise
- Types of noise include thermal noise, impulse noise, and inter-modulation noise

- Analog signal is characterised by being continuously variable along amplitude and frequency. Digital signal is characterised by series of discrete pulses, representing one and zero bits.

6.0 TUTOR-MARKED ASSIGNMENT

- a. What do you understand by transmission impairment?
- b. What is the difference between *FEXT* and *NEXT*?
- c. What do you understand by the term “noise”?

7.0 REFERENCES/FURTHER READING

Friedhelm, Hillebrand (2002). *GSM and UMTS - The Creation of Global Mobile Communications* (Ed.). John Wiley & Sons.

"Brazil, Russia, India and China to Lead Internet Growth Through 2011". Clickz. com.
<http://clickz.com/showPage.html?page=3626274>. Retrieved 2009-05-28.

MODULE5 NETWORKTECHNOLOGIES

Unit1	IntegratedServiceDigitalNetwork(ISDN)
Unit2	Digital SubscriberLine(DSL)
Unit3	SONET
Unit4	Packet Switching
Unit5	InternetandTCP/IP

UNIT1 INTEGRATEDSERVICEDIGITALNETWORK (ISDN)

CONTENTS

1.0	Introduction
2.0	Objectives
3.0	MainCONTENTS
3.1	WhatisISDN?
3.2	History
3.3	ComponentsofISDN
	3.3.1ISDNAccessTypes
3.4	ISDNDevices
3.5	ISDNInterfaces
	3.5.1 TheR Interface
	3.5.2 TheSInterface
	3.5.3 TheTInterface
	3.5.4 TheU interface
3.6	HowDoesISDNwork?

3.6 The Usefulness of *ISDN*

4.0 Conclusion

5.0 Summary

6.0 Tutor-Marked Assignment

7.0 References/Further Reading

1.0 INTRODUCTION

Having read through the course guide, you will have a general understanding of what this unit is about and how it fits into the course as a whole. This unit will describe the general concept of *ISDN*.

2.0 OBJECTIVES

At the end of this unit, you should be able to:

- explain the term *ISDN*
- describe *ISDN* interfaces
- discuss how *ISDN* works
- state the usefulness of *ISDN*

3.0 MAIN CONTENT

3.1 What is Integrated Services Digital Network (ISDN)?

Integrated Services Digital Network (*ISDN*) is a digital communications technology that enables a small business or an individual to connect directly to the internet and other sites/users (for video conferencing, for instance). *ISDN* provides a standard interface for voice, fax, video, graphics, and data—all on a single telephone line.

“**Integrated services**” refers to *ISDN*'s ability to deliver two simultaneous connections, in any combination of voice, fax, data, and video, over a single line. Multiple devices can be attached to the line, and used as needed.

“**Digital**” refers to the fact that it is a purely digital transmission, as opposed to the analog transmission method used by conventional telephoned lines.

“**Network**” refers to the fact that *ISDN* is not simply a point-to-point connection like a leased telephoned line – *ISDN* networks extend from the local telephone exchange to the remote user, and include all the switching equipment in between. If your *ISDN* equipment includes analog capabilities, you can also connect to telephones, fax machines, and analog modems – even though they may be connected to standard analog telephoned lines.

ISDN service is provided by the same companies that provide telephone service – you get much faster, more dependable connections for voice, fax, data, and video – all through a single connection. While not new (*ISDN* has been in use for over 15 years), the advent of international standards has made *ISDN* viable as telephone companies around the world have upgraded their equipment to these *ISDN* standards. It is now commonly available in Europe, Japan, Australia, and from most major North American telephone companies – *AT&T*, *MCI*, and *Sprint* can provide long-distance *ISDN* lines for global connections. One of the

reasons for its widespread use is that it works on the ordinary copper wire already in place in the telephone system.

One advantage of *ISDN* over other digital communication technologies is its ability to handle all types of information such as voice, computer data, studio-quality sound, and video.

3.2 History

The early phone network consisted of a pure analog system that connected telephone users directly by a mechanical interconnection of wires. This system was very inefficient, was very prone to breakdown and noise, and did not lend itself easily to long-distance connections. Beginning in the 1960s, the telephone system gradually began converting its internal connections to a packet-based, digital switching system. Today, nearly all voice switching in the United States (U.S.) is digital within the telephone network. Still, the final connection from the local central office to the customer equipment was, and still largely is, an analog Plain-Old Telephone Service (POTS) line.

A standards movement was started by the International Telephone and Telegraph Consultative Committee (CCITT), now known as the International Telecommunications Union (ITU). The ITU is a United Nations organisation that coordinates and standardises international telecommunications. Original recommendations of *ISDN* were in *CCITT*

Recommendation I.120(1984) which described some initial guidelines for implementing *ISDN*.

Local phone networks, especially the regional Bell operating companies, have long hailed the system, but they had been criticised for being slow to implement *ISDN*. One good reason for the delay is the fact that the two major switch manufacturers, Northern Telecom (now known as Nortel Networks), and AT&T (whose switch business is now owned by Lucent Technologies), selected different ways to implement the *CCITT* standards. These standards didn't always inter-operate. This situation has been likened to that of earlier 19th century railroading - "people had different gauges, different tracks... nothing worked well".

In the early 1990s, an 'industry-wide' effort was initiated to establish a specific implementation for *ISDN* in the U.S. Members of the industry agreed to create the National *ISDN* 1 (**NI-1**) standard so that end users would not have to know the brand of switch they are connected to in order to buy equipment and software compatible with it. However, there were problems agreeing on this standard. In fact, many western states would not implement *NI-1*. Both Southwestern Bell and U.S. West (now Qwest) said that they did not plan to deploy *NI-1* software in their

central office switches due to incompatibilities with their existing *ISDN* networks.

Ultimately, all the Regional Bell Operating Companies (RBOCs) did support *NI-1*. A more comprehensive standardisation initiative, National *ISDN2 (NI-2)*, was later adopted. Some manufacturers of *ISDN* communication equipment—such as Motorola and U.S. Robotics (now owned by 3Com), worked with the RBOCs to develop configuration standards for their equipment. These kinds of actions, along with more competitive pricing, inexpensive *ISDN* connection equipment, and the desire for people to have relatively low-cost high-bandwidth internet access have made *ISDN* more popular in recent years.

Most recently, *ISDN* service has largely been displaced by broadband internet service, such as *xDSL* and cable modem service. These services are faster, less expensive, and easier to set up and maintain than *ISDN*. Still, *ISDN* has its place, as a backup to dedicated lines, and in locations where broadband service is not yet available.

3.3 Components of *ISDN*

While individual operating companies and ministries will define the specific services, within the *ISDN* architecture the *ITU* standards define a number of component parts and functions as listed below.

- ISDN Channels

- Access Types

-

Devices

-

Interfaces

-

Protocols

ISDN Channels

A CHANNEL is the basic unit of ISDN service. The ISDN Standards define three basic types of channels:

- Bearer channels (B channels)
- Delta (or "Demand") channels (D channels)
- High-capacity channels (H channels)

B Channel

A B channel is a 64-Kbps unit of clear digital bandwidth. Based on the data rate required to carry one digital voice conversation, a B channel can carry any type of digital information (voice, data, or video) with no restrictions on format or protocol imposed by the ISDN carrier.

D Channel

A *D channel* is a signaling channel. It carries the information needed to connect or disconnect calls and to negotiate special calling parameters (i.e., automatic number *ID*, call waiting, data protocol). The *D channel* can also carry packet-switched data using the *X.25 protocol*. The *D channel* is not a clear channel. It operates according to a well-defined pair of layered protocols:

- Q.921 (LAPD) at the data link layer (Layer 2)
- Q.931 at the upper layers (Layers 3 and above)

The data rate of a *D channel* varies according to the type of access it serves—a basic rate access *D channel* operates at 16 Kbps and a primary rate access *D channel* operates at 64 Kbps.

Signaling on the D Channel

The *ISDN D channel* carries all signaling between the customer's terminal device and the carrier's switching office.

Signaling information with end-to-end significance (i.e., which must be received by the terminal device at a call's destination, such as Automatic Calling Number Identification information) travels between the carrier's switching offices on the carrier's common-channel signaling network and onto the destination terminal, through the receiving user's *D channel*.

H Channel

An *H* channel is a special, high-speed clear channel. *H* channels, designed primarily for full-motion color video, are not yet in common use. There are currently three kinds of *H* channel:

- *H*0("H-zero")
- *H*11("H-one-one")
- *H*12("H-one-two")

An *H*0 channel operates at 384 Kbps (roughly one-fourth of a North American primary rate access or one-fifth of a European primary rate access). An *H*1 channel operates at 1.536 Mbps and occupies one whole North American primary rate access. An *H*12 channel occupies an entire European primary rate access.

3.3.1 ISDN Access Types

ISDN offers two general types of access:

- BASIC RATE ACCESS (BRA)
- PRIMARY RATE ACCESS (PRA)

These differ from one another by the amount of information they can carry.

Basic rate access

Basic rate access is based on new technology conceived especially for *ISDN*. Designed to provide service to individual users or small businesses, basic rate access provides two 64-Kbps *B* channels and one 16-Kbps *D* channel (referred to as 2B+D). In other words, it provides transmission facilities for one voice conversation (one *B* channel), one medium-speed data session (the other *B* channel), and the signaling exchanges needed to make them work (the *D* channel).

Two *B* channels at 64 Kbps plus one *D* channel at 16 Kbps equals 144 Kbps. The *ISDN* basic rate transmission protocol uses an additional 48 Kbps of bandwidth for maintenance and synchronisation, so an *ISDN* basic rate access actually uses 192 Kbps.

Primary rate access

Primary rate access, which is based on pre-*ISDN* digital carrier technology, is designed to provide high-capacity service to large customers for applications such as PBX-to-PBX *trunking*. There are two kinds of primary rate access: 23B+D and 30B+D. Each depends on the kind of digital carrier available in a given country.

In North America and Japan, 23B+D primary rate access operates at 1.544 Mbps and offers 23 *B channels* plus 1 64-Kbps *D channel* (usually located in time-slot 23), or 4 *H0* channels, or 1 *H11* channel. In most of the rest of the world, 30B+D primary rate access operates at 2.048 Mbps and offers 30 *B channels* plus 1 64-Kbps *D channel* (located in time-slot 16), or 5 *H0* channels, or 1 *H12* channel.

3.4 *ISDN* Devices

In the context of *ISDN* standards, standard devices refer not to actual hardware, but to standard collections of functions that can usually be performed by individual hardware units. The *ISDN* standard devices are listed below.

- Terminal Equipment (TE)
- Terminal Adapter (TA)
- Network Termination 1 (NT1)
- Network Termination 2 (NT2)
- Exchange Termination (ET)

Terminal Equipment (TE)

A *TE* is any piece of communicating equipment that complies with the *ISDN* standards. Examples included digital telephones, *ISDN* data terminals, Group IV Fax machines, and *ISDN*-equipped computers.

In most cases, a *TE* should be able to provide a full basic rate access (2B+D), although some *TE*s may use only 1B+D or even only a *D channel*.

Terminal Adapter (TA)

A *TA* is a special interface-conversion device that allows communicating devices that don't conform to *ISDN* standards to communicate over the *ISDN*.

The most common *TAs* provide basic rate access and have one *RJ*-type modular jack for voice and one *RS-232* or *V.35* connector for data (with each portable to connect to either of the available *B* channels). Some *TA*s have a separate data connector for the *D channel*.

Network Termination (NT1 and NT2)

The NT devices, NT1 and NT2, form the physical and logical boundary between the customer's premises and the carrier's network. NT1 performs the logical interface functions of switching and local-device control (local signaling). NT2 performs the physical interface conversion between the dissimilar customer and network sides of the interface.

In most cases, a single device, such as a PBX or digital multiplexer, performs both physical and logical interface functions. In ISDN terms, such a device is called NT12 ("NT-one-two") or simply NT.

Exchange Termination (ET)

The ET forms the physical and logical boundary between the digital local loop and the carrier's switching office. It performs the same functions at the end of the office that the NT performs at the customer's premises.

In addition, the *ET*:

- separates *B channels*, placing them on the proper interoffice trunk to their ultimate destinations
- terminates the signaling path of the customer's *D channel*, converting any necessary end-to-end signaling from the *ISDN D-channel* signaling protocol to the carrier's switch-to-switch trunk signaling protocol

3.5 ISDN Interfaces (Standard Reference Points)

The *ISDN* standards specify four distinct interfaces in the customer's connection to the network: *R*, *S*, *T*, and *U*, as discussed in subsection 3.5.1 through 3.5.4

From the standards viewpoint, these are not "real" physical interfaces, but simply 'standard reference points' where physical interfaces may be necessary. However, in common practice, the names of reference points are used to refer to physical interfaces.

3.5.1 The *R* Interface

The interface at reference point *R* is the physical and logical interface between a non-*ISDN* terminal device and a Terminal Adapter (TA). The *R* interface is not really part of the *ISDN*; it can conform to any of the common telephone or data interface standards.

3.5.2 The S Interface

The interface at reference point *S* is the physical and logical interface between a *TE* (or *TA*) and an *NT*. The *S* interface uses four wires and employs a bipolar transmission technique known as Alternate Mark Inversion (AMI).

A special feature of the *S* interface is the "short passive bus" configuration, which allows up to eight *ISDN* devices (*TE* or *TA*) to contend for packet access to the *D* channel in a prioritised, round-robin fashion. Only one device at a time can use a given *B* channel.

3.5.3 The T Interface

The interface at reference point *T* is the physical and logical interface between *NT1* and *NT2*, whenever the two *NT*s are implemented as separate pieces of hardware. The specification for the *T* interface is identical to the specification for the *S* interface.

In most implementations, *NT1* and *NT2* exist in the same physical device, so there is no real *T* interface.

3.5.4 The *U* Interface

The interface at reference point *U* is the physical and logical interface between *NT* (or *NT2*) and the *ISDN* carrier's local transmission loop. It is also the legal demarcation between the carrier's loop and the customer's premises.

The *U* interface is implemented with two wires and uses a special quaternary signal format (i.e., four possible electrical states, with one pulse encoding a predefined combination of 2 bits) called 2B1Q. Quaternary encoding allows the *U* interface to carry data with a logical bitrate of 192 Kbps over a signal with a physical pulse rate of only 96 Kbps. The slower pulse rate is better suited to the less-predictable environment of the outside-plant loop carrier system.

3.6 How Does *ISDN* Work?

The simplest *ISDN* connection (called Basic Rate or BRI) consists of two 64 Kbps (kilobits-per-second) data channels (called *B-channels*) plus a 16 Kbps control channel (called the *D-channel*). This is sometimes referred to as "2B+D." On the other end of the spectrum is primary rate *ISDN* (called PRI) with 23 *B-channels* plus a *D-channel* (i.e.: "23B+D").

To connect to the *ISDN* line, you need a black-box called an *NT1* Network Terminator—a power supply (which you also need) is often built-in. You will also need a Terminal Adapter (often called a "TA") to connect non-*ISDN* equipment (such as your computer or fax machine) to

the line—these are also available as plug-in cards for PC's. Some TA's work as ethernet bridges so that you can connect your LAN directly to the ISDN line.

3.7 The Usefulness of ISDN

One of the most common uses for this technology today is videoconferencing. By using from one to four BRI lines, a videoconference can be established between two or more sites—the more lines, the faster the connection. For a videoconference application, higher connection speed translates to higher resolution and video frame rates. The telephone company infrastructure allows these connections to be made in a similar fashion to dialing a telephone. While videoconferencing has been around for a long time, in the past it has primarily been confined to large corporations. The ability to transmit quality voice and video over long distances used to require expensive equipment and costly leased lines—these could only be justified by the largest of companies. Due to this dependency on leased lines, videoconferences were point-to-point (e.g. headquarters may be permanently linked to a

manufacturing plant). Videoconferencing on the scale of teleconferencing was simply impractical.

The advent of new low-cost videoconference hardware that can utilise *ISDN* is rapidly changing this. Both desktop conferencing (a participant uses a PC equipped with a microphone, a small video camera, and an *ISDN* interface) and true videoconferencing (where more sophisticated equipment and remote control cameras allow group participation) have become as easy to set up as voice conferencing. Due to *ISDN*'s versatility, videoconferences can include the sharing of graphic images and presentations, computer applications, documents, and computer files. This capability is proving popular for telecommuting, long distance meetings, workgroup collaboration, security and surveillance, and dozens of other innovative applications.

4.0 CONCLUSION

In this unit, you have been introduced to the fundamental concepts of *ISDN*. You have also learnt the different components and interfaces of *ISDN*.

5.0 SUMMARY

In this unit, you have learnt about:

- the term *ISDN*
- *ISDN* interfaces
- how *ISDN* works
- the usefulness of *ISDN*.

6.0 TUTOR-MARKED ASSIGNMENT

- a. According to ITU, describe ISDN architecture. b. Describe ISDN access types.

7.0 REFERENCES/FURTHER READING

Friedhelm, Hillebrand (2002). *GSM and UMTS: The Creation of Global Mobile Communications* (Ed.). John Wiley & Sons.

Sodiya, A.S. (2008). *Digital Communication and Computer Network – An Introduction* (A handbook).

"Brazil, Russia, India and China to Lead Internet Growth Through 2011". Clickz.com.
<http://clickz.com/showPage.html?page=3626274>. Retrieved 2009-05-28.

UNIT2 DIGITALSUBSCRIBERLINE(DSL)

CONTENTS

- 1.0 Introduction
- 2.0 Objectives
- 3.0 MainContent
 - 3.1 Whatis Digital SubscriberLine(*DSL*)?
 - 3.2 Historyof*DSL*
 - 3.3 Featuresof*DSL*
 - 3.4 Benefitsof*DSL*
 - 3.5 How*DSL*works
- 4.0 Conclusion
- 5.0 Summary
- 6.0 Tutor-MarkedAssignment
- 7.0 References/FurtherReading

1.0 INTRODUCTION

DSL is a specialised and problem-oriented language. Contrary to a GeneralPurpose Language(GPL) (e.g.*UML*,*Java*or*C#*),a*DSL*serves to accuratelydescribeadomainofknowledge.Theinteresttocombinea *DSL*andatransformation functionistoraisetheabstractionlevelof software.A*DSL*Userconcentratesher/his efforts ondomaindescription whilecomplexity, designandimplementation decisionsanddetailsare hidden.Thestakeis to improveproductivityandsoftwarequality.

However, what is the next consensus beyond this general definition? An experience consists in starting the development of a *DSL* editor coupled to a generator. Quickly, the issue of the variants of *DSL* editors and generator emerges. Regarding the language, is it a tree-based *DSL* or a set of data without real structure? Is it a graphical or textual notation? Is it a declarative or imperative style? Answers to most of these questions will be provided in this unit.

2.0 OBJECTIVES

At the end of this unit, you should be able to:

- explain the term *DSL*
- state the features of *DSL*
- illustrate how *DSL* works
- highlight the advantages and disadvantages of *DSL*

3.0 MAIN CONTENT

3.1 What is DSL?

DSL is a telephone loop technology that uses existing copper phone lines, and provides a dedicated, high-speed internet connection. The big advantage of some *DSLs* (notably *ADSL*), is that they can co-exist on the same line with traditional voice services such as "POTS" (Plain Old Telephone Service), and even *ISDN*. This is accomplished by utilising different frequency ranges above the voice range (voice is up to 4KHz). Essentially, this gives two lines in one - one for voice, and one for internet connectivity. When all is working normally, there should be no interference between the two "lines". This gives *DSL* a potentially broad consumer base, and helps minimise costs for service providers.

DSL is positioned for the Home and Small Office (SOHO) market that is looking for high-speed internet access at reasonable prices. Since it also typically provides dedicated, "always on" access, it can be used for interconnecting low to mid-range bandwidth servers, and provides a great access solution for small LANs. It is also great for those Linux power users that just want a fat pipe).

Phone companies, and other independent telecommunications providers (CLECs), are now deploying *DSL* to stay ahead of the cable companies - the main consumer and SOHO competition for *DSL* providers. This mad rush to get "a piece of the pie" is bringing much competition (a good

thing!), much diversity, and some confusion, into the consumer market. The DSL provider (often, but not always, the phone company) will provide the DSL infrastructure. This will include your line, the DSLAM, and physical connection to the outside world. From there, it is typically picked up by an ISP, who provides the traditional internet services.

Consumer DSL plans are typically "best effort" services. While boasting speeds approaching T1, and even surpassing that in some cases, it is not necessarily as reliable as T1, however. Business class DSL offers more reliability at a higher cost than consumer plans, and is a good compromise where both reliability and bandwidth are at a premium. All in all, the cost of DSL compared to traditional telecom services, such as T1, is attractive and substantially more affordable for home and small business users.

DSL providers often do not have service contracts for home users, while business class DSL services typically do include similar SLAs (Service Level Agreements) to that offered for a T1 line.

The downside is that *DSL* is not available everywhere. Availability, and available bitrate (speed), are purely a function of where you live, where the telecom firm has installed the prerequisite hardware, how far you are from the *DSLAM/CO*, and the quality of your phone line (loop). Not all loops are created equal, unfortunately. The primary limitation is distance.

3.2 History of DSL

Implementation of digital subscriber line technology originally was part of the Integrated Services Digital Network (ISDN) specification published in 1984 by the CCITT and ITU as part of recommendation *I.120*, later reused as ISDN Digital Subscriber Line (IDSL). Engineers have developed higher-speed *DSL* facilities such as High (bit rate) Digital Subscriber Line (HDSL) and Symmetric Digital Subscriber Line (SDSL) to provide traditional Digital Signal 1 (DS1) services over standard copper pair facilities. Consumer-oriented Asymmetric Digital Subscriber Line (ADSL), first tested at Bellcore in 1988, was designed to operate on existing lines already conditioned for *BRI* ISDN services. This itself is as switched digital service (non-*IP*), though most Incumbent Local Exchange Carriers (ILECs) provide Rate-Adaptive Digital Subscriber Line (RADSL) to work on virtually any available copper pair facility—whether conditioned for *BRI* or not.

The development of *DSL*, like many other forms of communication, can be traced back to Claude Shannon's 1948 seminal paper—'A Mathematical Theory of Communication'. Employees at Bellcore (now

Telcordia Technologies) developed ADSL in 1988 by placing wide-band digital signals above the existing (baseband) analog voice signal carried between Telephone Company (central offices) and customers on conventional (twisted pair) cabling facilities. ADSL circuit provides "digital service". The underlying technology of transport across DSL facilities uses high-frequency (sinusoidal-carrier wave) modulation, which is an analog signal transmission.

ADSL circuit terminates at each end in a modem which modulates patterns of Binary digit (bits) into certain high-frequency impulses for transmission to the opposing modem. Signals received from the far-end modem are demodulated to yield a corresponding bit pattern that the modem retransmits, in digital form, to its interfaced equipment, such as a computer, router, switch, etc. Unlike traditional dial-up modems, which modulate bits into signals in the 300–3400 Hz baseband (voice service), DSL modems modulate frequencies from 4000 Hz to as high as 4 MHz. This frequency band separation enables DSL service and plain old telephone service (POTS) to coexist on the same copper pair facility. Generally, higher bit rate transmissions

require a wider frequency band, though the ratio of bit rate to bandwidth are not linear due to significant innovations in digital signal processing.

Early *DSL* service required a dedicated dry loop, but when the U.S. Federal Communications Commission (FCC) required *ILEC* to lease their lines to competing *DSL* service providers, shared-line *DSL* became available. Also known as *DSL* over (Unbundled Network Element), this unbundling of services allows a single subscriber to receive two separate services from two separate providers on one cable pair. The *DSL* service provider's equipment is co-located in the same central office as that of the *ILEC* supplying the customer's pre-existing voice service. The subscriber's circuit is then rewired to interface with hardware supplied by the *ILEC* which combines a *DSL* frequency and *POTS* frequency on a single copper pair facility.

On the subscriber's end of the circuit, in-line low-pass *DSL* filters (splitters) are installed on each telephone to filter the high-frequency "hiss" that would otherwise be heard. Conversely, High-pass filters already incorporated in the circuitry of *DSL* modems filter out voice frequencies. Although *ADSL* and *RADSL* modulation do not use the voice-frequency band, nonlinear elements in the phone could otherwise generate audible inter-modulation and may impair the operation of the data modem in the absence of low-pass filters.

3.3 Features of *DSL*

Let us look at these in relation to the following

- Languagefeature
- Transformationfeature
- Toolfeature
- Processfeature

At the root level, language and transformation are mandatory features because they are part of the *DSL* definition. Tool is also mandatory because it serves to automate transformation from a domain, the problem space, down to lower abstraction levels, the solution space. Process is optional because it can be undefined or implicit.

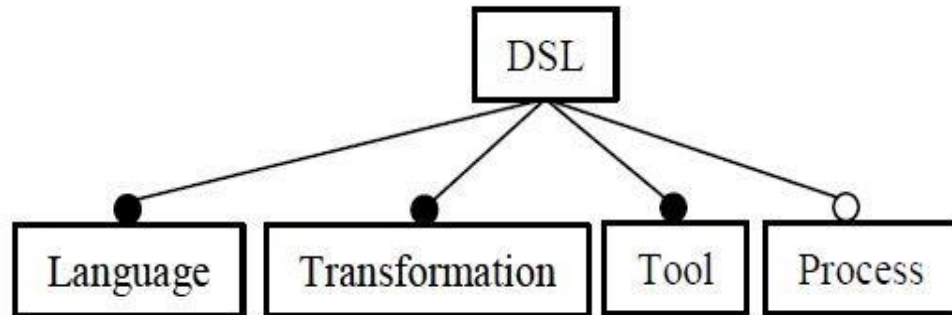


Figure 2.1: Root of the DSL Feature Model

3.4 Benefits of DSL

Fast DSL access allows you to download images, videos and other large files at lightning-fast speeds.

It is always on—there is no waiting to get connected, no busy signals, no dialing required with DSL.

No extra phone line necessary—DSL technology uses your existing phone line, allowing you to share phone and Internet on the same line—at the same time. You can also use other devices as usual on your telephone line. DSL won't interfere.

Fast internet connection- you will be able to download information, graphics and video from the internet at speeds up to 2.3Mb/s; and, you can upload information as fast as 2.3Mb/s.

Dedicated line- you will be the only person using your connection and line. Unlike cable or modems, you do not share your connection with other users, so your Internet speed will be more private, stable and will have fewer delays.

Available in limited areas- DSL users must be within a certain radius of the area telephone switch.

In essence, DSL has the following advantages:

- the speed is much higher than a regular modem
- DSL does not necessarily require new wiring
- the company that offers DSL will usually provide the modem as part of the installation

Disadvantages of DSL

- The service is not available everywhere
- The connection is fast for receiving data than it is for sending data over the internet
- A DSL works better when you are close to the provider's central office

3.5 How DSL Works

DSL is an innovative technology that allows your computer to transmit information over existing phone line, but at a higher frequency than telephones, fax machines and other devices. Since it utilizes the higher frequencies, you can use the same phone line for telephone calls and your internet connection at the same time, without any interference. Your telephone calls will still be clear and crisp, and your Internet connection will seem like lightning compared to a 56K modem connection.

4.0 CONCLUSION

ADSL is a problem-oriented language, which combined to transformation tools, such as generators, servers to raise the abstraction level of software and ease software development; but beyond this general definition, DSL and DSL tool variants are numerous. The reason of a DSL feature model is to formalise DSL and DSL tool variants. A first application of this feature model is a DSL tool factory, which applies variations during production of DSL tools. A second application

is the selection of pertinent *DSL* families among all possible families from the feature model. A third application is the definition of *DSL* tool foundations. A fourth usage is the selection of *DSL* tools. The feature model, combined with classification criteria, contains needed information to evaluate *DSL* tools. *DSL* feature model is in the scope of domain analysis of *DSL*s. Its clarification becomes a prerequisite for long-term and large scale *DSL* developments.

5.0 SUMMARY

In this unit, you have learnt that:

- the implementation of Digital Subscriber Line technology originally was part of the Integrated Services Digital Network (ISDN) specification published in 1984 by the CCITT and ITU as part of recommendation I.120, later reused as *ISDN* Digital Subscriber Line (IDSL)

- *DSL* is a telephone loop technology that uses existing copper phone lines, and provides a dedicated, high-speed Internet connection
- *DSL* feature model consists of language, transformation, tool, process
- *DSL* is an innovative technology that allows your computer to transmit information over existing phone lines, but at a higher frequency than telephones, fax machines and other devices
- *ADSL* is a problem-oriented language, which—combined with transformation tools, such as generators, servers or a set of the abstraction level of software—eases software development

6.0 TUTOR-MARKED ASSIGNMENT

What are the advantages and disadvantages of *DSL*?

7.0 REFERENCES/FURTHER READING

ANT, Apache Project, <http://ant.apache.org/>

ATL, Atlas Transformation Language, official web-site.
<http://www.sciences.univnantes.fr/lina/atl>.

Bass, L.; Clements, P. & Kazman, R. (2003). *Software Architecture in Practice*. Addison Wesley.

Czarnecki, K. & Eisenecker, U. W. (2000). *Generative Programming*. Addison-Wesley.

Czarnecki, K. & Helsen, S. (2003). "Classification of Model Transformation Approaches". OOPSLA 2003, Workshop on Generative Techniques in the Context of Model-Driven Architecture.

'Eclipse'. <http://www.eclipse.org/gmf/>

Greenfield, J.; Short, K.; Cook, S. & Kent, S. (2004). *Software Factories, Assembling applications with Patterns, Models, Framework, and Tools*. Wiley.

IEEE. 'IEEE Recommended Practice for Architectural Description of Software-Intensive Systems'. *IEEE Std 1471-2000*, 21 September 2000.

UNIT3 SYNCHRONOUS OPTICAL NETWORK (SONET)

CONTENT

S

- 1.0 Introduction
- 2.0 Objectives
- 3.0 MainContent
 - 3.1 Definitionof*SONET*
 - 3.2 Inceptionof*SONET*
 - 3.3 *SONET*Standards
 - 3.4 *SONET*Topology
 - 3.4.1 Pointto PointConfiguration
 - 3.4.2 *Hubbed*Configuration
 - 3.4.3 Linear Add/Dropconfiguration
 - 3.4.4 Self-healingringconfiguration
 - 3.5 *SONET*EquipmentLayers
 - 3.6 TheImportanceofSynchronisedTiming
- 4.0 Conclusion
- 5.0 Summary
- 6.0 Tutor-MarkedAssignment
- 7.0 References/FurtherReading

1.0 INTRODUCTION

In today's business world, each industry is looking for different ways to create competitive advantages to deliver information, products and services in a more timely and cost-effective manner. End-to-end SONET (Synchronous Optical Networking) network solutions are one important ingredient in creating a competitive edge. As a convergence technology, SONET provides for the unification of voice, data and video over the same transport service.

This unit gives you an operational overview of SONET (Synchronous Optical Networking). Here, an important point to remember is this- SONET is a powerful, highly scalable technology. Although it may appear to be complex, most of what goes on in a SONET network is transparent to the user. Also, this unit briefly discusses Wave Division Multiplexing (WDM) for awareness purposes only, since WDM is another high-performance transport technology that also leverages fiber optics.

SONET is a transport technology, designed to provide enterprise and government users – as well as service providers – a network infrastructure with survivability characteristics, so that business operations can continue uninterrupted. SONET's self-healing fiberoptic

ring functionality enables automatic network recovery due to failures that can be caused by a fiber optic cable cut, lost signal, or degraded signal (e.g. due to aging laser) or node/system failure. *SONET* is also a technology that is designed to ensure network traffic is restored within 60 milliseconds in the event of a failure.

2.0 OBJECTIVES

At the end of this unit, you should be able to:

- define the term *SONET*
- state the standards of *SONET*
- discuss the topology of *SONET*
- describe the *SONET* equipment layers
- explain *SONET* synchronisation and timing

3.0 MAIN CONTENT

3.1 Definition of *SONET*

Synchronous optical network (*SONET*) is a standard for optical telecommunications transport formulated by the Exchange Carriers Standards Association (ECSA) for the American National Standards Institute (ANSI), which sets industry standards in the U.S. for

telecommunications and other industries. The comprehensive *SONET* standard is expected to provide the transport infrastructure for worldwide telecommunications for at least the next two or three decades.

3.2 Inception of *SONET*

SONET was conceived of and written about back in the early 1980's, when submitted to the members of American National Standard Institute (ANSI) *T1* Committee as a universal transport system. In the mid-1980's, the *T1* Committee further enhanced the standard to arrive at the Synchronous Transport Signal One (STS-1) as the base-signaling rate. Around this time, the ITU-T (International Telecommunication Union-Telecommunications standard) (formerly CCITT) adopted *SONET* as the basis for its international standard referred to as SDH (Synchronous Digital Hierarchy) transport system, where the *STS-1* rate (51.84 Mbps) was to be a factor of 3 in terms of the European base rate of 155.52 Mbps.

3.3 SONET Standards

The base standard for SONET is the T1.105-1991 American National Standard (ANSI) for Telecommunications-Digital Hierarchy-Optical Interface Rates and Formats Specification (SONET). SONET standards define rates and formats as well as optical interfaces. The following American National Standards Institute (ANSI) specifications provide the primary standards, which define SONET:

- *ANSI T1.106-1998* specification for optical parameters
- *ANSI T1.102-1993* specification for electrical parameters
- *ANSI T1.105-1991* specification for multiplexing methods to map existing Digital Signals (e.g. DS1) into SONET payload signals
- *ANSI T1.105-1991* specification for criteria for optical line automatic
- Protection switching
- *ANSI T1.105-1991* specification for overhead channels to support standard
- Operation, Administration and Provisioning (OAM&P)

SELF-ASSESSMENT EXERCISE 1

Describe the structure of SONET

3.4 SONET Topology

SONET technology enables a number of different network topologies to solve networking requirements, including survivability, cost, and bandwidth efficiencies. The following provides a description of 3 different *SONET* configurations, which are deployed in a variety of enterprise situations. *SONET* configurations include:

- point-to-point configuration
- *hubbed* configuration
- linear add/drop configuration
- ring configuration

3.4.1 Point-to-Point Configuration

SONET point-to-point configurations create a simple topology that terminates a *SONET* payload at each point of a fiber optic cable span. Point-to-point configurations are typically deployed in transport applications, which require a single *SONET* multiplexer in a single route. Point-to-point configurations can be enhanced to increase

survivability by deploying a protection path (second fiber span) over a different path between two or more SONET multiplexers.

3.4.2 Hubbed Configuration

Hubbed configurations consolidate traffic from multiple sites onto a single optical channel, which then can be forwarded to another site. This topology is often used in applications where the user wants to consolidate traffic from multiple satellite sites to a single site such as corporate headquarters, before extending it, in some cases to a central office. This topology helps to reduce the number of hops as well as the equipment required to create a multisite topology.

3.4.3 Linear Add/Drop Configuration

In the asynchronous digital signal hierarchy environment, every time a digital signal is accessed the entire signal needs to be multiplexed/demultiplexed, costing time and money at each site along a given path. However, a linear add/drop configuration enables direct access to VTS/STS channels at each intermediate site along a fiber optic path. Therefore the linear add/drop configuration eliminates the need to process (multiplex/demultiplex) the entire optical signal for pass-through traffic.

3.4.4 Self-Healing Ring Configuration

In a self-healing ring configuration, a mechanism referred to as automatic protection switching is employed. There are two types of protection ring—the first is UPSR (Unidirectional Path Switched Ring), the other is BLSR (Bi-Directional Line Switched Ring). Each of these is discussed later. The self-healing ring configuration is the most commonly deployed *SONET* topology in mission-critical government and enterprise backbones, due to its survivability characteristics. Automatic protection switching is a mechanism provided within the *SONET* specification that is designed to provide duplicate fiber span paths. In this configuration, a backup fiber span (protection ring) is enabled when and if there is a failure within the fiber span currently carrying traffic on a *SONET* network. It should be noted that during normal operating conditions, both fiber spans are always active, and a *SONET* multiplexer selects which fiber span to receive traffic, based on an internal algorithm (e.g., based on which fiber module was installed in the multiplexer first). The *SONET* standard specifies that the protection ring should automatically become the fiber span (ring) the *SONET* multiplexer receives traffic from within 60 milliseconds (unnoticeable to the user) in the event of a failure on the other fiber span.

3.5 SONET Equipment Layers

SONET defines the end-to-end connection as being made up of 3 different equipment layers, including Path Terminating Equipment (PTE), Line Terminating Equipment (LTE), and Section Terminating Equipment (STE).

a. Path Terminating Equipment (PTE)

STS (Synchronous Transport Signal) path terminating equipment provides the multiplexing and demultiplexing functions within a *SONET* network. Path terminating equipment can originate access, modify, or terminate path overhead in any combination.

b. Line Terminating Equipment (LTE)

SONET line terminating equipment provides the function that originates and terminates line signals. *SONET* line terminating equipment can originate, access, modify, or terminate line overhead in any combination.

c. Section Terminating Equipment (STE)

A *SONET* "section" is any two neighboring *SONET* network elements. *SONET* section terminating equipment can be a network element or a *SONET*

regenerator. *SONET* section terminating equipment can originate, access, modify, or terminate section overhead in any combination

3.6 Importance of Synchronised Timing

Network timing between *SONET* devices is an integral part of maintaining accurate information transmitted over a *SONET* network. In the earlier days of networking, the method used in timing was Asynchronous. In Asynchronous timing, each switch runs its own clock. In Synchronous timing, switches can use a single common clock to maintain timing. This single common clock is referred to as a Primary Reference Source (PRS) or Master Clock. Synchronous timing maintains better accuracy than Asynchronous timing because it uses a single common clocking scheme to maintain timing. This accurate timing often becomes important to government and enterprise applications, particularly when they are running time-sensitive applications (e.g. video streaming applications). There are three methods typically used in obtaining synchronous timing in *SONET* multiplexers (e.g. Lucent *DDM-2000* Multiplexer), they are:

- timing from an onboard internal oscillator
- timing from an incoming optical signal from a high-speed interface
- timing from an external source coming from a DS1 timing reference that can be stratum 3 or higher clocking

SELF-ASSESSMENT EXERCISE 2

What do you understand by synchronous timing?

4.0 CONCLUSION

In this unit, you have been introduced to the fundamental concepts of *SONET*. You also learnt the historical background of *SONET*, the various standards of *SONET*. You were also introduced to various topologies of *SONET*.

5.0 SUMMARY

In this unit, you have learnt that:

- *SONET* is a transport technology, designed to provide enterprise and government users as well as service providers
- *SONET* is a network infrastructure with survivability characteristics, so that business operations continue uninterrupted

- Synchronous optical network (SONET) is a standard for optical telecommunications transport formulated by the Exchange Carriers Standards Association (ECSA) for the American National Standards Institute (ANSI), which sets industry standards in the U.S. for telecommunications and other industries
- States the standards of SONET
- The base standard for SONET is the T1.105-1991 American National Standard (ANSI) for Telecommunications- Digital Hierarchy-Optical Interface Rates and Formats Specification (SONET)
- SONET consists of 3 different equipment layers, including Path Terminating Equipment (PTE), Line Terminating Equipment (LTE), and Section Terminating Equipment (STE).

6.0 TUTOR-MARKED ASSIGNMENT

Explain SONET topologies.

7.0 REFERENCES/FURTHER READING

"About Rand". *Paul Baran and the Origins of the Internet*.

<http://www.rand.org/about/history/baran.html>. Retrieved January 14, 2006.

"How does the Internet Work?". [http://tldp.org/HOWTO/Unix-and-](http://tldp.org/HOWTO/Unix-and-Internet-Fundamentals-HOWTO/internet.html)

[Internet-Fundamentals-HOWTO/internet.html](http://tldp.org/HOWTO/Unix-and-Internet-Fundamentals-HOWTO/internet.html). Retrieved June 15, 2009.

Johna Till Johnson. "Net was born of economic necessity, not fear".

[http://www.networkworld.com/columnists/2004/0607johnson.ht](http://www.networkworld.com/columnists/2004/0607johnson.html)
[ml](http://www.networkworld.com/columnists/2004/0607johnson.html). Retrieved June 15, 2009. URL is sufficient attribution

Leonard Kleinrock (2005). *The history of the Internet*.

http://www.lk.cs.ucla.edu/personal_history.html. Retrieved 2009-05-28.

UNIT4 SWITCHING

CONTENTS

- 1.0 Introduction
- 2.0 Objectives
- 3.0 MainContent
 - 3.1 Overview
 - 3.2 Taxonomy of switched networks
 - 3.2.1Circuit-Switched Networks
 - 3.2.2Packet-switched Networks
 - 3.2.3Message-Switched Networks
- 4.0 Conclusion
- 5.0 Summary
- 6.0 Tutor-MarkedAssignment
- 7.0 References/FurtherReading

1.0 INTRODUCTION

Network switching is the process of channeling data received from any number of input ports to another designated port that will transmit the data to its desired destination. The device through witch the input data passes is called a switch. Data entering a port is referred to as ingress, while data leaving the port is referred to as egress. The switch represents the medium through which the data is routed to its final destination. There are different types of network switching. To better understand this let us examine the following illustration.

2.0 OBJECTIVES

At the end of this unit, you should be able to:

- define the term circuit-switched network
- discuss the three phases in circuit-switched network
- differentiate between datagram network and virtual-circuit network.
- state the advantages and disadvantages of packet switching

3.0 MAIN CONTENT

3.1 Overview

A network is a set of connected devices. Once 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 the majority 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 3.1 shows a switched network.

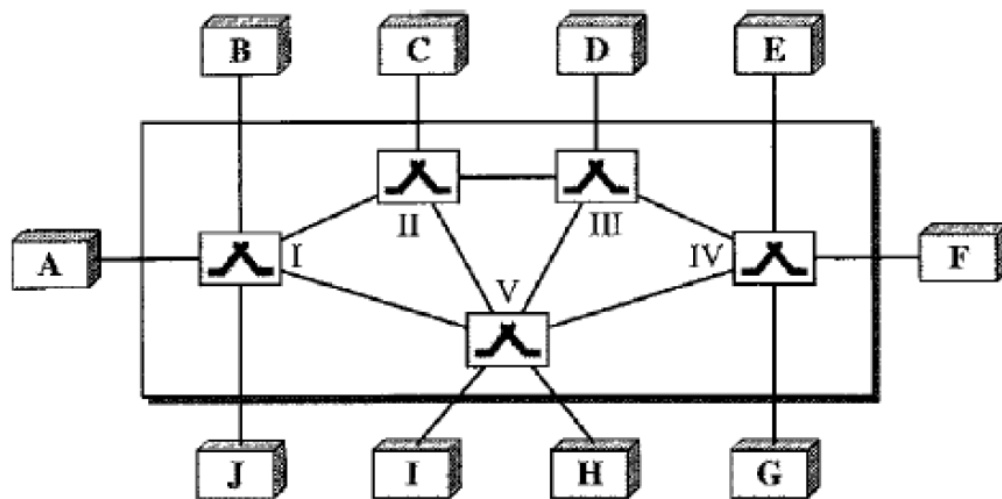


Figure 3.1 Switched network

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.

3.2 Taxonomy of switched networks

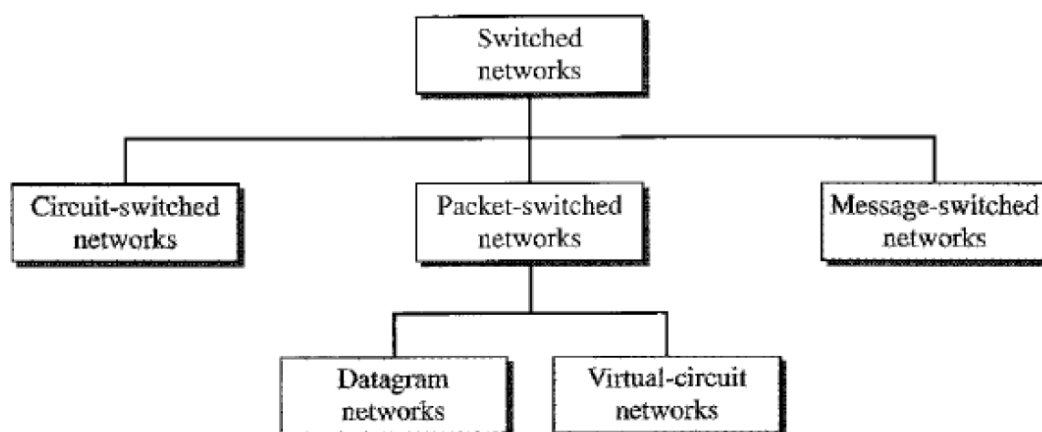


Figure 3.2: switched networks Taxonomy

3.2.1 CIRCUIT-SWITCHED NETWORKS

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

Figure 3.2.1 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.

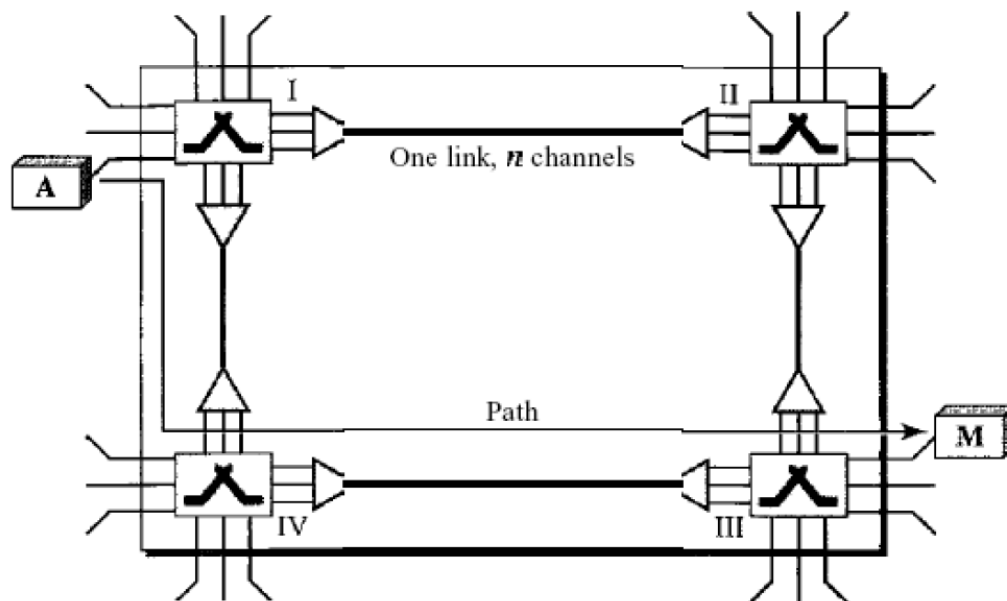


Fig 3.2.1: A trivial circuit-switched network

The Phases of circuit-switched network

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

1. 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 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 sends the request to switch IV, which finds a dedicated channel between itself and switch III. Switch III informs system M of system A's intention at this time.

In the next step to making a connection, an acknowledgment from system M needs to be sent in the opposite direction to system A. Only after system A receives this acknowledgment is the connection established. 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

2. Data Transfer Phase:

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

3. Teardown Phase:

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

Efficiency of circuit-switched networks:

It can be argued that circuit-switched networks are not as efficient as the other two types of networks because resources are allocated during the entire duration of the connection. These resources are unavailable to other connections. In a telephone network, people normally terminate the communication when they have finished their conversation. However, in computer networks, a computer can be connected to another computer even if there is no activity for a long time. In this case, allowing resources to be dedicated means that other connections are deprived.

Delay of circuit-switched networks:

Although a circuit-switched network normally has low efficiency, the delay in this type of network is minimal. During data transfer the data are not delayed at each switch; the resources are allocated for the duration of the connection. Figure 3.2.2 shows the idea of delay in a circuit-switched network when only two switches are involved. As Figure shows, there is no waiting time at each switch. The total delay is due to the time needed to create the connection, transfer data, and disconnect the circuit.

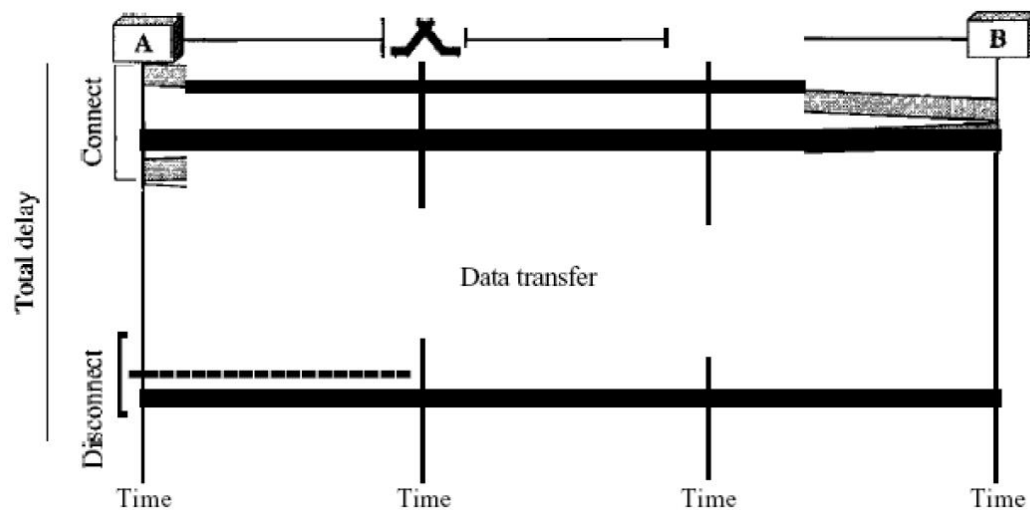


Fig3.2.2: Delay in a circuit-switched network

The delay caused by the setup is the sum of four parts. The propagation time of the source computer request (slope of the first gray box), the request signal transfer time (height of the first gray box), the propagation time of the acknowledgment from the destination computer (slope of the second gray box), and the signal transfer time of the acknowledgment (height of the second gray box). The delay due to data transfer is the sum of two parts: the propagation time (slope of the colored box) and data transfer time (height of the colored box), which can be very long. The third box shows the time needed to tear down the circuit. We have shown the case in which the receiver requests disconnection, which creates the maximum delay.

3.2.2 Packet Switching

Packet switching is dividing of messages into *packets* before they are sent, transmitting each packet individually, and then reassembling them into the original message once all of them have arrived at the intended destination.

Packets are the fundamental unit of information transport in all modern computer networks, and increasingly in other communications networks as well. Each packet, which can be of fixed or variable size depending on the protocol, consists of a *header*, body (also called a *payload*) and a *trailer*. The body contains a segment of the message being transmitted.

3.2.2.1 Advantages of Packet Switching

The following are some advantages of packet switching:

- Share link usage and greater link efficiency
- Links can transmit at different rates
- Under heavy load, calls can still be accepted but with greater delay
- Can prioritize packet transmission

3.2.2.2 Disadvantages of Packet Switching

Despite its great advantages, packet switching has the following disadvantages:

- Packet delay at each node (Processing, Queuing, Transmission)
- A packet can acquire jitter while passing through the network
- Packets may contain metadata i.e Higher overhead than circuit Switching (e.g., telephone system)
- More processing required at each node in the network.

3.2.3 Packet Switching Methods

3.2.3.1 DATAGRAM NETWORKS (Connectionless)

In a datagram network, each packet is treated independently of all others. Even if a packet is part of a multi-packet transmission, the network treats it as though it existed alone. Packets in this approach are referred to as datagrams.

Datagram switching is normally done at the network layer. We briefly discuss datagram networks here as a comparison with circuit-switched and virtual-circuit switched networks

Figure 3.2.3.1 shows how the datagram approach is used to deliver four packets from station A to station X. The switches in a datagram network are traditionally referred to as routers. That is why we use a different symbol for the switches in the figure.

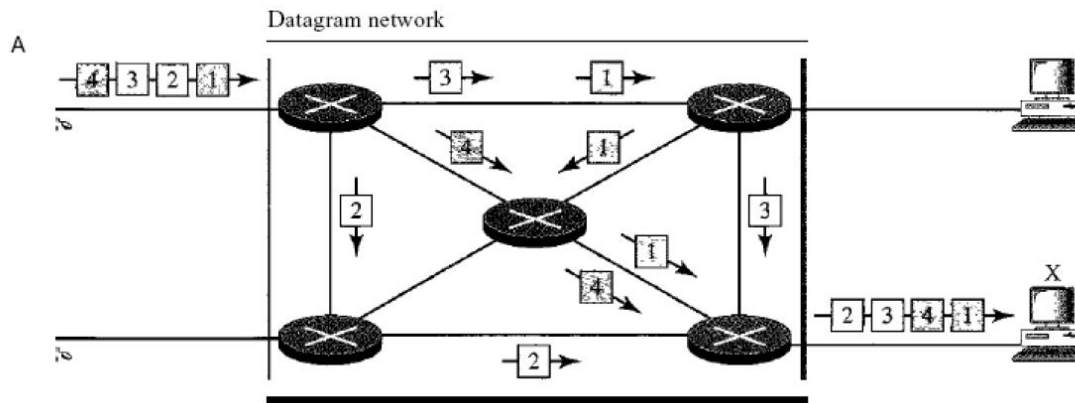


Fig 3.2.3.1: A datagram network with four switches (routers)

In this example, all four packets (or datagrams) belong to the same message, but may travel different paths to reach their destination. This is so because the links may be involved in carrying packets from other sources and do not have the necessary bandwidth available to carry all the packets from A to X. This approach can cause the datagrams of a transmission to arrive at their destination out of order with different delays between the packets. Packets may also be lost or dropped because of a lack of resources. In most protocols, it is the responsibility of an upper-layer protocol to reorder the datagrams or ask for lost datagrams before passing them on to the application.

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

Routing Table

If there are no setup or teardown phases, how are the packets routed to their destinations in a datagram network? In this type of network, each switch (or packet switch) has a routing table which is based on the destination address. The routing tables are dynamic and are updated periodically. The destination addresses and the corresponding forwarding output ports are recorded in the tables. This is different from the table of a circuit switched network in which each entry is created when the setup phase is completed and deleted when the teardown phase is over. Figure 3.2.3.2 shows the routing table for a switch.

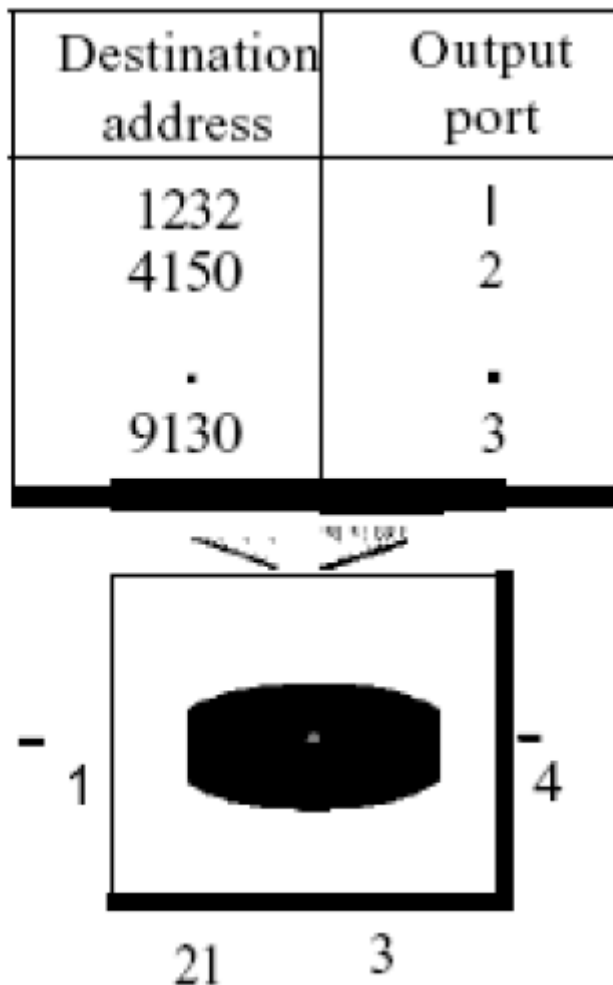


Fig 3.2.3.2: Routing table in a datagram network

Destination Address of datagram network

Every packet in a datagram network carries a header that contains, among other information, the destination address of the packet. When the switch receives the packet, this destination address is examined; the routing table is consulted to find the corresponding port through which the packet should be forwarded. This address, unlike the address in a virtual-circuit-switched network, remains the same during the entire journey of the packet.

Efficiency of datagram network

The efficiency of a datagram network is better than that of a circuit-switched network; resources are allocated only when there are packets to be transferred. If a source sends a packet and there is a delay of a few minutes before another packet can

be sent, the resources can be reallocated during these minutes for other packets from other sources.

Delay of datagram network

There may be greater delay in a datagram network than in a virtual-circuit network. Although there are no setup and teardown phases, each packet may experience a wait at a switch before it is forwarded. In addition, since not all packets in a message necessarily travel through the same switches, the delay is not uniform for the packets of a message.

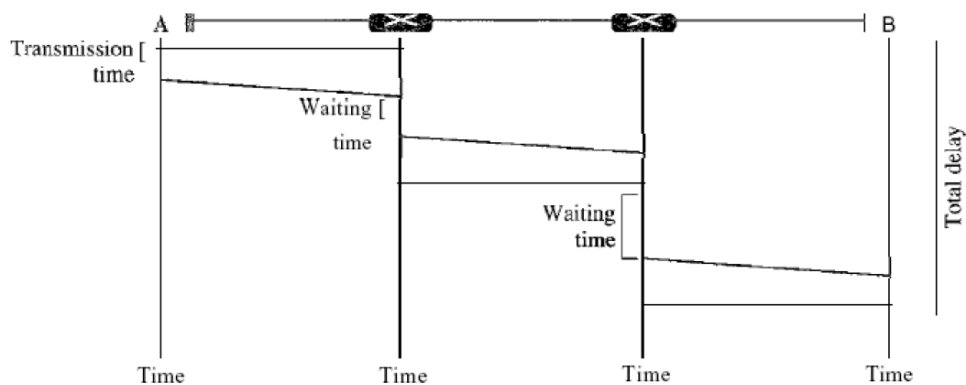


Fig 3.2.3.3: Delay in a datagram network

The packet travels through two switches. There are three transmission times ($3T$), three propagation delays (slopes $3t$ of the lines), and two waiting times ($W_1 + W_2$)

We ignore the processing time in each switch. The total delay is

$$\text{Total delay} = 3T + 3t + W_1 + W_2$$

3.2.3.2 VIRTUAL-CIRCUIT NETWORKS (Connection-Oriented):

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

1. As in a circuit-switched network, there are setup and teardown phases in addition to the data transfer phase.
2. Resources can be allocated during the setup phase, as in a circuit-switched network, or on demand, as in a datagram network.
3. As in a datagram network, data are packetized and each packet carries an address in the header. However, the address in the header has local jurisdiction (it defines what should be the next switch and the channel on

which the packet is being carried), not end-to-end jurisdiction. The reader may ask how the intermediate switches know where to send the packet if there is no final destination address carried by a packet. The answer will be clear when we discuss virtual-circuit identifiers in the next section.

4. As in a circuit-switched network, all packets follow the same path established during the connection.
5. A virtual-circuit network is normally implemented in the data link layer, while a circuit-switched network is implemented in the physical layer and a datagram network in the network layer. But this may change in the future. Figure 3.2.3.2 is an example of a virtual-circuit network. The network has switches that allow traffic from sources to destinations. A source or destination can be a computer, packet switch, bridge, or any other device that connects other networks.

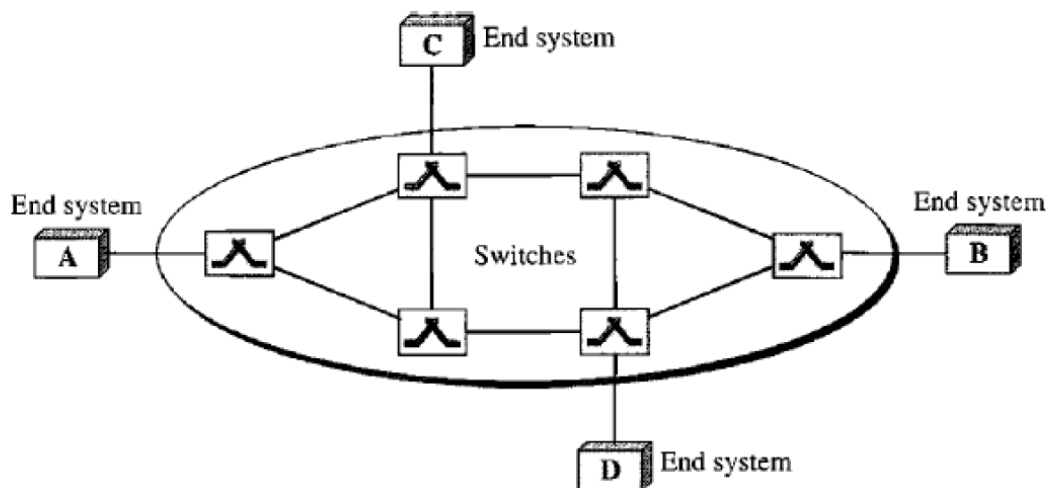


Fig 3.2.3.2 : Virtual-circuit network

Addressing of a virtual-circuit network

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

Global Addressing: A source or a destination needs to have a global address—an address that can be unique in the scope of the network or internationally if the network is part of an international network. However, we will see that a global address in virtual-circuit networks is used only to create a virtual-circuit identifier, as discussed next.

Virtual-Circuit Identifier: The identifier that is actually used for data transfer is called the virtual-circuit identifier (VCI). A VCI, unlike a global address, is a small number that has only switch scope; it is used by a frame between two switches. When a frame arrives at a switch, it has a VCI; when it leaves, it has a different VCI. Fig 3.2.3.2 shows how the VCI in a data frame changes from one switch to another. Note that a VCI does not need to be a large number since each switch can use its own unique set of VCIs.

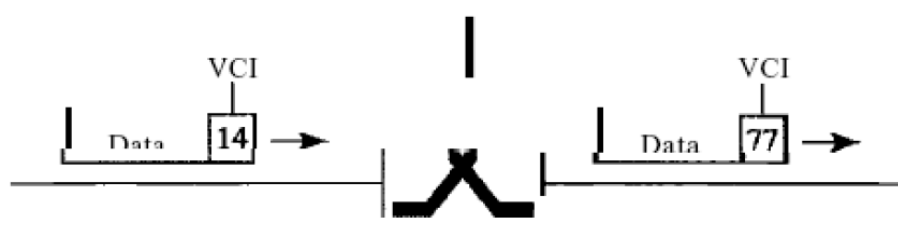


Fig 3.2.3.3: Virtual-Circuit Identifier

Three Phases in circuit-switched network

As in a circuit-switched network, a source and destination need to go through three phases in a virtual-circuit network: setup, data transfer, and teardown. In the setup phase, the source and destination use their global addresses to help switches make table entries for the connection. In the teardown phase, the source and destination inform the switches to delete the corresponding entry. Data transfer occurs between these two phases. We first discuss the data transfer phase, which is more straightforward; we then talk about the setup and teardown phases.

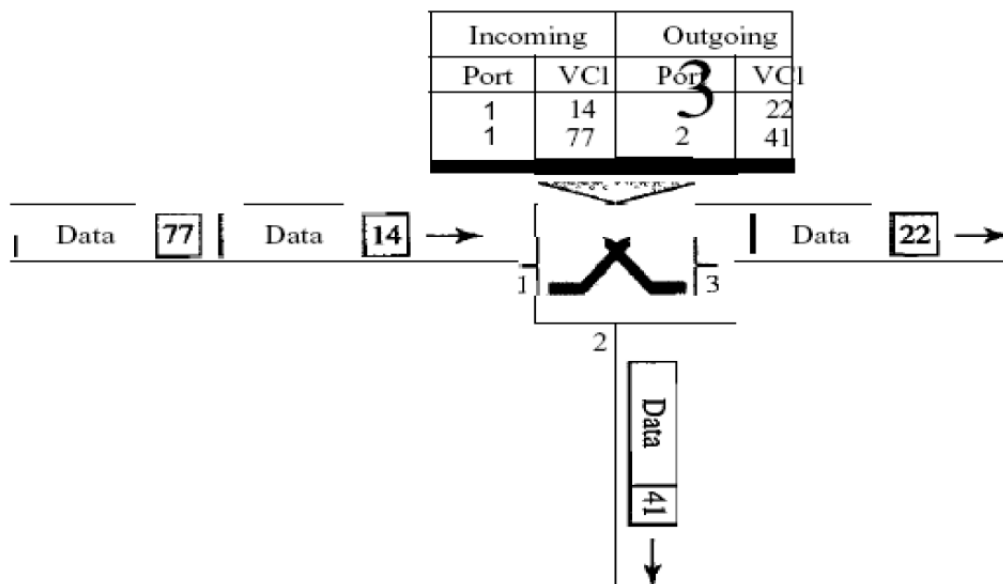
1. Data Transfer Phase

To transfer a frame from a source to its destination, all switches need to have a table entry for this virtual circuit. The table, in its simplest form, has four columns. This means that the switch holds four pieces of information for each virtual circuit that is already set up. We show later how the switches make their table entries, but for the moment we assume that each switch has a table with entries for all active virtual circuits. Figure 3.2.3.3 shows such a switch and its corresponding table. And also shows a frame arriving at port 1 with a VCI of 14. When the frame arrives, the switch looks in its table to find port 1 and a VCI of 14. When it is found, the switch

knows to change the VCI to 22 and send out the frame from port 3. Figure 3.2.3.4 shows how a frame from source A reaches destination B and how its VCI changes during the trip. Each switch changes the VCI and routes the frame. The data transfer phase is active until the source sends all its frames to the destination. The procedure at the switch is the same for each frame of a message. The process creates a virtual circuit, not a real circuit, between the source and destination.

2.Setup Phase

In the setup phase, a switch creates an entry for a virtual circuit. For example, suppose source A needs to create a virtual circuit to B. Two steps are required: the setup request and the acknowledgment.



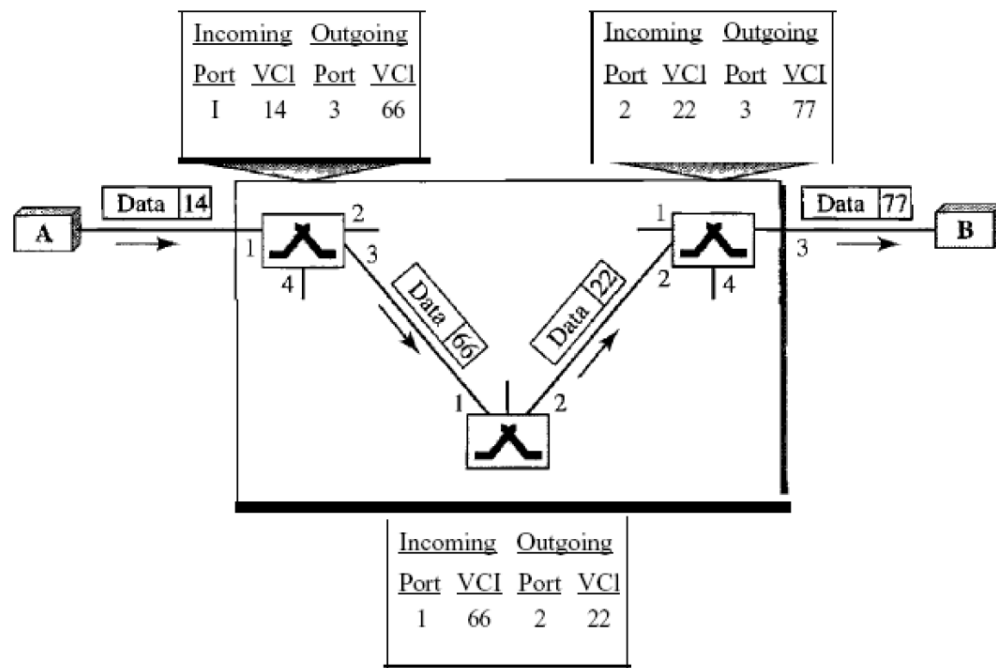


Fig 3.2.3.3: Switch and tables in a virtual-circuit network

Setup Request

A setup request frame is sent from the source to the destination.

Figure 3.2.3.4 shows the process.

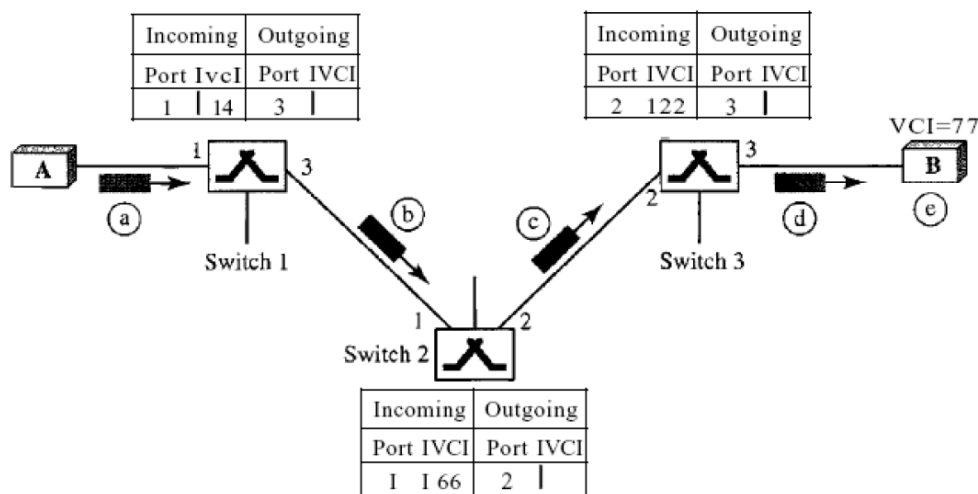


Fig 3.2.3.4: Setup request process

- i. Source A sends a setup frame to switch 1.
- ii. Switch 1 receives the setup request frame. It knows that a frame going from A to B goes out through port 3. The switch, in the setup phase, acts as a

packet switch; it has a routing table which is different from the switching table. For the moment, assume that it knows the output port. The switch creates an entry in its table for this virtual circuit, but it is only able to fill three of the four columns. The switch assigns the incoming port (1) and chooses an available incoming VCI (14) and the outgoing port (3). It does not yet know the outgoing VCI, which will be found during the acknowledgment step. The switch then forwards the frame through port 3 to switch 2.

- iii. Switch 2 receives the setup request frame. The same events happen here as at switch 1; three columns of the table are completed: in this case, incoming port (1), incoming VCI (66), and outgoing port (2).
- iv. Switch 3 receives the setup request frame. Again, three columns are completed: incoming port (2), incoming VCI (22), and outgoing port (3).
- v. Destination B receives the setup frame, and if it is ready to receive frames from A, it assigns a VCI to the incoming frames that come from A, in this case 77. This VCI lets the destination know that the frames come from A, and not other sources.

Acknowledgment A special frame, called the acknowledgment frame, completes the entries in the switching tables. Figure 3.2.3.5 shows the process.

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

- iv. Finally switch 1 sends an acknowledgment to source A that contains its incoming VCI in the table, chosen in the previous step.
- v. The source uses this as the outgoing VCI for the data frames to be sent to destination B.

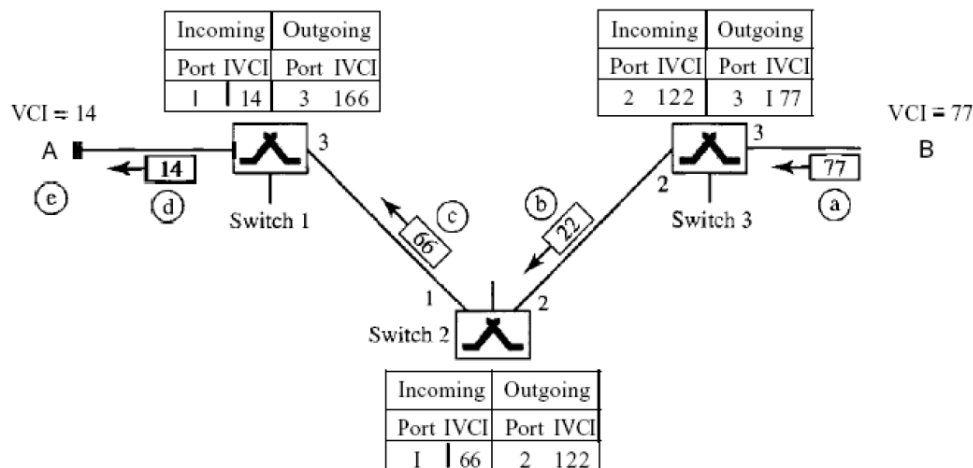


Fig 3.2.3.5: Setup acknowledgments in a virtual-circuit network

3. Teardown Phase

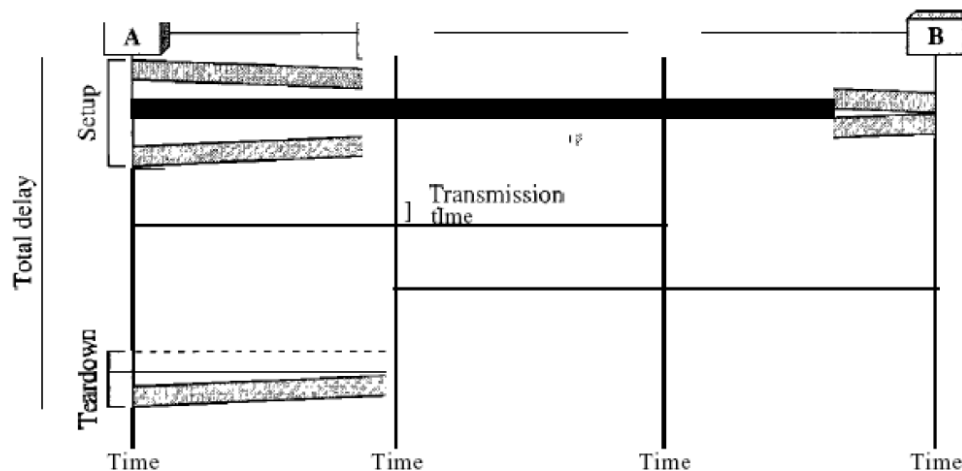
In this phase, source A, after sending all frames to B, sends a special frame called a *teardown request*. Destination B responds with a teardown confirmation frame. All switches delete the corresponding entry from their tables.

Efficiency of a virtual-circuit network

As we said before, resource reservation in a virtual-circuit network can be made during the setup or can be on demand during the data transfer phase. In the first case, the delay for each packet is the same; in the second case, each packet may encounter different delays. There is one big advantage in a virtual-circuit network even if resource allocation is on demand. The source can check the availability of the resources, without actually reserving it. Consider a family that wants to dine at a restaurant. Although the restaurant may not accept reservations (allocation of the tables is on demand), the family can call and find out the waiting time. This can save the family time and effort.

Delay in Virtual-Circuit Networks

In a virtual-circuit network, there is a one-time delay for setup and a one-time delay for teardown. If resources are allocated during the setup phase, there is no wait time for individual packets. Below Figure shows the delay for a packet traveling through two switches in a virtual-circuit network.



The packet is traveling through two switches (routers). There are three transmission times ($3T$), three propagation times ($3t$), data transfer depicted by the sloping lines, a setup delay (which includes transmission and propagation in two directions), and a teardown delay (which includes transmission and propagation in one direction). We ignore the processing time in each switch. The total delay time is

$$\text{Total delay} = 3T + 3t + \text{setup delay} + \text{teardown delay}$$

3.2.3 Message-switched Networks

Message switching is a connectionless network switching technique where the entire message is routed from the source node to the destination node, one hop at a time. It was a precursor of packet switching.

Packet switching treats each message as an individual unit. Before sending the message, the sender node adds the destination address to the message. It is then delivered entirely to the next intermediate switching node. The intermediate node stores the message in its entirety, checks for transmission errors, inspects the destination address and then delivers it to the next node. The process continues till the message reaches the destination. In the switching node, the incoming message is not discarded if the required outgoing circuit is busy. Instead, it is stored in a queue

for that route and retransmitted when the required route is available. This is called store and forward network.

The following diagram represents routing of two separate messages from the same source to same destination via different routes, using message switching

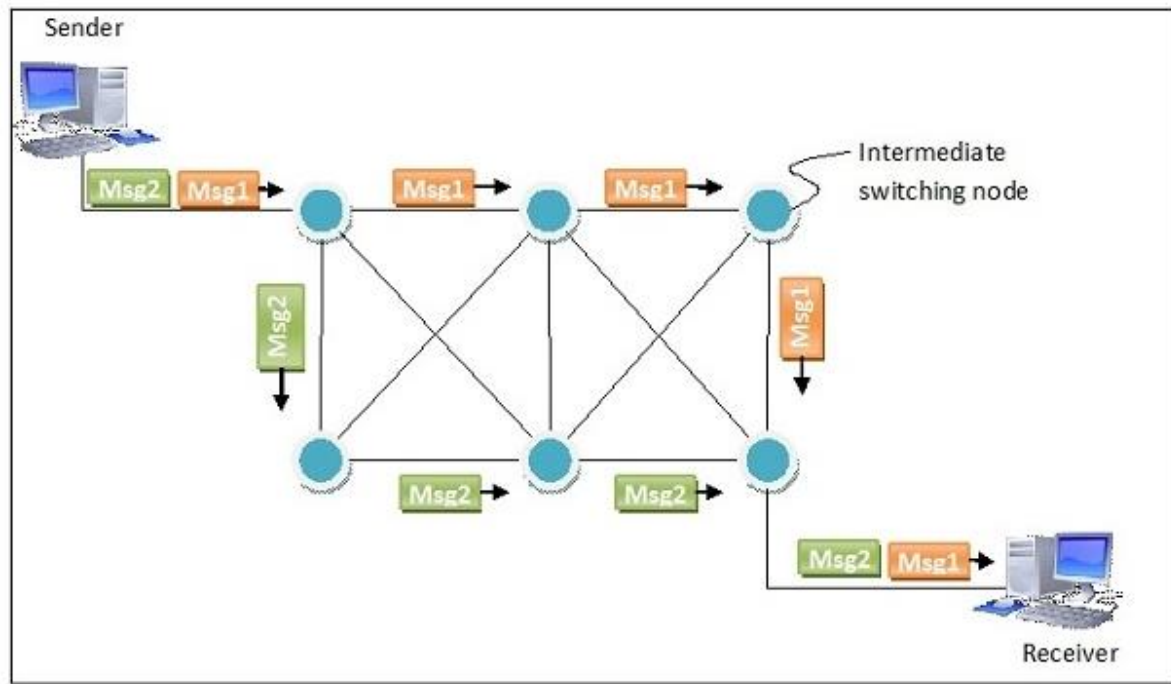


Fig 3.2.3: Routing of Messages

Advantages and Disadvantages of Message Switching

Advantages

- Sharing of communication channels ensures better bandwidth usage.
- It reduces network congestion due to store and forward method. Any switching node can store the messages till the network is available.
- Broadcasting messages requires much less bandwidth than circuit switching.
- Messages of unlimited sizes can be sent.
- It does not have to deal with out of order packets or lost packets as in packet switching.

Disadvantages

- In order to store many messages of unlimited sizes, each intermediate switching node requires large storage capacity.
- Store and forward method introduces delay at each switching node. This renders it unsuitable for real time applications.

Exercise 1: What are the advantages and disadvantages of the packet switching methods?

4.0 CONCLUSION

In this unit, you have been introduced to the fundamental concepts of Switching. You also learnt the advantages and disadvantages of packet switching, Message switching. You were also introduced to various packet switching methods.

5.0 SUMMARY

In this unit, you have learnt that:

- The concept of switching was first explored
- The taxonomy of network switching was discussed
- Packet switching is dividing of messages into packets before they are sent, transmitting each packet individually, and then reassembling them into the original message once all of them have arrived at the intended destination
- Packet switching methods datagram, virtual circuit and source routing.

6.0 TUTOR-MARKED ASSIGNMENT

1. State the advantages and the disadvantages of Packet Switching.
2. What do you understand by the term “Virtual-circuit network”

c

7.0 REFERENCES/FURTHER READING

- Samuel Sam (2018). "Message Switching"<https://www.tutorialspoint.com/message-switching>
- Stephen Cooper (2018). "A guide to message switched networks"<https://www.comparitech.com/net-admin/message-switched-networks-guide/>
- Vishnu Universal Learning. "Computer Networks"<https://www.coursehero.com/file/13176989/Docs-CSECDLNotes2013/>
- Ronda Hauben (2001). *From the ARPANET to the Internet*.
http://www.columbia.edu/~rh120/other/tcpdigest_paper.txt. Retrieved 2009-05-28.
- "NORSAR and the Internet". NORSAR. <http://www.norsar.no/pc-5-30-NORSAR-and-the-Internet.aspx>. Retrieved 2009-06-05.
- "History of X.25, CCITT Plenary Assemblies and Book Colors". Itu.int.
<http://www.itu.int/ITU-T/studygroups/com17/history.html>. Retrieved 2009-06-05.
- Prasad, K. V. (2009). "Principles of Digital Communication Systems and Computer Networks", Dreamtech Press.

UNIT5 INTERNETANDTCP/IP

CONTENTS

- 1.0 Introduction
- 2.0 Objectives
- 3.0 MainContent
 - 3.1 HistoricalBackground
 - 3.2.1 InternetStandardProcess
 - 3.3 TCP/IPTerminologies
 - 3.4 TCP/IPProtocolSuite
 - 3.5 NetworkInterfaceLayer
 - 3.6 InternetLayerInterface
 - 3.7 ApplicationLayerInterface
- 4.0 Conclusion
- 5.0 Summary
- 6.0 Tutor-MarkedAssignment
- 7.0 References/FurtherReading

1.0 INTRODUCTION

The "internet protocolsuite" is the set of communication protocols used for the internet and other similar networks. It is commonly also known as "TCP/IP", named from two of the most important protocols in it - the Transmission Control Protocol (TCP) and the Internet Protocol (IP), which were the first two networking protocols defined in this standard. Modern IP networking represents a synthesis of several developments

that began to evolve in the 1960s and 1970s, namely the internet and local area networks, which emerged during the 1980s, together with the advent of the World Wide Web (WWW) in the early 1990s. The internet protocol suite consists of four abstraction layers. From the lowest to the highest layer, these are the link layer, the internet layer, the transport layer, and the application layer.

The link layer contains communication technologies for the local network where the host is connected to directly, the link. It provides the basic connectivity functions interacting with the networking hardware of the computer and the associated management of interface-to-interface messaging. The internet layer provides communication methods between multiple links of a computer and facilitates the interconnection of networks. As such, this layer establishes the internet. It contains primarily the internet protocol, which defines the fundamental addressing namespaces, *IPv4* Internet Protocol Version 4 (IPv4) and *IPv6* Internet Protocol Version 6 (IPv6) used to identify and locate hosts on the network. Direct host-to-host communication tasks are handled in the transport layer, which provides a general framework to transmit data

between hosts using protocols like the transmission control protocol and the user datagram.

2.0 OBJECTIVES

At the end of this unit, you should be able to:

- describe the internet standards process
- define common terms used in *TCP/IP*
- state the advantages of including *TCP/IP* components in windows
- explain how the *TCP/IP* protocol suite maps to the Department of Defense Advanced Research Projects Agency (DARPA) and Open System Interconnection (OSI) models.

3.0 MAIN CONTENT

3.1 Historical Background

The internet protocol suite resulted from research and development conducted by the Defense Advanced Research Projects Agency (DARPA) in the early 1970s. After initiating the pioneering *ARPANET* in 1969, DARPA started work on a number of other data transmission technologies. In 1972, Robert E. Kahn joined the DARPA information processing technology office, where he worked on both satellite packet

networks and ground-based radiopacket networks, and recognised the value of being able to communicate across both. In the spring of 1973, Vinton Cerf, the developer of the existing ARPANET Network Control Program (NCP) protocol, joined Kahn to work on open-architecture interconnection models with the goal of designing the next protocol generation for the ARPANET.

By the summer of 1973, Kahn and Cerf had worked out a fundamental reformulation, where the differences between network protocols were hidden by using a common [[internet network protocol]], and, instead of the network being responsible for reliability, as in the ARPANET, the hosts became responsible.

The design of the network included the recognition that it should provide only the functions of efficiently transmitting and routing traffic between end nodes and that all other intelligences should be located at the edge of the network, in the end nodes. Using a simple design, it became possible to connect almost any network to the ARPANET, irrespective of their local characteristics, thereby solving Kahn's initial problem. One popular expression is that *TCP/IP*, the eventual product of Cerf and Kahn's work, will run over "two tin cans and a string". A computer, called a router, is provided with an interface to each network.

It forwards packets back and forth between them. Originally a router was called "gateway", but the term was changed to avoid confusion with other types of gateway (computer networking) computers.

The idea was worked out in more detailed form by Cerf's networking research group at Stanford within 1973–74, resulting in the first *TCP* specification. The internet protocols suite resulted from research and development conducted by the Defense Advanced Research Projects Agency (DARPA) in the early 1970s. After initiating the pioneering *ARPANET* in 1969, DARPA started work on a number of other data transmission technologies. In 1972, Robert E. Kahn joined the DARPA information processing technology office, where he worked on both satellite packet networks and ground-based radiopacket networks, and recognized the value of being able to communicate across both. In the spring of 1973, Vinton Cerf, the developer of the existing *ARPANET* Network Control Program (NCP) protocol, joined Kahn to work on open-architecture interconnection models with the goal of designing the next protocol.

3.2 Internet Standard Process

TCP/IP is the protocol of the Internet, it has evolved based on fundamental standards that have been created and adopted over more than 30 years. The future of *TCP/IP* is closely associated with the advances and administration of the Internet as additional standards continue to be developed. Although, no one organisation owns the Internet or its technologies, several organizations oversee and manage

these new standards, such as the Internet Society and the Internet Architecture Board.

The Internet Society (ISOC) was created in 1992 and is a global organisation responsible for the internet networking technologies and applications of the internet. Although the society's principal purpose is to encourage the development and availability of the internet, it is also responsible for the further development of the standards and protocols that allow the Internet to function. The ISOC sponsors the Internet Architecture Board (IAB), a technical advisory group that sets internet standards, publishes *RFCs*, and oversees the Internet standards process. The IAB governs the following bodies:

- the Internet Assigned Number Authority (IANA) oversees and coordinates the assignment of protocol identifiers used on the Internet.
- the Internet Research Task Force (IRTF) coordinates all TCP/IP-related research projects.

- the Internet Engineering Task Force (IETF) solves technical problems and needs as they arise on the internet and develops internet standards and protocols. *IETF* working groups define standards known as *RFC*.

3.3 TCP/IP Terminology

The internet standards use a specific set of terms when referring to network elements and concepts related to *TCP/IP* networking. Common terms and concepts in *TCP/IP* are defined as follows:

- Node**-any device, including routers and hosts, which runs an implementation of *IP*.
- Router**- a node that can forward *IP* packets not explicitly addressed to itself. On an *IPv6* network, a Router also typically advertises its presence and host configuration information.
- Host**- a node that cannot forward *IP* packets not explicitly addressed to itself (a non-router). A host is typically the source and the destination of *IP* traffic. A host silently discards traffic that it receives but that is not explicitly addressed to it.
- Upper-layer protocol**: a protocol above *IP* that uses *IP* as its transport. Examples include internet layer protocols such as the Internet Control Message Protocol (*ICMP*) and transport layer protocols such as the Transmission Control Protocol (*TCP*) and User Datagram Protocol (*UDP*). (However, application layer protocols that use *TCP* and *UDP* as their transports are not

considered upper-layer protocols. File Transfer Protocol [FTP] and Domain Name System [DNS] fall into this category).

- **LAN segment**-a portion of a subnet consisting of a single medium that is bounded by bridges or layer 2 switches.
- **Subnet**-one or more LAN segments that are bounded by routers and use the same IP address prefix. Other terms for subnet are network segment and link.
- **Network**-two or more subnets connected by routers. Another term for network is internet network.
- **Neighbour**-a node connected to the same subnet as another node.
- **Interface**-the representation of a physical or logical attachment of a node to a subnet. An example of a physical interface is a network adapter. An example of a logical interface is a tunnel.

interface that is used to send IPv6 packets across an IPv4 network.

- **Address-** an identifier that can be used as the source or destination of IP packets and that is assigned at the Internet layer to an interface or set of interfaces.
- **Packet-** The protocol Data Unit (PDU) that exists at the Internet layer and comprises an IP head payload. a common [[Internet network protocol]], and, instead of the network being responsible for reliability, as in the ARPANET, the hosts became responsible.

Table 4.1: The Advantages of the TCP/IP Protocol Suite and the Inclusion of TCP/IP Components in Windows.

Advantages of the TCP/IP protocol suite	Advantages of TCP/IP components in Windows
A standard, routable enterprise networking protocol that is the most complete and accepted protocol available. All modern operating systems support TCP/IP, and most large private networks rely on TCP/IP for much of their traffic.	TCP/IP components in Windows enable enterprise networking and connectivity for Windows and non-Windows-based computers.
A technology for connecting dissimilar systems. Many TCP/IP application protocols were designed to access and transfer data between dissimilar systems. These protocols include HTTP, FTP, and Telnet.	TCP/IP components in Windows allow standards-based connectivity to other operating system platforms.
A robust, scalable, cross-platform client/server framework.	TCP/IP components in Windows support the Windows Sockets application programming interface, which developers use to create client/server applications.
A method of gaining access to the Internet.	Windows-based computers are Internet-ready.

3.4 The TCP/IP Protocol Suite

The *TCP/IP* protocol suite maps to a four-layer conceptual model known as the *DARPA* model, which was named after the U.S. government agency that initially developed *TCP/IP*. The four layers of the *DARPA* model are: Application, Transport, Internet, and Network Interface. Each layer in the *DARPA* model corresponds to one or more layers of the seven-layer *OSI* model.

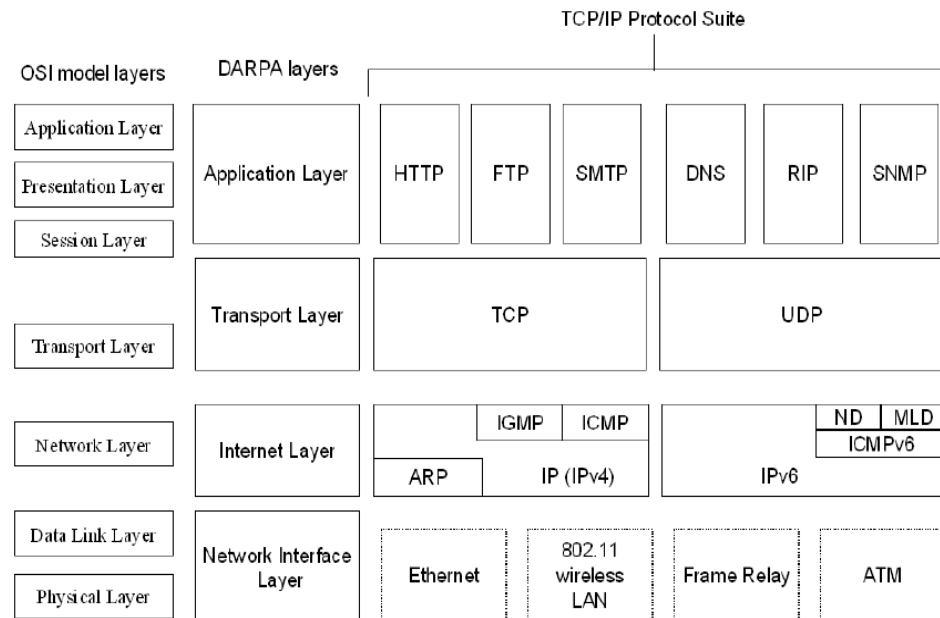


Figure 4.1: The Architecture of the TCP/IP model

The *TCP/IP* protocol suite has two sets of protocols at the internet layer:

- IPv4*, also known as *IP*, is the Internet layer in common use today on private intranets and the Internet.
- IPv6* is the new Internet layer that will eventually replace the existing *IPv4* Internet layer.

SELF-ASSESSMENT EXERCISE 1

What are the specific problems with *IPv4*?

3.5 Network Interface Layer

The network interface layer (also called the network access layer) sends *TCP/IP* packets on the network medium and receives *TCP/IP* packets off the network medium. *TCP/IP* was designed to be independent of the network access method, frame format, and medium. Therefore, you can use *TCP/IP* to communicate across differing network types that use LAN technologies—such as Ethernet and 802.11 wireless LAN—and WAN technologies—such as Frame Relay and Asynchronous Transfer Mode (ATM). By being independent of any specific network technology, *TCP/IP* can be adapted to new technologies. The network interface layer of the *DARPA* model encompasses the data link and physical layers of the *OSI* model. The internet layer of the *DARPA* model does not take advantage of sequencing and acknowledgment services that might be present in the data link layer of the *OSI* model. The internet layer

assumes an unreliable network interface layer and that reliable communications through session establishment and the sequencing and acknowledgment of packets is the responsibility of either the transport layer or the application layer.

3.6 Internet Layer

The internet layer responsibilities include addressing, packaging, and routing functions. The internet layer is analogous to the network layer of the OSI model. The core protocols for the IPv4 internet layer consist of the following:

- i. the Address Resolution Protocol (ARP) resolves the internet layer address to a network interface layer address such as a hardware address.
- ii. the Internet Protocol (IP) is a routable protocol that addresses, routes, fragments, and reassembles packets.
- iii. the Internet Control Message Protocol (ICMP) reports errors and other information to help you diagnose unsuccessful packet delivery.

SELF-ASSESSMENT EXERCISE 2

What is the difference between the internet layer and the network layer of the OSI model?

3.7 Application Layer Interface

The application layer allows applications to access the services of the other layers, and it defines the protocols that applications use to exchange data. The application layer contains many protocols, and more are always being developed.

The most widely known application layer protocols help users exchange information. Some of these are:

- i. the Hypertext Transfer Protocol (HTTP) transfers files that make up pages on the World Wide Web (WWW).
- ii. the File Transfer Protocol (FTP) transfers individual files, typically for an interactive user session
- iii. the Simple Mail Transfer Protocol (SMTP) transfers mail messages and attachments.

Additionally, the following application layer protocols help you use and manage *TCP/IP* networks:

- i. the Domain Name System (DNS) protocol resolves a hostname, such as 'www.microsoft.com', to an IP address and copies name information between DNS servers.
- ii. the Routing Information Protocol (RIP) is a protocol that routers use to exchange routing information on an IP network., called a router, is provided with an interface to each network.

4.0 CONCLUSION

In this unit, you have been introduced to the fundamental concepts of internet *TCP/IP*. You also learnt the internet standard process and some *TCP/IP* terminologies. You were also introduced to various *TCP/IP* protocol suites.

5.0 SUMMARY

In this unit, you have learnt that:

- the "internet protocol suite" is the set of communications protocols used for the internet and other similar networks
- the *TCP/IP* is the protocol of the internet; it has evolved based on fundamental standards that have been created and adopted over more than 30 years

- the *TCP/IP* protocol suite maps to a four-layer conceptual model known as the *DARPA* model, which was named after the U.S. government agency that initially developed *TCP/IP*
- the *TCP/IP* protocol suite has two sets of protocols at the internet layer—namely *IPv4* and *IPv6*
- the network interface layer (also called the network access layer) sends *TCP/IP* packets on the network medium and receives *TCP/IP* packets off the network medium
- the internet layer responsibilities include addressing, packaging, and routing functions
- the application layer allows applications to access the services of the other layers, and it defines the protocols that applications use to exchange data.

6.0 TUTOR-MARKED ASSIGNMENT

1. Discuss the history of the *TCP/IP* model.
2. What is the function of the internet and the application layer?

7.0 REFERENCES/FURTHER READING

"About Rand". *Paul Baran and the Origins of the Internet*.

<http://www.rand.org/about/history/baran.html>. Retrieved January 14, 2006.

<http://tldp.org/HOWTO/Unix-and-Internet-Fundamentals-HOWTO/internet.html>. Retrieved June 15, 2009.

John, Till, Johnson. "Net was born of economic necessity, not fear".

<http://www.networkworld.com/columnists/2004/0607johnson.html>. Retrieved June 15, 2009.

Prasad, K. V. (2009). *Principles of Digital Communication Systems and Computer Networks*. Dreamtech Press.