



**SHRI RAMDEOBABA COLLEGE OF ENGINEERING AND
MANAGEMENT,
KATOL ROAD, NAGPUR, MAHARASHTRA, INDIA**

Communication System Analysis (ECTH41)

Dr. (Mrs.) Mridula Korde

Associate Professor, Department of Electronics and
Communication Engineering, RCOEM, Nagpur

Honors Scheme

Sr No	Semester	Course Code	Course Title	Hours per Week	Credits	Maximum Marks			ESE Duration in Hours
						Continuous Evaluation	End Sem Exams	Total	
01	IV	ECTH41	Communication System Analysis	4	4	40	60	100	3
02	V	ECTH51	Radio Frequency Circuit Design	4	4	40	60	100	3
03	VI	ECTH61	Multimedia Networks	4	4	40	60	100	3
04	VII	ECTH71	Cryptography and Information Security	4	4	40	60	100	3
05	VIII	ECTH81	Evolution of Air Interface towards 5G	4	4	40	60	100	3

Syllabus for Semester IV, B. E. (Electronics and Communication Engineering – Honors)

Course Code: ECTH41

**Course: Communication System Analysis
(Honors Course)**

L: 4 Hrs, T: 0 Hr, P: 0 Hrs. Per week

Total Credits: 04

Course Objectives:

The Objective of this course is to make students aware of:

1. Advanced concepts in communication systems.
2. Various advanced modulation techniques.
3. Advanced concepts like synchronization, channel estimation

Course Outcomes:

After completion of this course, the students will be able to:

1. Understand the advanced concepts in communication systems.
 2. Understand advanced modulation techniques.
 3. Know advanced concepts like synchronization, channel estimation
 4. Analyze the behavior of ATM traffic in presence of congestion
-

Unit I

Spread Spectrum Communications: Spreading sequences- Properties of Spreading Sequences, Pseudo- noise sequence, Gold sequences, Kasami sequences, Walsh Sequences, Orthogonal Variable Spreading Factor Sequences, Barker Sequence, Complementary Codes Direct sequence spread spectrum – DS-CDMA Model, Conventional receiver, Rake Receiver, Synchronization in CDMA, Power Control, Soft handoff

Unit II

Orthogonal Frequency Division Multiplexing: Basic Principles of Orthogonality, Single vs Multicarrier Systems, OFDM Signal Mathematical Representation, Selection parameter for Modulation, Pulse shaping in OFDM Signal and Spectral Efficiency, Window in OFDM Signal and Spectrum, Synchronization in OFDM, Pilot Insert in OFDM Transmission and Channel Estimation

Unit III

MIMO Systems: Introduction, Space Diversity and System Based on Space Diversity, Smart Antenna system and MIMO, MIMO Based System Architecture, MIMO Exploits Multipath, Space – Time Processing, Antenna Consideration for MIMO, MIMO Channel Modelling, MIMO Channel Measurement, MIMO Channel Capacity, Cyclic Delay Diversity (CDD), Space Time Coding, Advantages and Applications of MIMO in Present Context, MIMO Applications in 3G Wireless System and Beyond, MIMO-OFDM

Unit IV

SONET/SDH: Architecture, SONET Layers, SONET Frames, STS Multiplexing, SONET Networks, Virtual Tributaries.

Unit V

ATM: Overview, Virtual channels, Virtual paths, VP and VC switching, ATM cells, Header format, Generic flow control, Header error control, Transmission of ATM cells, Adaptation layer, AAL services and protocols.

Unit VI

ATM Traffic and congestion Control: Requirements for ATM Traffic and Congestion Control, Cell Delay Variation, ATM Service Categories, Traffic and Congestion Control Framework, Traffic Control, Congestion Control

Dr. Mridula Korde, EC, RCOEM

Text Books:

1. Gary J. Mullett, “Introduction to Wireless Telecommunications Systems and Networks”, CENGAGE
2. Upena Dalal, “Wireless Communication”, Oxford University Press, 2009
3. William Stallings, “ISDN and Broadband ISDN with Frame Relay and ATM” Prentice Hall, 4th edition

Reference books:

- 1) Ke-Lin Du & M N S Swamy, “Wireless Communication System”, Cambridge University Press, 2010
- 2) Behrouz A Forouzan, “Data Communications and Networking”, 4th Edition, McGraw Hill.



Spread Spectrum Techniques

A Short History

- Spread-spectrum communications technology was first described on paper by an actress and a musician!
- In 1941 Hollywood actress Hedy Lamarr and pianist George Antheil described a secure radio link to control torpedos. They received U.S. Patent #2.292.387.
- The technology was not taken seriously at that time by the U.S. Army and was forgotten until the 1980s, when it became active.
- Since then the technology has become increasingly popular for applications that involve radio links in hostile environments.
- Typical applications for the resulting short-range data transceivers include satellite-positioning systems (GPS), 3G mobile telecommunications, W-LAN (IEEE® 802.11a, IEEE 802.11b, IEEE 802.11g), and Bluetooth®.
- Spread-spectrum techniques also aid in the endless race between communication needs and radio-frequency availability—situations where the radio spectrum is limited and is, therefore, an expensive resource.

Theoretical Justification for Spread Spectrum

- Spread-spectrum is apparent in the Shannon and Hartley channel-capacity theorem:

$$C = B \times \log_2 (1 + S/N)$$

In this equation,

- C is the channel capacity in bits per second (bps), which is the maximum data rate for a theoretical bit-error rate (BER).
- B is the required channel bandwidth in Hz.
- S/N is the signal-to-noise power ratio.
- To be more explicit, one assumes that C, which represents the amount of information allowed by the communication channel, also represents the desired performance. Bandwidth (B) is the price to be paid, because frequency is a limited resource. The S/N ratio expresses the environmental conditions or the physical characteristics (i.e., obstacles, presence of jammers, interferences, etc.).

Some more facts of Spread Spectrum

- Different spread-spectrum techniques are available, but all have one idea in common: the key (also called the code or sequence) attached to the communication channel.
- The term "spread spectrum" refers to the expansion of signal bandwidth, by several orders of magnitude in some cases, which occurs when a key is attached to the communication channel.
- The ratio (in dB) between the spread baseband and the original signal is called processing gain.
- Typical spread-spectrum processing gains run from 10dB to 60dB.
- Spread Spectrum refers to a system originally developed for military applications, to provide secure communications by spreading the signal over a large frequency band.

Some more facts of Spread Spectrum

- The idea behind spread spectrum is to use more bandwidth than the original message while maintaining the same signal power.
- A spread spectrum signal does not have a clearly distinguishable peak in the spectrum.
- This makes the signal more difficult to distinguish from noise and therefore more difficult to jam or intercept.

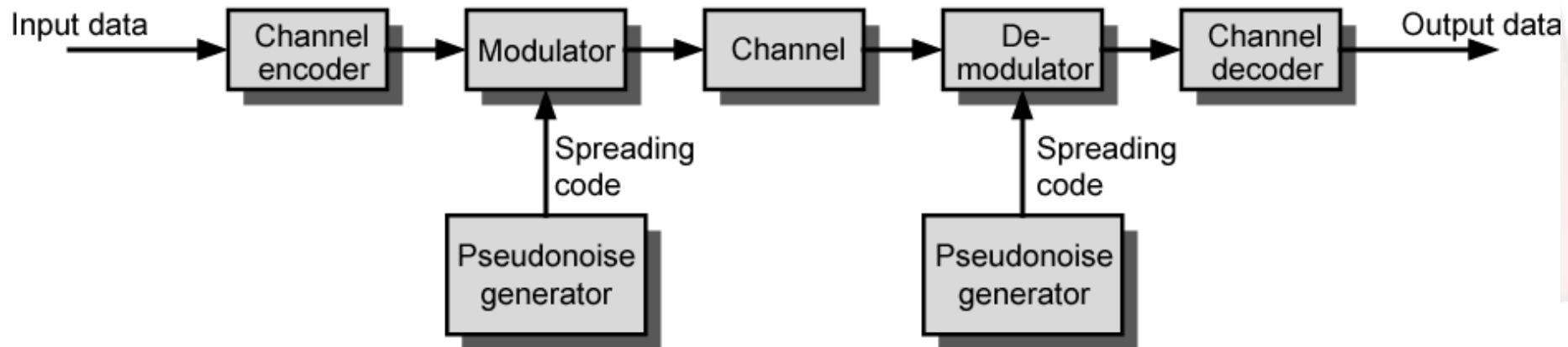
How Spread Spectrum Works

- Spread Spectrum uses wide band, noise-like signals.
- Because Spread Spectrum signals are noise-like, they are hard to detect.
- Spread Spectrum signals are also hard to Intercept or demodulate.
- Further, Spread Spectrum signals are harder to jam (interfere with) than narrowband signals.
- These Low Probability of Intercept (LPI) and anti-jam (AJ) features are why the military has used Spread Spectrum for so many years.
- Spread signals are intentionally made to be much wider band than the information they are carrying to make them more noise-like.
- Spread Spectrum signals use fast codes that run many times the information bandwidth or data rate.
- These special "Spreading" codes are called "Pseudo Random" or "Pseudo Noise" codes. They are called "Pseudo" because they are not real Gaussian noise.

What Spread Spectrum Does?

- The use of these special pseudo noise codes in spread spectrum (SS) communications makes signals appear wide band and noise-like.
- It is this very characteristic that makes SS signals possess the quality of Low Probability of Intercept.
- SS signals are hard to detect on narrow band equipment because the signal's energy is spread over a bandwidth of maybe 100 times the information bandwidth.
- Processing gain is essentially the ratio of the RF bandwidth to the information bandwidth.
- A typical commercial direct sequence radio, might have a processing gain of from 11 to 16 dB, depending on data rate.
- It can tolerate total jammer power levels of from 0 to 5 dB stronger than the desired signal.

General Model of Spread Spectrum System



- Input data is fed into a channel encoder Which Produces analog signal with narrow bandwidth.
- Signal is further modulated using sequence of digits.
- Spreading code or spreading sequence Generated by pseudo noise, or pseudo-random number generator.
- Effect of modulation is to increase bandwidth of signal to be transmitted.
- On receiving end, digit sequence is used to demodulate the spread spectrum signal.
- Signal is fed into a channel decoder to recover data.

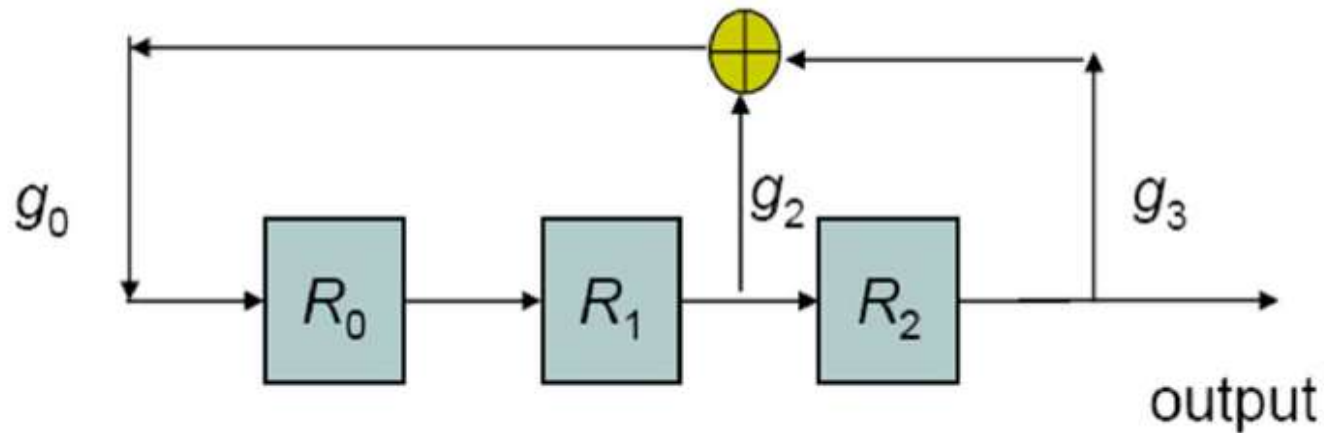
Processing Gain

- An important characteristic of a spread spectrum system is the processing gain (G_p).
- It is defined as the ratio of the spread bandwidth (BW) to the bandwidth of the information (R), which can be written as follows:
- $G_p = \text{Bandwidth of Spread Signal (BW)} / \text{Bandwidth of digital information signal (R)}$
- It's a measure of difference between the system performance when using Spread Spectrum techniques & the system performance when not using Spread Spectrum techniques.
- Processing gain of Spread Spectrum is also known as Bandwidth Expansion Factor.
- This is because it gives the factor, by which the bandwidth of digital message signal is expanded.

Pseudo noise (PN) Sequences:

- PN sequences are codes that appear to be random but are not, strictly speaking, random since they are generated using predetermined circuit connections.
- PN sequences are usually generated using Linear Feedback Shift Registers.
- The length of the PN sequence depends on the number of shift register stages.
- If there are n shift registers used, the maximum possible PN sequence length, p is given as in below equation:
$$p = 2^n - 1$$
- Such a sequence are also referred to as a maximal length sequence.

- As an example consider the case of a 3-stage LFSR used to generate PN sequences.



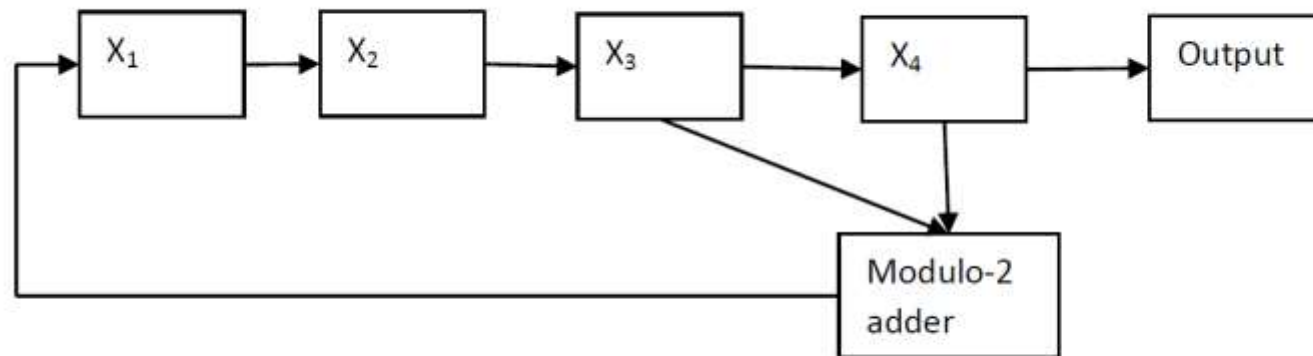


Figure .2 Connections in a maximal length sequence 4-stage LFSR

The initial states of X1, X2, X3 and X4 can be any value but 0000 respectively. An initial state of 0000 would lock the output to 0. Assuming that the initial state was 1000, the output sequence for 15 clock pulses is:

0 0 0 1 0 0 1 1 0 1 0 1 1 1 1

After 15 clock pulses the same sequence would again repeat.

Properties of PN sequences:

- In order for PN sequences to be considered random they exhibit a number of randomness properties.

Balance Property:

- In the balance property, the number of output binary ones and the number of binary output zeros in a single period differs by at most one.

Run Length Property:

- A run is defined as a continuous sequence of the same type of binary digit.
- A new run commences with the appearance of a different binary digit.
- The length of a run is the number of digits contained in a given run.
- In PN sequences about half of the runs are of length 1, about a quarter of the runs are of length 2, about an eighth of the runs are of length three and so on.

Correlation Property:

- Based on the correlation property if any PN output sequence is compared with any cyclic shift of itself, the number of agreements differs from the number of disagreements by at most one count.
- Therefore, if the cross correlation is done for different shifts, there will be maximum correlation when there is no shift and minimum correlation when the cyclic shift is one or more.
- The correlation property makes synchronization easier since during synchronization, by correlating the transmitter PN sequence with the receiver PN sequence the receiver PN sequence can be continuously delayed until a set threshold of the correlation under which acquisition can be declared is attained.

m-Sequences

These codes (DSSS codes) will all be treated as pseudo noise (PN) sequences because it resembles random sequences of bits with a flat noise like spectrum.

This sequence appears to have a random pattern but in fact can be recreated by using the shift register structure in Figure 4 with $M=4$, polynomial $x^4 + x^3 + 1$ and initial state '1 1 0 0'.

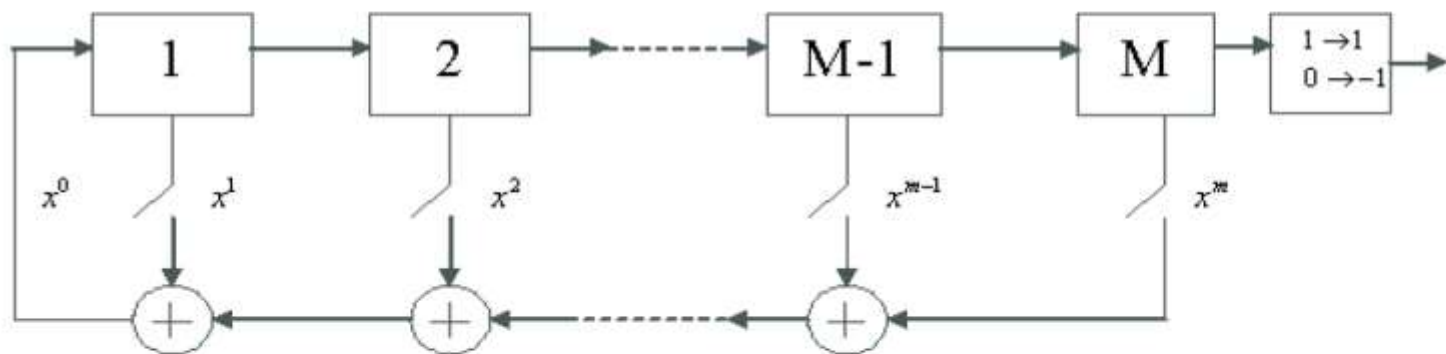


Figure 5: Shift register structure for m-sequence

Where ' $\oplus$ ' represents modulo 2 addition.

Gold Sequences

Gold sequences help generate more sequences out of a pair of m-sequences giving now many more different sequences to have multiple users. Gold sequences are based on preferred pairs m-sequences. For example, take the polynomials $1+x^2+x^5$ and $1+x+x^2+x^4+x^5$:

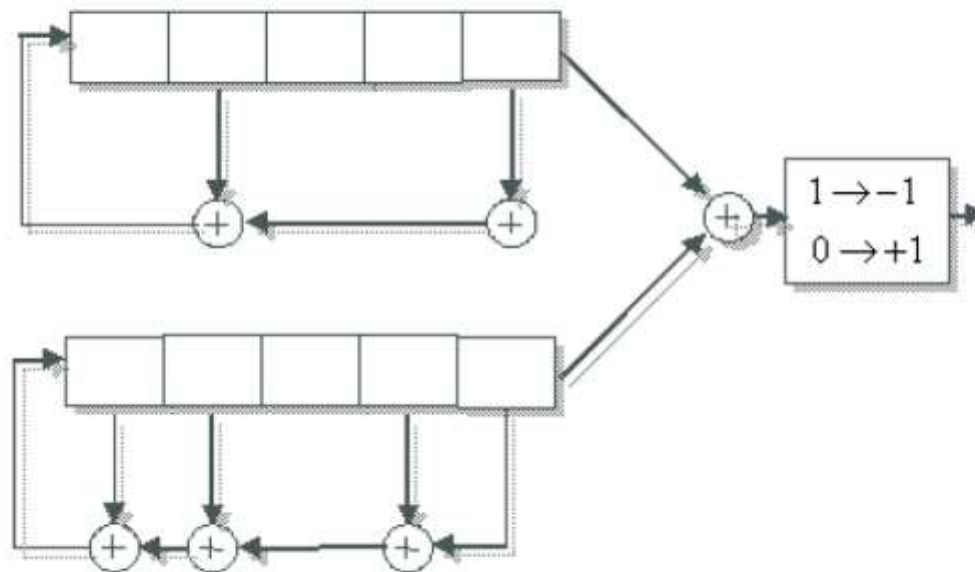


Figure 11: Example of gold sequence generator using one preferred pair of m-sequences: $1+x^2+x^5$ and $1+x+x^2+x^4+x^5$

Remember m-sequences gave only one sequence of length $2^5 - 1$. By combining two of these sequences, we can obtain up to 31 ($2^5 - 1$) plus the two m-sequences themselves, generate 33 sequences (each one length $2^5 - 1$) that can be used to spread different input messages (different users CDMA).

The m-sequence pair plus the $2^m - 1$ Gold sequences form the $2^m + 1$ available sequences to use in DSSS. The wanted property about Gold codes is that they are balanced (i.e. same number of 1 and -1s).

Walsh Hadamard Codes

Other common sequences are Walsh-Hadamard sequences currently used in CDMA systems. These sequences are orthogonal (i.e. $\sum b_i b_j = 0$ where b is a row of the matrix), convenient properties for multiple users. The sequences are the rows of the Hadamard matrix H_M defined for $M = 2$ as:

$$H_2 = \begin{bmatrix} +1 & +1 \\ +1 & -1 \end{bmatrix}$$

For larger matrices use the recursion:

$$H_{2M} = \begin{bmatrix} H_M & H_M \\ H_M & -H_M \end{bmatrix}$$

$$H_4 = \begin{bmatrix} +1 & +1 & +1 & +1 \\ +1 & -1 & +1 & -1 \\ +1 & +1 & -1 & -1 \\ +1 & -1 & -1 & +1 \end{bmatrix}$$

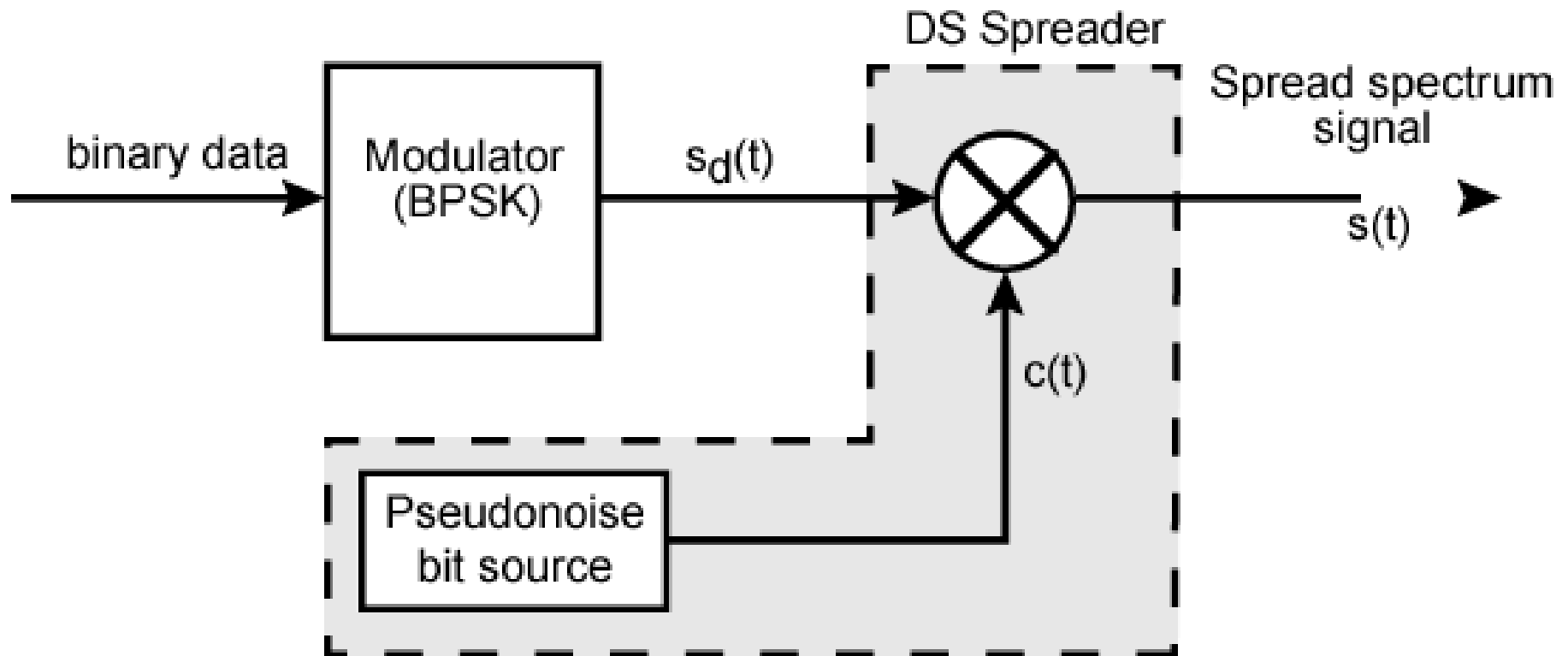
Example for

Spread Spectrum Access

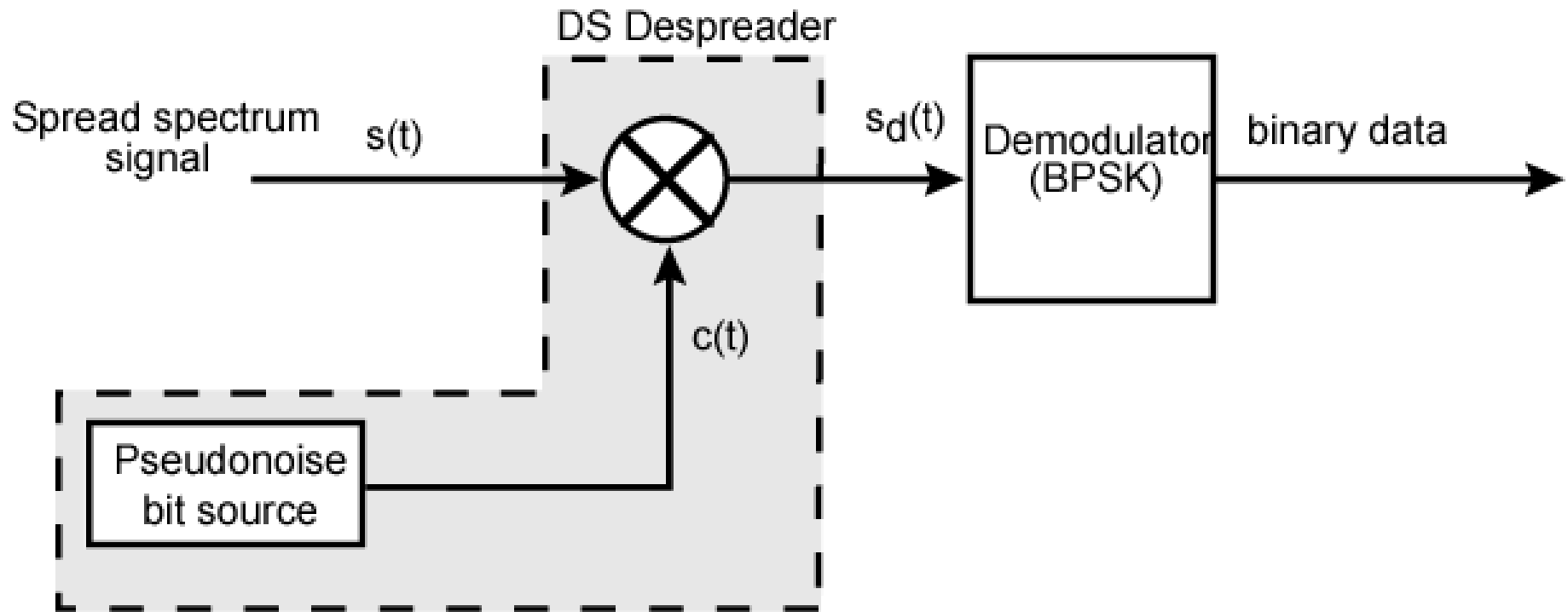
- Two techniques
 - Frequency Hopped Multiple Access (FHMA)
 - Direct Sequence Multiple Access (DSMA)
 - Also called Code Division Multiple Access – CDMA

- **CHIP:** The time it takes to transmit a bit or single symbol of a PN code.
- **CODE:** A digital bit stream with noise-like characteristics.
- **CORRELATOR:** The SS receiver component that demodulates a Spread Spectrum signal.
- **DE-SPREADING:** The process used by a correlator to recover narrowband information from a spread spectrum signal.
- **PN:** Pseudo Noise - a digital signal with noise-like properties.

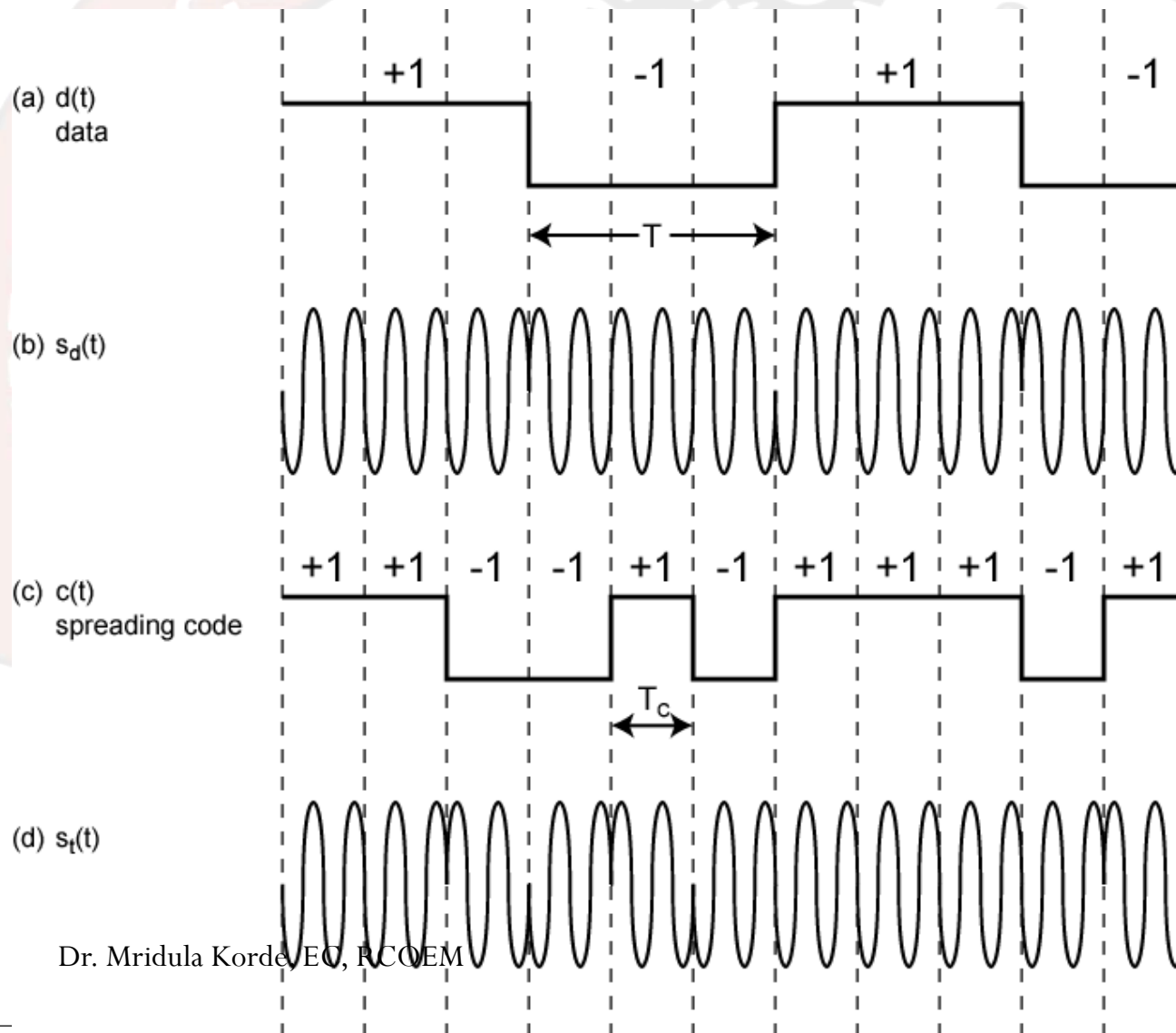
Direct Sequence Spread Spectrum Transmitter



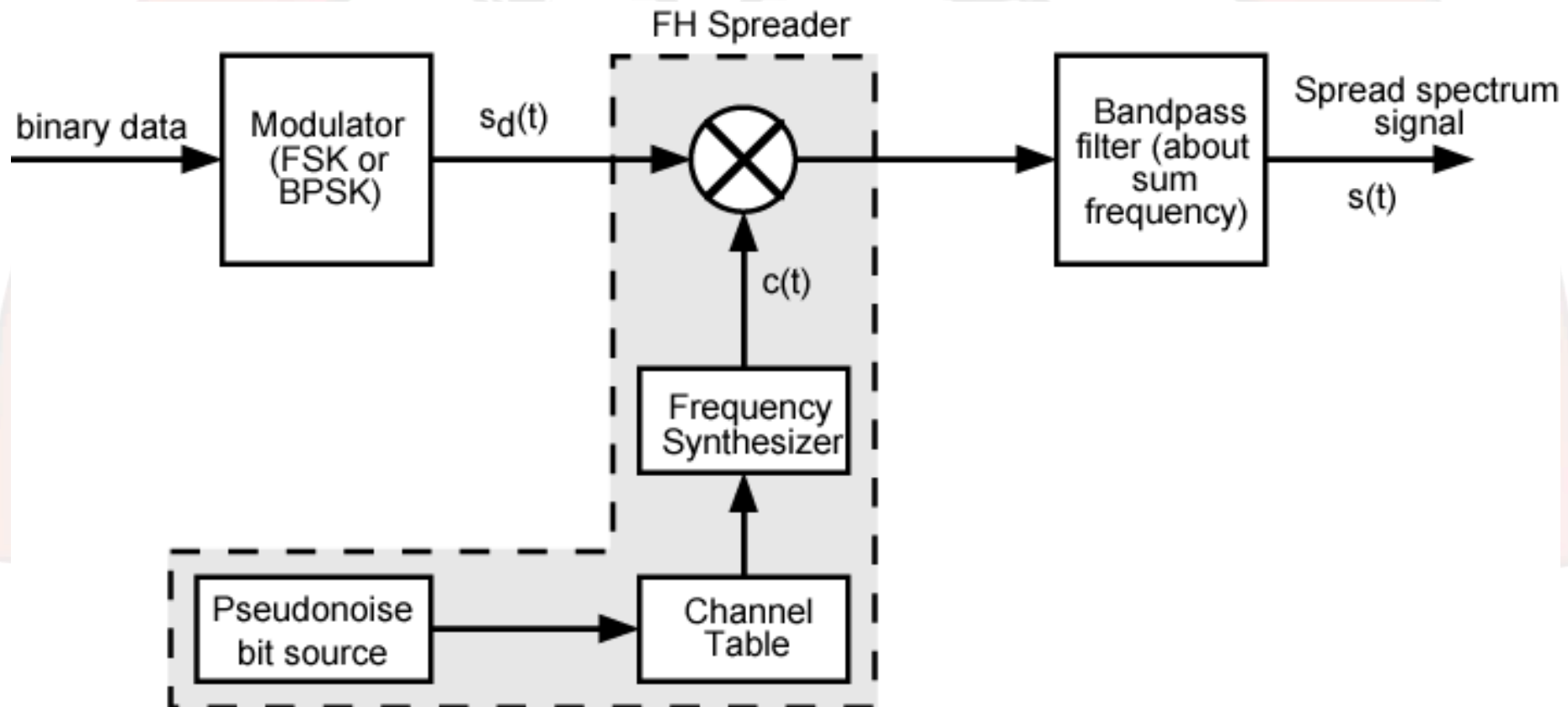
Direct Sequence Spread Spectrum Transmitter



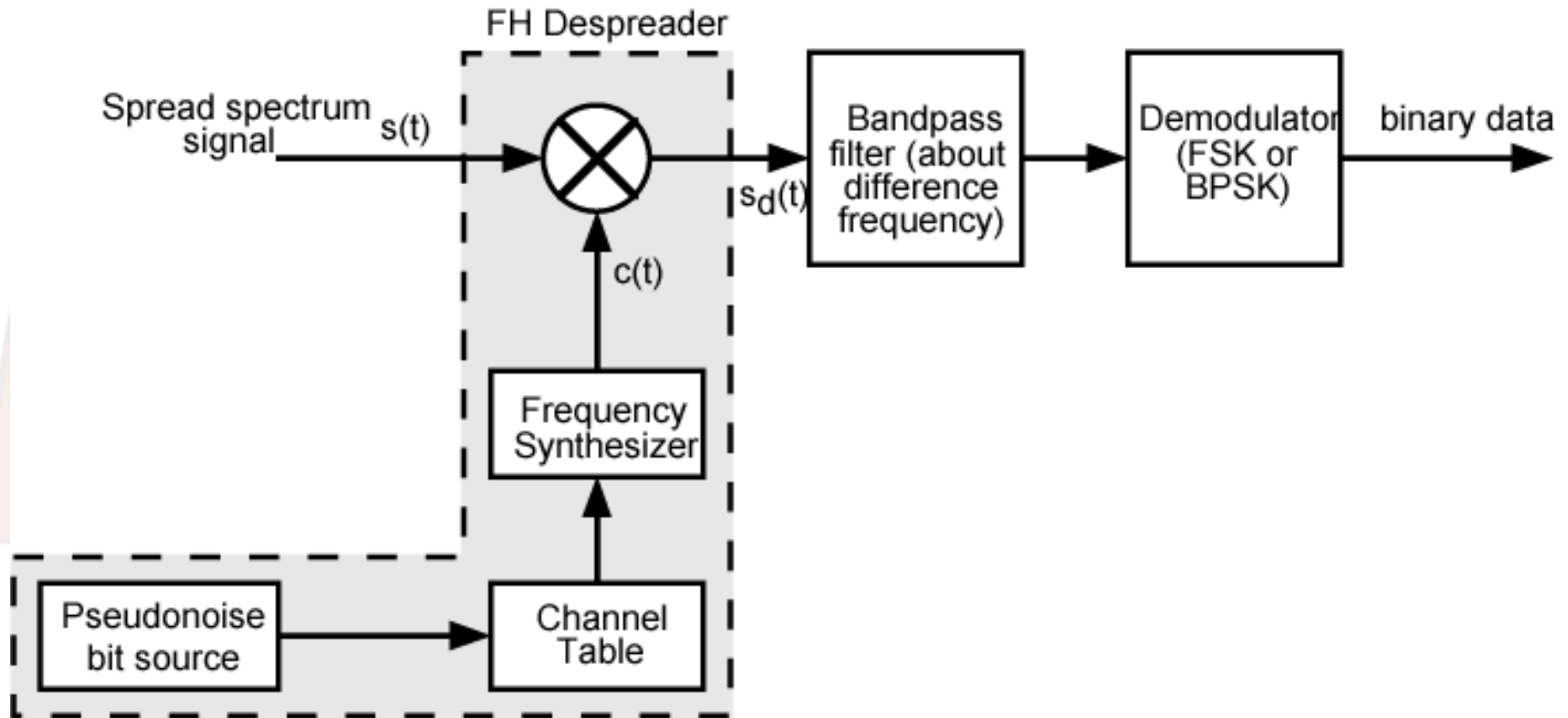
Direct Sequence Spread Spectrum Using BPSK Example



Frequency Hopping Spread Spectrum System (Transmitter)



Frequency Hopping Spread Spectrum System (Receiver)



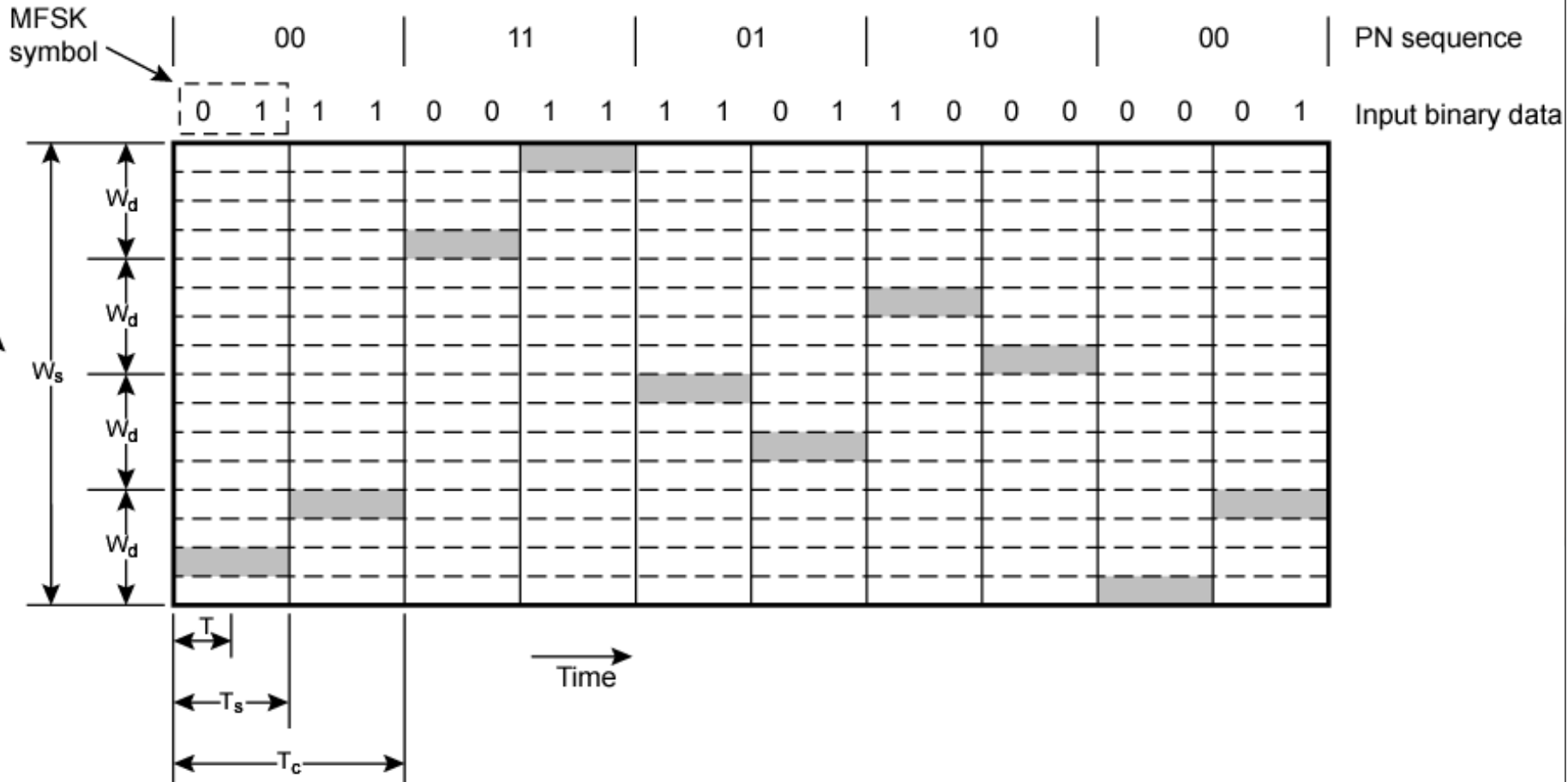
Frequency Hopping Spread Spectrum (FHSS)

- Carrier frequencies of the individual users are varied in pseudo random manner.
- FHMA allows multiple users to simultaneously occupy the same spectrum at the same time, where each user dwells at a specific narrowband channel at a particular instance of time based on particular PN code of the user.
- The pseudorandom change of the channel frequencies of the user randomizes the occupancy of specific channel at a given time, thereby allowing for multiple access over a wide range of frequencies.

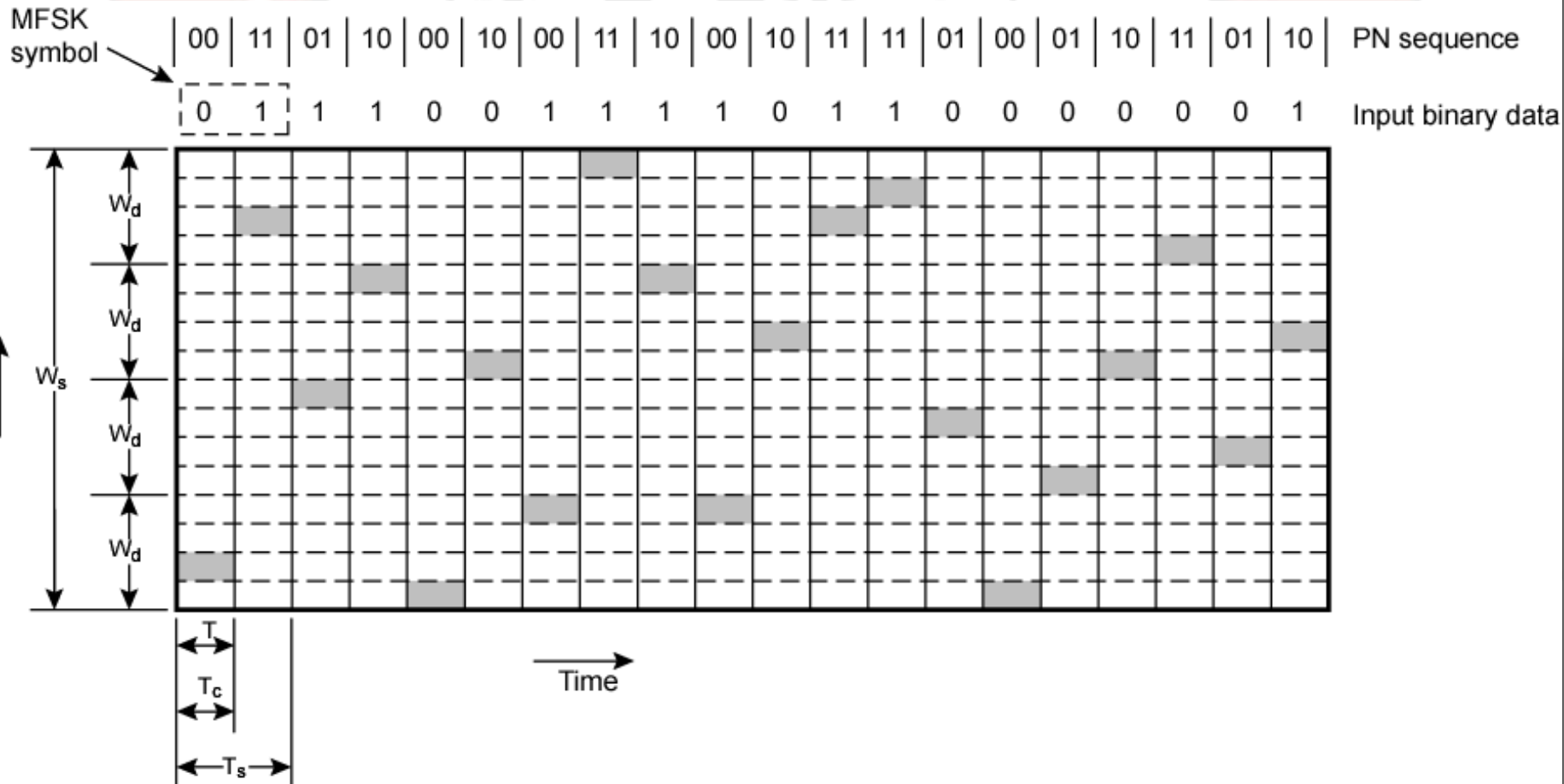
Frequency Hopping Spread Spectrum (FHSS)

- In FH receiver, locally generated PN code is used to synchronize the receiver's instantaneous frequency with that of transmitter.
- Difference between FHMA and FDMA is that the frequency hopped signal changes channels at rapid intervals.
- If the rate of change of carrier frequency is greater than the symbol rate, then the system is called as **fast frequency hopping system**.
- IF the channel changes at a rate less than or equal to the symbol rate, it is called as **slow frequency hopping system**.

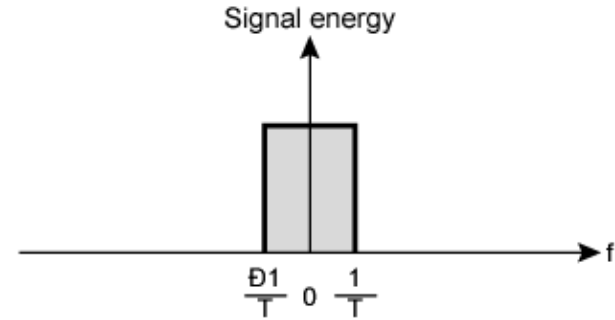
Slow Frequency Hop Spread Spectrum Using MFSK ($M=4$, $k=2$)



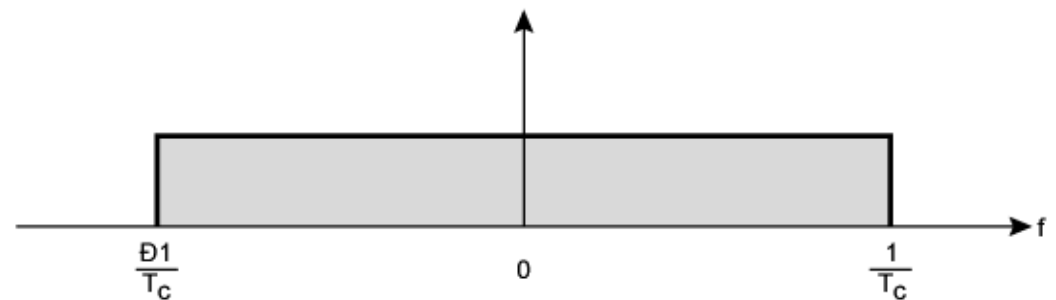
Fast Frequency Hop Spread Spectrum Using MFSK (M=4, k=2)



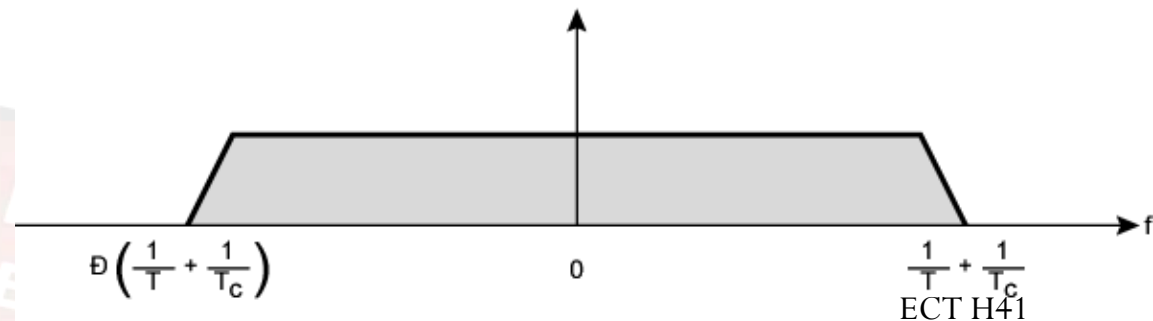
Approximate Spectrum of DSSS Signal



(a) Spectrum of data signal



(b) Spectrum of pseudonoise signal



(c) Spectrum of combined signal

Advantages of direct sequence symbols

1. This system has best noise and antijam performance.
2. Unrecognized receivers find it most difficult to detect direct sequence signals.
3. It has best discrimination against multipath signals.

Disadvantages of direct sequence systems

1. It requires wideband channel with small phase distortion.
2. It has long acquisition time.
3. The pseudo-noise generator should generate sequence at high rates.
4. This system is distance relative.

Advantages of frequency hopping system

1. These systems bandwidth (spreads) are very large.
2. They can be programmed to avoid some portions of the spectrum.
3. They have relatively short acquisition time.
4. The distance effect is less.

Disadvantages of frequency hopping systems

1. Those systems need complex frequency synthesizers.
2. They are not useful for range and range-rate measurement.
3. They need error correction.

Comparison of DSSS and FHSS

Sr. No.	Parameter	Direct sequence spread spectrum	Frequency hop spread spectrum
1	Definition	PN sequence of large bandwidth is multiplied with narrowband data signal.	Data bits are transmitted in different frequency slots which are changed by PN sequence.
2	Spectrum of signal	Data sequence is spread over entire bandwidth of spread spectrum signal.	Data sequence is spread over small frequency slots of the spread spectrum signal.
3	Chip rate R_c	Chip rate is fixed. It is the rate at which bits of PN sequence occur. $R_c = \frac{1}{T_c}$	Chip rate is maximum of hop rate or symbol rate. $R_c = \max(R_h, R_s)$

4	Modulation technique	Normally uses BPSK modulation.	Normally uses M-ary FSK modulation.
5	Processing gain	$PG = \frac{T_b}{T_c} = N$	$PG = 2^t$, here t is the bits in PN sequence.
6	Probability of error	$P_e = \frac{1}{2} \operatorname{erfc} \sqrt{\frac{E_b}{JT_c}}$	$P_e = \frac{1}{2} e^{-\gamma_b R_c/2}$ Here $\gamma_b = \frac{E_b}{J_0}$
7	Effect of distance	This system is distance relative.	Effect of distance is less in this system.
8	Acquisition time	Acquisition time is long.	Acquisition time is short.

Benefits of Spread Spectrum Communications:

- Just like a road in the movement of vehicles bandwidth is an important resource in communication.
- In fact, it is more desirable to transmit more information using limited bandwidth.
- This begs the question why spread spectrum communications is employed yet they appear to contradict this basic tenet of economizing a limited resource.
- There is no doubt that spreading the spectrum of transmitted information has some benefits.
- From ensuring secure military communication to its applications in commercial mobile telephony, the following are some of the benefits of using a spread spectrum communication system:

Benefits of Spread Spectrum Communications:

Interference suppression:

- For a signal with a bandwidth W and a duration T the total number of dimensions is $2WT$.
- Spreading offers no performance improvement in the case of white Gaussian noise.
- However, in the case of a jammer with finite power, there is a remarkable improvement in the performance.

Low power spectral density:

- This makes it possible to achieve hiding of the signal and is the basis for low probability of intercept (LPI) communications.
- It also makes it possible to conform to acceptable regulator transmission power spectral density levels.

Multiple access:

- Spread spectrum techniques offer a platform for simultaneous system access by several users in a coordinated manner thus allowing channel sharing.

Alleviation of multipath propagation:

- The good performance of spread spectrum systems in multipath channels enables them to be useful in fading channels.
- Both transmit as well as receive antenna diversity is possible with spread spectrum communication systems.

Applications of Spread Spectrum Modulation

1. Antijamming capacity for military applications

The spread spectrum has the ability to resist the effect of intentional jamming. Previously this antijam capability was used in military applications. Some commercial applications also use spread spectrum because of its antijam capability.

2. Low probability of intercept

Low probability of intercept is an application of spread spectrum in military. In this case, the signal spectral density is kept small such that the presence of the signal is not detected easily.

4. Secured communication

The spread spectrum communications are 'secure'. Since pseudo-noise codes are used to generate spread spectrum signals, unwanted receivers cannot recognize the spread spectrum signals. The secrecy capability of spread spectrum is used in military as well as in many commercial applications.

7. CDMA communication

Spread spectrum is used in code division multiple access (CDMA) communication. In this system many users communicate at the same time and through the same channel. Each user is allotted a particular code. The signals are separated at the receiver because of these codes. Those codes are allotted by pseudo-noise sequence. The important advantage of CDMA is that the message of some other user is not decoded. And number of users can communicate simultaneously through the same channel.

Code Division Multiple Access (CDMA)

- In CDMA, the narrowband message signal is *multiplied* by a very large bandwidth signal called spreading signal (code) before modulation and transmission over the air. This is called spreading.
- CDMA is also called DSSS (Direct Sequence Spread Spectrum). DSSS is a more general term.
- Message consists of symbols
 - Has symbol period and hence, symbol rate

Code Division Multiple Access (CDMA)

- Spreading signal (code) consists of chips
 - Has Chip period and hence, chip rate
 - Spreading signal use a pseudo-noise (PN) sequence (a pseudo-random sequence)
 - PN sequence is called a codeword
 - Each user has its own cordword
 - Codewords are orthogonal. (low autocorrelation)
 - Chip rate is oder of magnitude larger than the symbol rate.
- The receiver correlator distinguishes the senders signal by examining the wideband signal with the same time-synchronized spreading code
- The sent signal is recovered by despreading process at the receiver.

CDMA Advantages

- Low power spectral density.
 - Signal is spread over a larger frequency band
 - Other systems suffer less from the transmitter
- Interference limited operation
 - All frequency spectrum is used
- Privacy
 - The codeword is known only between the sender and receiver. Hence other users can not decode the messages that are in transit
- Reduction of multipath affects by using a larger spectrum

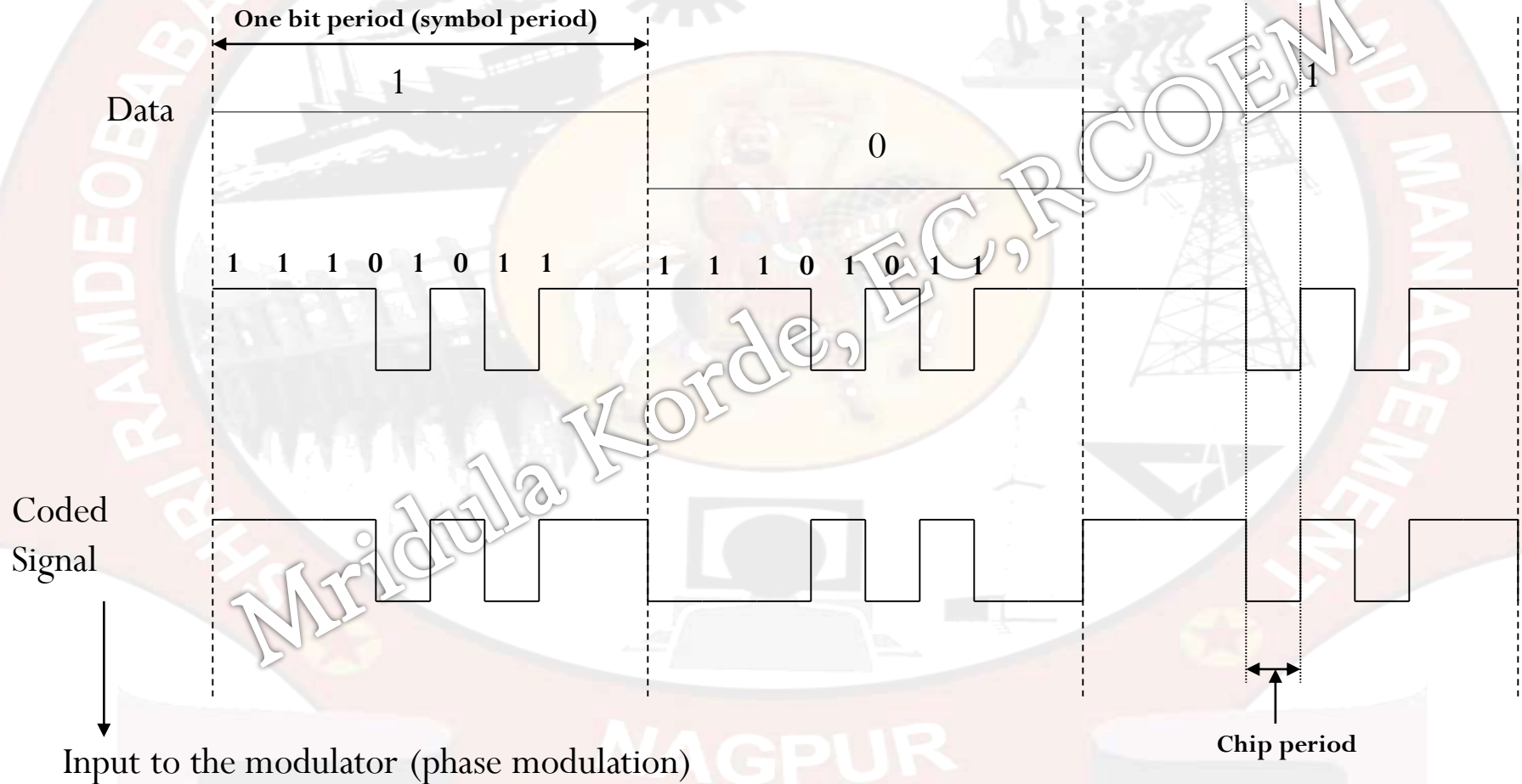
CDMA Advantages

- Random access possible
 - Users can start their transmission at any time
- Cell capacity is not concrete fixed like in TDMA or FDMA systems. Has soft capacity
- Higher capacity than TDMA and FDMA
- No frequency management
- No equalizers needed
- No guard time needed
- Enables soft handoff

CDMA Principle

Represent bit 1 with +1

Represent bit 0 with -1



Processing Gain

- Main parameter of CDMA is the processing gain that is defined as:

$$G_p = \frac{B_{spread}}{R} = \frac{B_{chip}}{R}$$

G_p : processing gain

B_{spread} : PN code rate

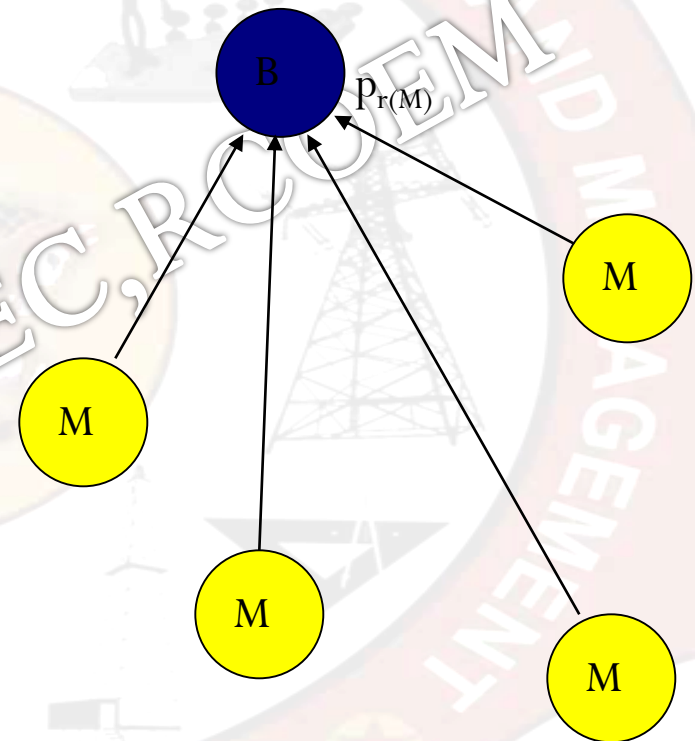
B_{chip} : Chip rate

R : Data rate

- IS-95 System (Narrowband CDMA) has a gain of 64. Other systems have gain between 10 and 100.
 - ❑ 1.228 Mhz chipping rate
 - ❑ 1.25 MHz spread bandwidth

Near Far Problem and Power Control

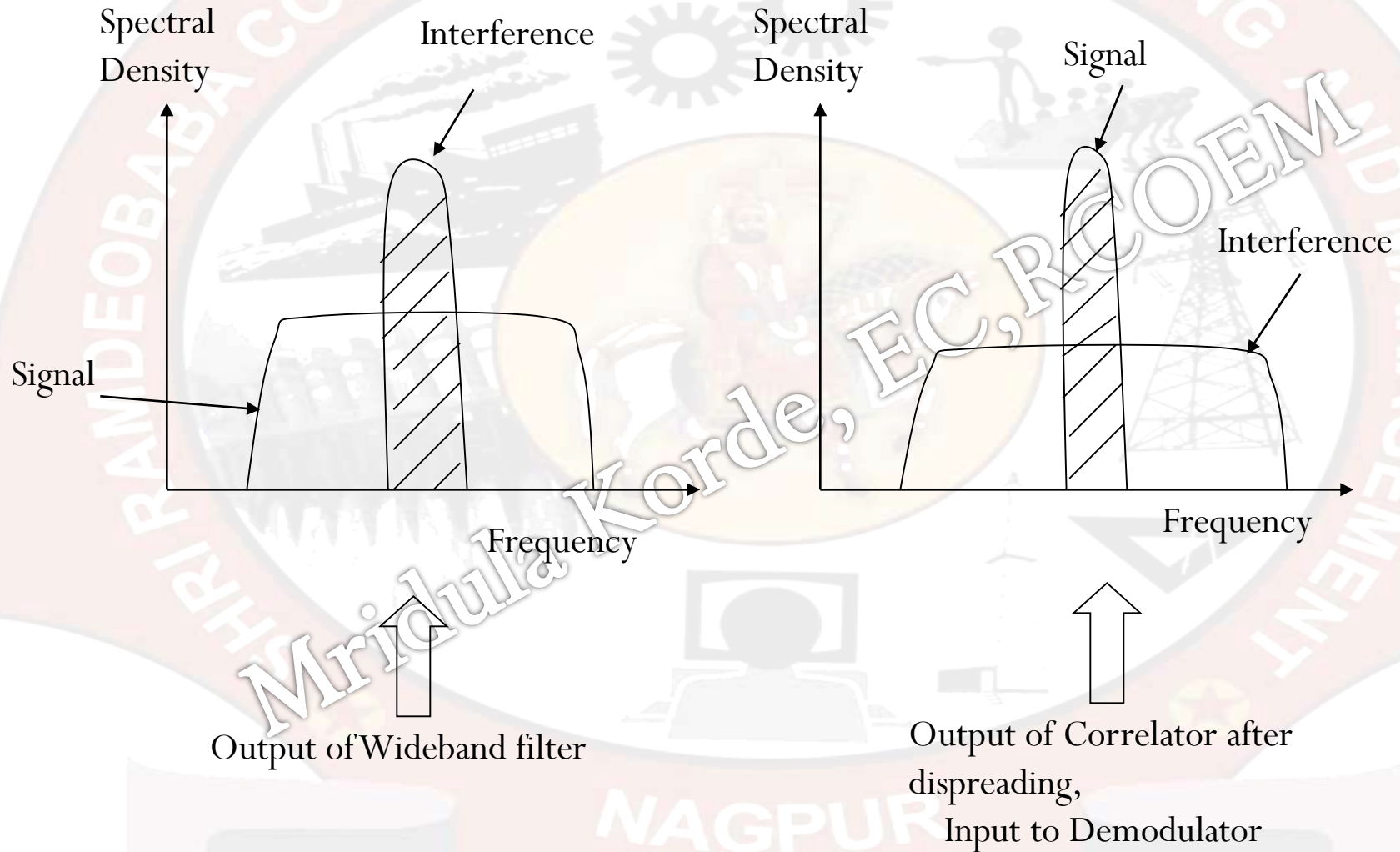
- At a receiver, the signals may come from various (multiple sources).
 - ❑ The strongest signal usually captures the modulator. The other signals are considered as noise
 - ❑ Each source may have different distances to the base station



Near Far Problem and Power Control

- In CDMA, we want a base station to receive CDMA coded signals from various mobile users at the same time.
 - Therefore the receiver power at the base station for all mobile users should be close to each other.
 - This requires power control at the mobiles.
- **Power Control:** Base station monitors the RSSI values from different mobiles and then sends power change commands to the mobiles over a forward channel. The mobiles then adjust their transmit power.

Spectra of Received Signal





**SHRI RAMDEOBABA COLLEGE OF ENGINEERING AND
MANAGEMENT,
KATOL ROAD, NAGPUR, MAHARASHTRA, INDIA**

Communication System Analysis(ECTH 41)

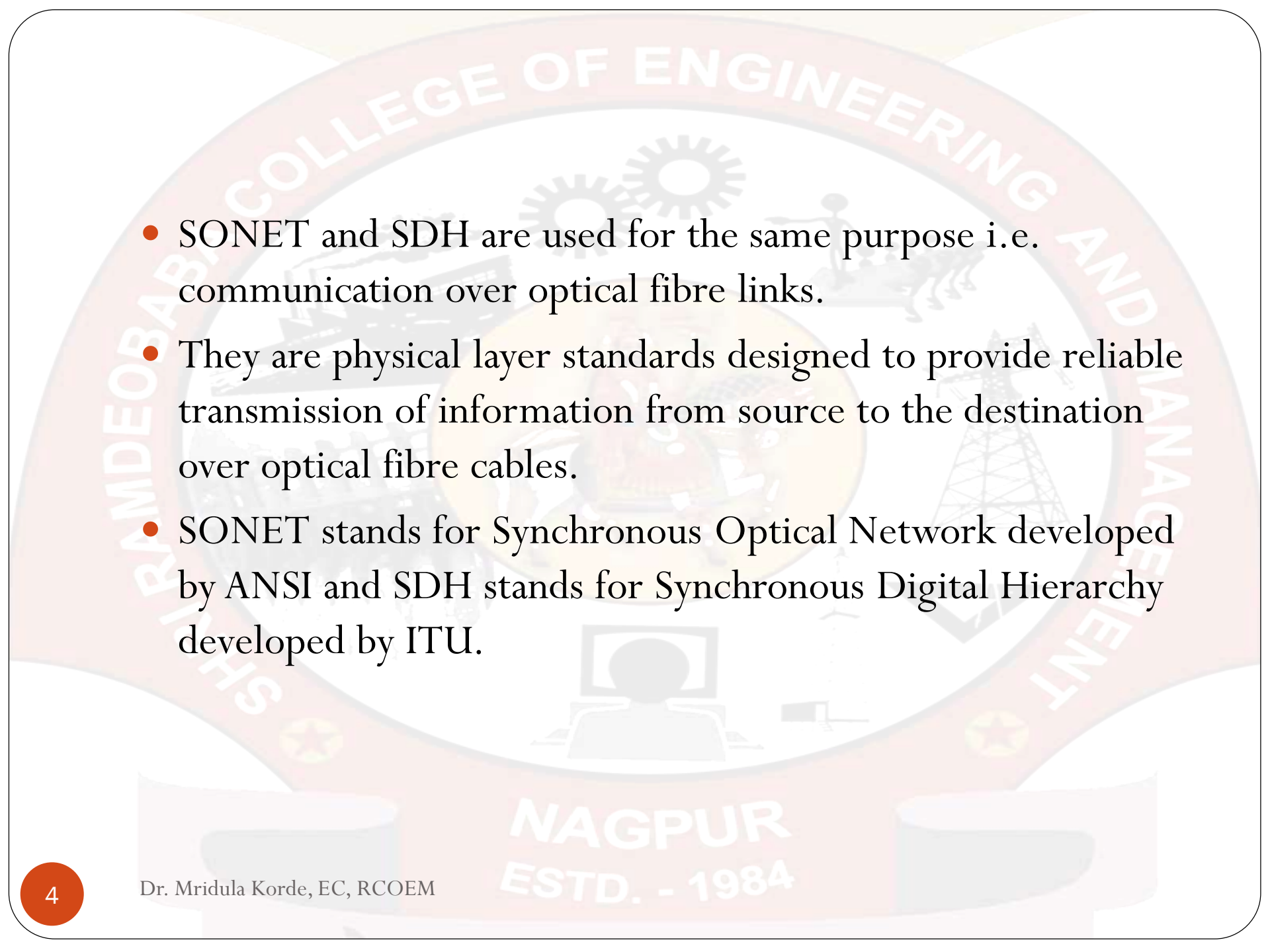
Dr. (Mrs.) Mridula Korde

Associate Professor, Department of Electronics and
Communication Engineering, RCOEM, Nagpur

SONET/SDH

Introduction to SONET

- SONET (Synchronous Optical NETwork) is a standard for optical telecommunications transport.
- It was 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.
- SDH is an ideal and particular network especially for network providers, with efficient delivery and economical network management system that can easily be adapted to accommodate the demands on BANDWIDTH for applications and services1

- 
- SONET and SDH are used for the same purpose i.e. communication over optical fibre links.
 - They are physical layer standards designed to provide reliable transmission of information from source to the destination over optical fibre cables.
 - SONET stands for Synchronous Optical Network developed by ANSI and SDH stands for Synchronous Digital Hierarchy developed by ITU.

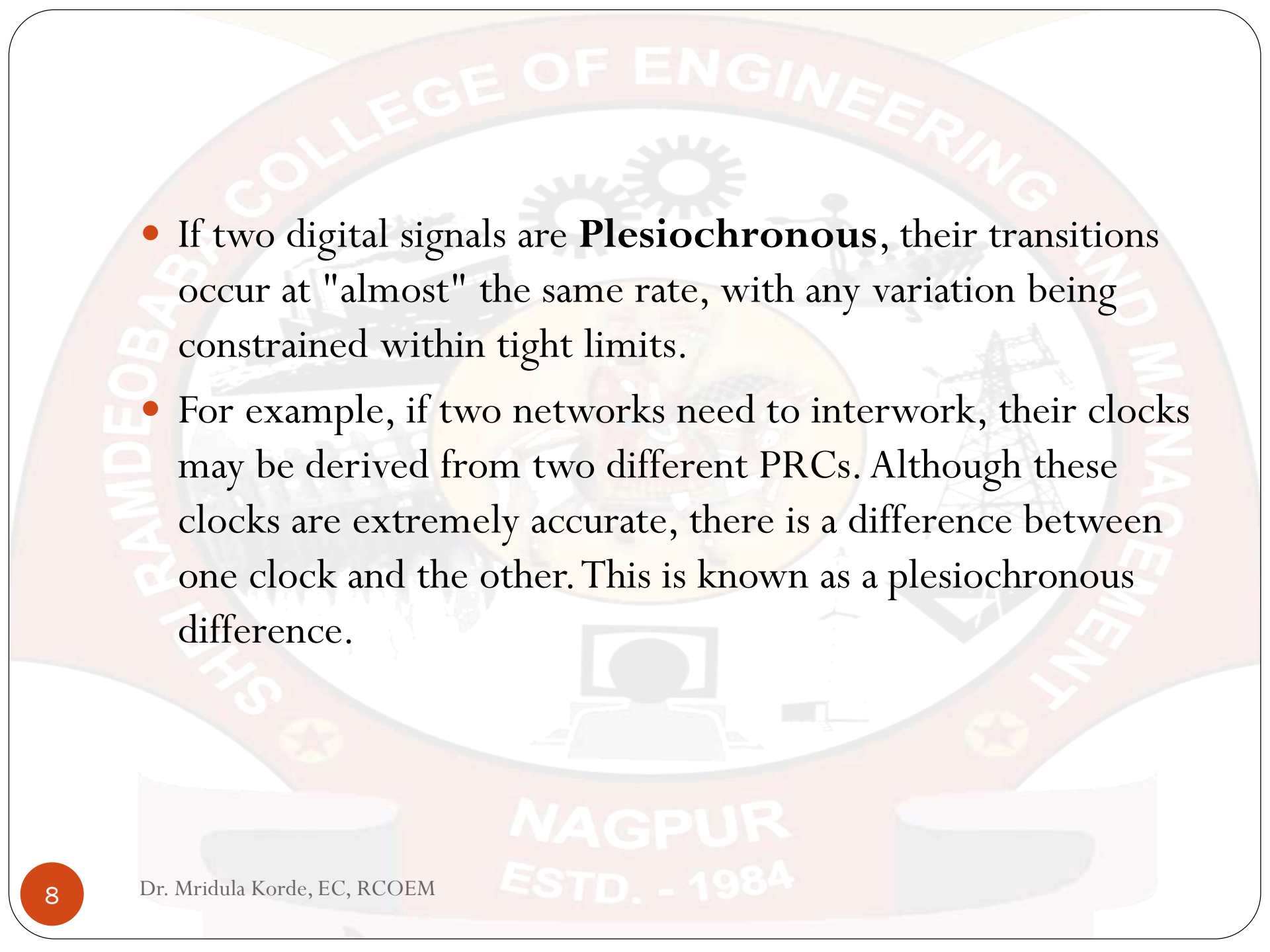
- SONET/SDH are needed to take care of following objectives in communication.
 - To carry multiple T or E carrier lines between two switching centers to take care of increased communication lines.
 - Data framing which is taken care by adopting HDLC frame format.
 - Managing information transfer between optical fibre nodes.
 - Multiplexing and de-multiplexing
 - Error checking
 - Signalling required to perform automatic switching in failure events of fibre cable or optical nodes.

Synchronization of Digital Signals

- Three Types:
 - **Synchronous**
 - **Plesiochronous**
 - **Asynchronous**

Synchronization of Digital Signals

- In a set of **Synchronous** signals, the digital transitions in the signals occur at exactly the same rate.
- There may, however, be a phase difference between the transitions of the two signals, and this would lie within specified limits.
- These phase differences may be due to propagation time delays or jitter introduced into the transmission network.
- In a synchronous network, all the clocks are traceable to one Primary Reference Clock (PRC).

- 
- If two digital signals are **Plesiochronous**, their transitions occur at "almost" the same rate, with any variation being constrained within tight limits.
 - For example, if two networks need to interwork, their clocks may be derived from two different PRCs. Although these clocks are extremely accurate, there is a difference between one clock and the other. This is known as a plesiochronous difference.

- In the case of **Asynchronous** signals, the transitions of the signals do not necessarily occur at the same nominal rate.
- Asynchronous, in this case, means that the difference between two clocks is much greater than a plesiochronous difference.
- For example, if two clocks are derived from free-running quartz oscillators, they could be described as asynchronous.
- *In a synchronous system, such as SONET, the average frequency of all clocks in the system will be the same (synchronous) or nearly the same (plesiochronous).*

SONET

- SONET can be expanded as Synchronous Optical Network and is designed mainly to use in optical network cable for transfer the data at the faster speed than the other network medium.
- SONET and SDH have been designed and implemented mostly for the same purpose and were designed for transport circuit mode communication from various sources to various destination.
- SONET and SDH is a transport protocol than it is a communication protocol.

SDH

- Since the internet connection and cell phone has been increased, the Bandwidth has been called for often for increase in reliability in connection and in high quality services.
- Above mentioned terms helps the Synchronous Digital Hierarchy and optical fiber is the medium of cable which is used mostly to transfer data from one to another.
- The main advantages of these optical fiber cables are they can transmit the data in very faster speed which no other medium can pass the data and no distraction will takes place in the data and no damages.
- The main disadvantages of this optical fiber cable, cost of installation is very high.

- The displayed figure.1 shows where SONET and SDH are placed in core networking areas and how the data is transmitted over the Optical fiber1,2

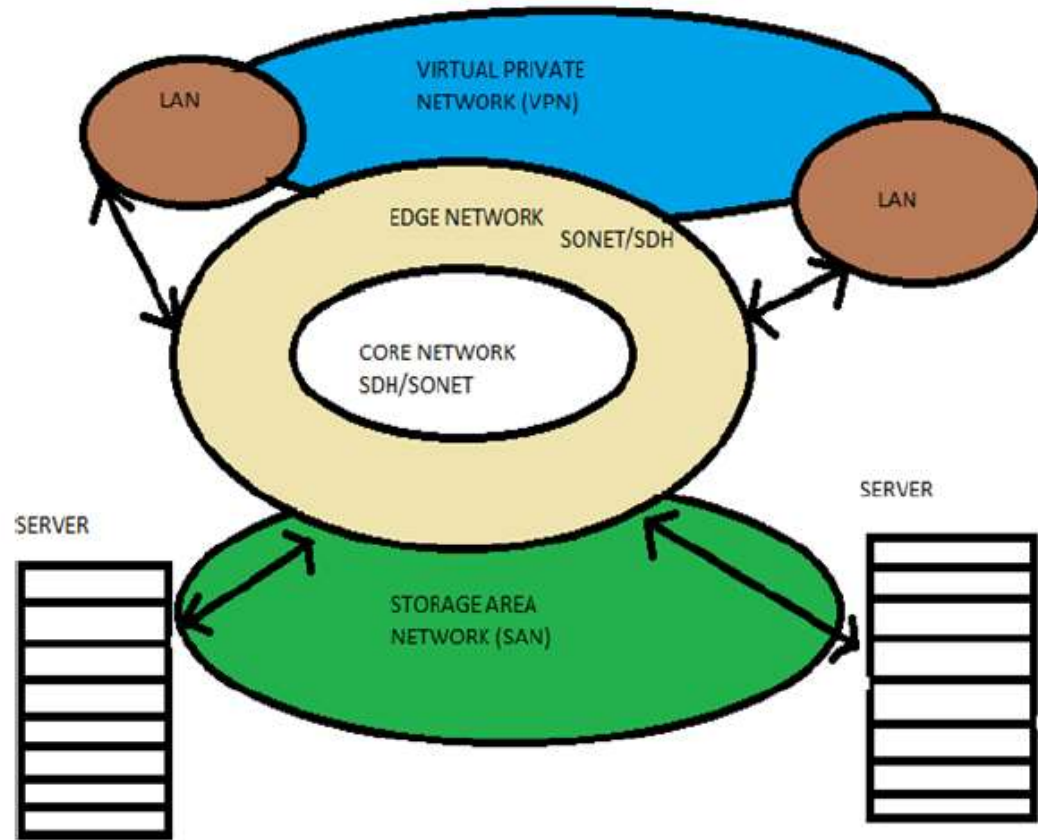


Figure 1. Diagrammatic Representation of SONET

SONET/SDH Architecture

- SONET devices: STS multiplexer/demultiplexer, regenerator, add/drop multiplexer, terminals

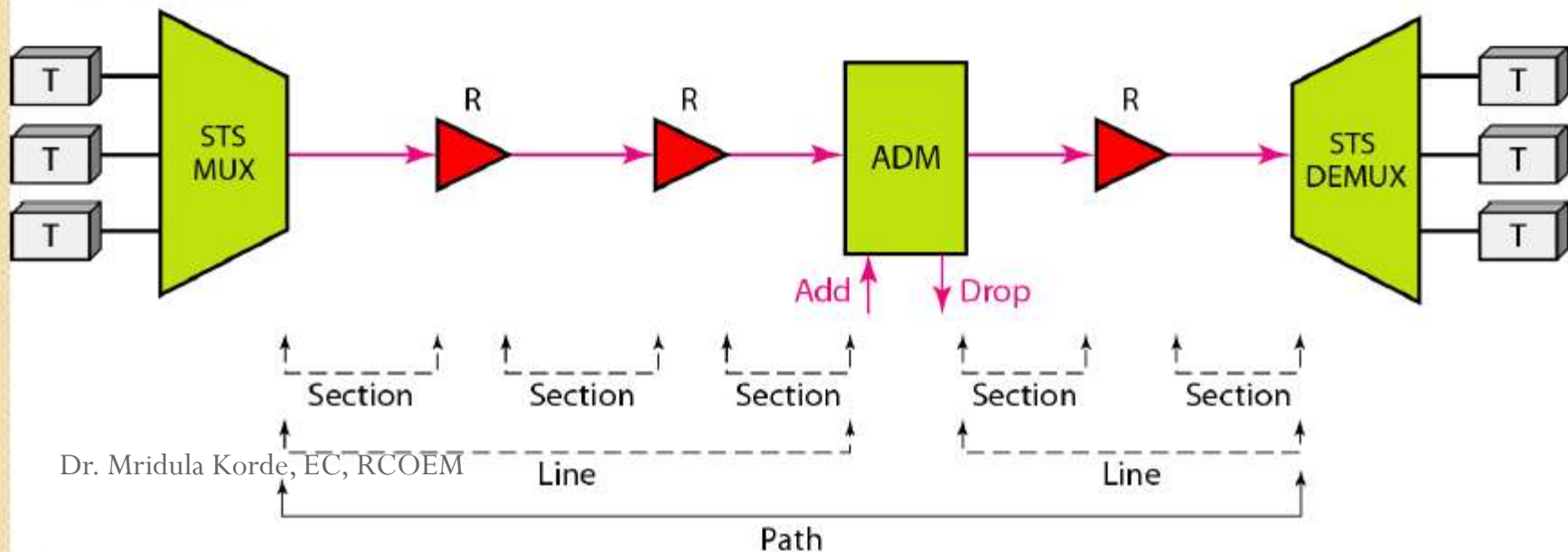
ADM: Add/drop multiplexer

STS MUX: Synchronous transport signal multiplexer

STS DEMUX: Synchronous transport signal demultiplexer

R: Regenerator

T: Terminal



SONET/SDH Architecture

- **Connections:** SONET devices are connected using *sections, lines, and paths*
- **Section:** optical link connecting two neighbor devices: mux to mux, mux to regenerator, or regenerator to regenerator
- **Lines:** portion of network between two multiplexers
- **Paths:** end-to-end portion of the network between two STS multiplexers

Network Elements of SONET

Terminal Multiplexor

- Terminal multiplexor which concentrates tributary DS-1 signals as well as other signals which derives it and transforms the electric signal in optical signal and vice versa. Simplest of the SONET links are made by two ends of fiber optics joined multiplexor with or without signal regenerator.

Signal Regenerator

- Signal regenerator is needed when the distance is too long between the two terminal multiplexor or the optic signal is too low. On receiving the signal, the signal regenerator closes and a header is added to the signal pattern before transmission. This way the information in the data is not affected.

Add/Drop Multiplexor (ADM)

- ADM gives access to new traffic from a particular point or implement the same, in addition it also absorbs a section of data traffic. On implementing ADM, it can download or insert into the main flow or on to other signals which can be altered.

SONET Network Configuration

Point to Point

- Point to point setup is formed with two terminals multiplexor which is connected with fiber optics cable and with the feature of using a single regenerator.

Point to Multipoint

- This architecture includes ADM network elements to the network. ADM connects intermediate network points to the network channels.

HUB Network

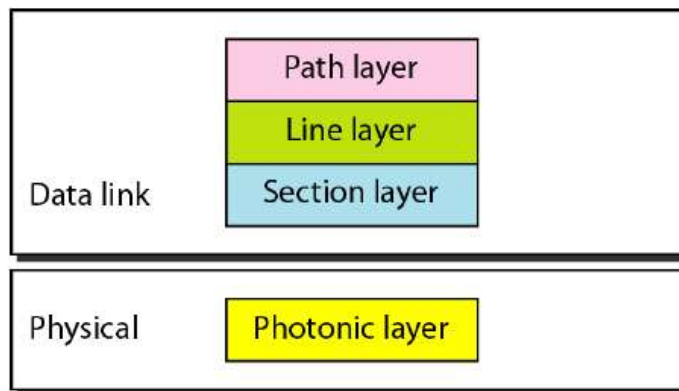
- A HUB deals with sudden data growth and network changes smoothly and efficiently over point to point networks.
- A HUB, distributes the signal traffic at the central point to various circuits.

RING Network

- The most valuable and important element of ring network is ADM.
- Many ADM can be placed in a ring structure for single way or two way data traffic.
- The ring network is advantageous because of its security.
- The working ring nodes can distribute the data traffic in case of damage to any fiber optic cable of a multiplexor.

SONET LAYERS

- SONET defines four layers: path, line, section, and photonic
- Path layer is responsible for the movement of a signal from its optical source to its optical destination
- Line layers is for the movement of a signal across a physical line
- Section layer is for the movement of a signal across a physical section, handling framing, scrambling, and error control
- Photonic layer corresponds to the physical layer of OSI model



SONET STS FRAMING FORMATS

- SONET defines its structure in Synchronous Transport Signal Levels.
- The Synchronous Transport Signal - Level 1 (STS-1) is the lowest level fundamental framing structure in SONET.
- The STS-1 frame structure is byte oriented and consist of a matrix of 810 bytes per frame.
- Currently, the fastest well-defined communication channel used in optical transmission of digital data is the SONET standard OC-768, which sends about 40 gigabits per second.

TABLE 1: SONET/SDH STANDARD DATA RATES

ELECTRICAL LEVEL	OPTICAL LEVEL	DATA RATE (MBPS)	PAYLOAD RATE (MBPS)	SDH EQUIVALENT
STS-1	OC-1	51.84	48.38	STM-0
STS-3	OC-3	155.52	149.76	STM-1
STS-12	OC-12	622.08	599.04	STM-4
STS-48	OC-48	2488.32	2396.16	STM-16
STS-192	OC-192	9953.28	9584.64	STM-64
STS-768	OC-768	39813.12	38486.02	STM-256

Network Using SONET

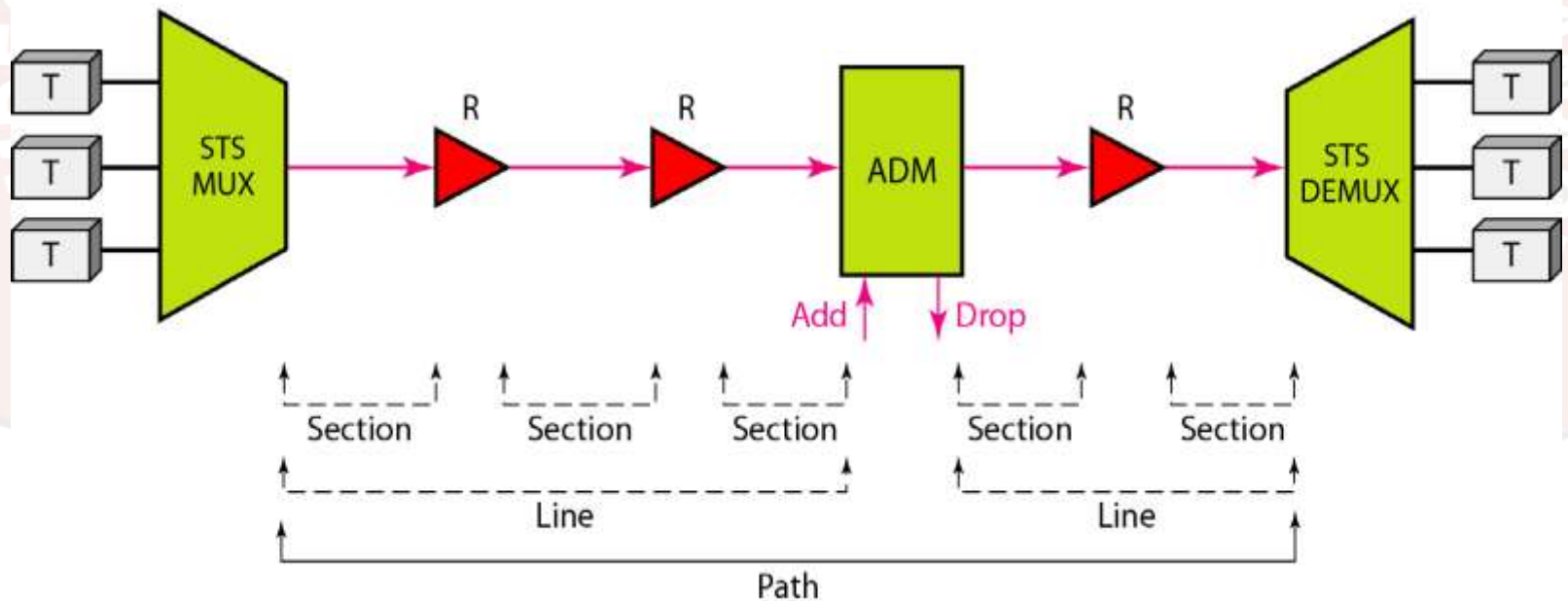
ADM: Add/drop multiplexer

STS MUX: Synchronous transport signal multiplexer

STS DEMUX: Synchronous transport signal demultiplexer

R: Regenerator

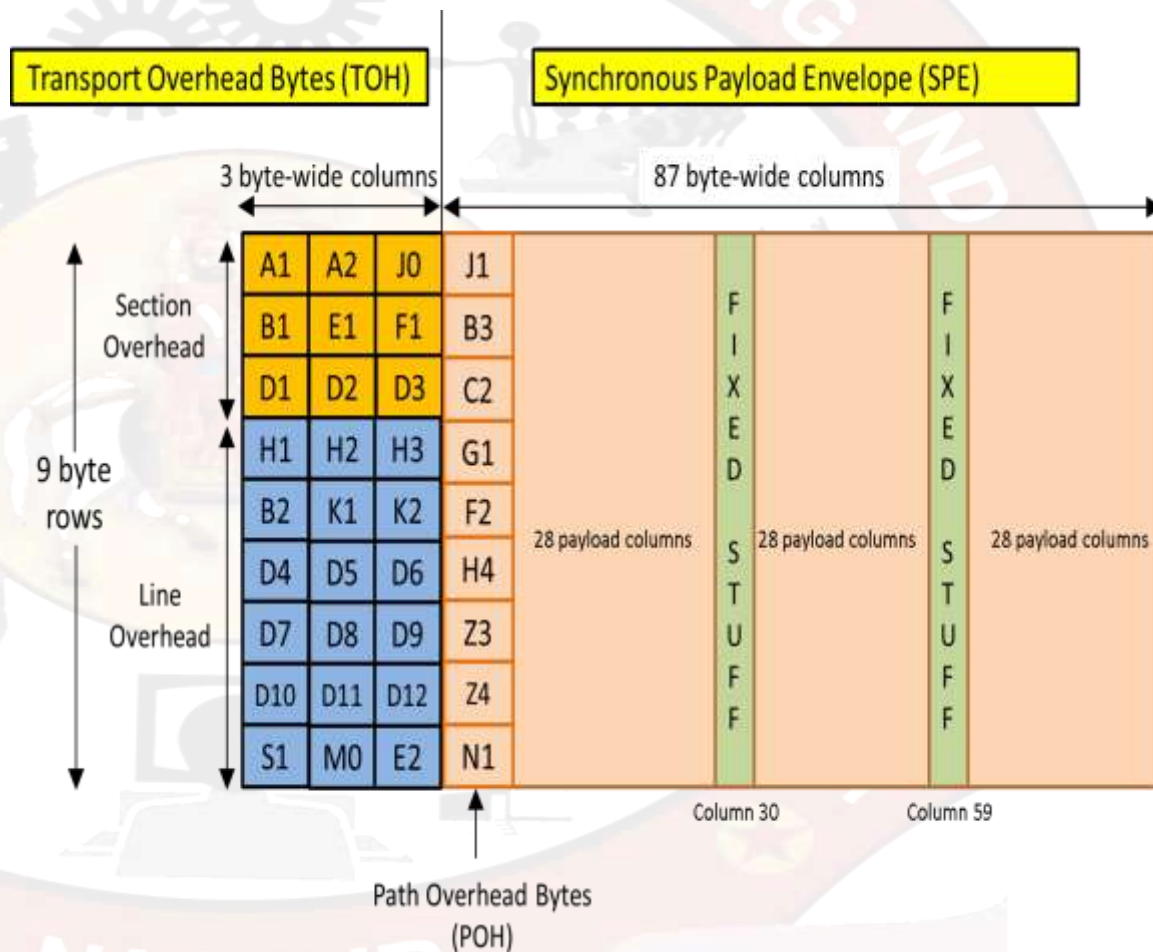
T: Terminal



- As shown in the figure.1 end to end connection between SDH/SONET networks is referred as path. Connection between major ADM nodes is referred as line. ADM stands for Add/Drop multiplexer. Connection between ADM and regenerator and also between regenerators is referred as section. Based on this following abbreviations are derived and are popularly used in SDH.
- - PTE -Path Terminating Equipment
 - LTE -Line terminating Equipment
 - STE -Section Terminating Equipment
 - ADM -Add/Drop Multiplexer

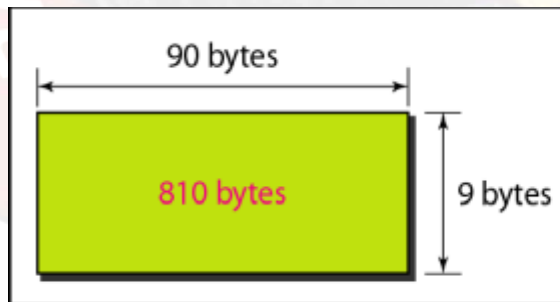
SONET STS-1 Basic Frame Structure

- Frame structure consists of 2 main components with a total of 90 bytes * 9 rows
- Transport Overhead** (TOH) - The Transport Overhead (TOH) section consists of the Section Overhead (SOH) layer and the Line Overhead (LOH) layer.
- Synchronous Payload Envelope (SPE)** - The Synchronous Payload Envelope (SPE) consists of the Path Overhead (POH) layer and the Payload.

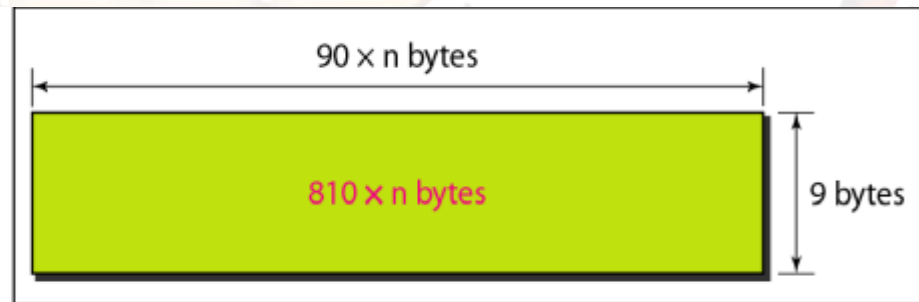


SONET FRAMES

Each synchronous transfer signal STS-n is composed of 8000 frames. Each frame is a two-dimensional matrix of bytes with 9 rows by $90 \times n$ columns.

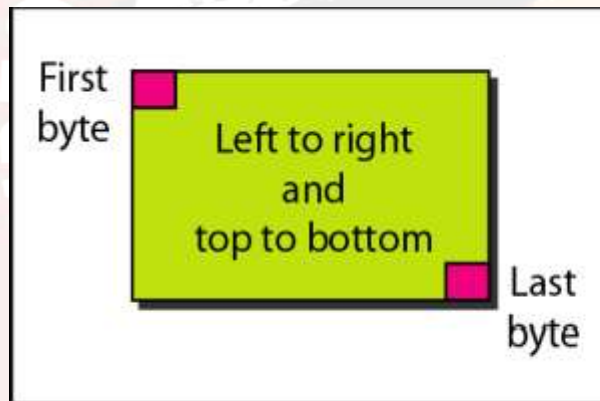


a. STS-1 frame

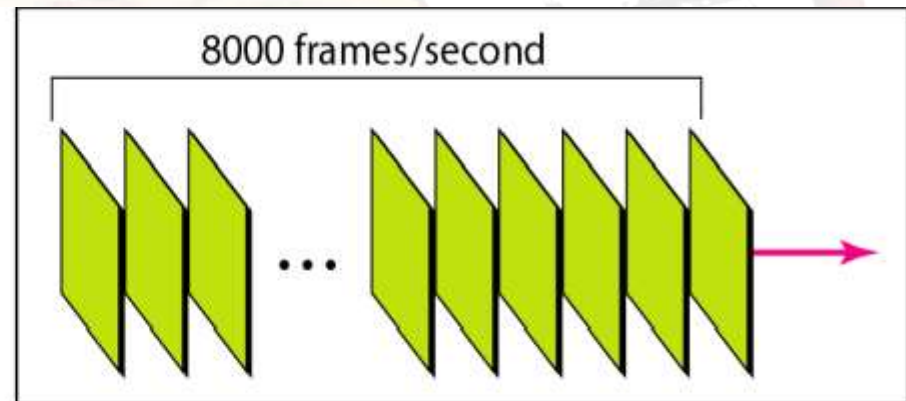


b. STS-n frame

Frame Transmission



a. Byte transmission



b. Frame transmission

Benefits of SONET / SDH

- Need for a digital transmission system faster and more sophisticated than T1 E1 systems
- Standardization
- High Speed
- Reliability
- Operations, Administration, Maintenance and Provisioning (OAM & P)
- Quality of Service (QoS)
- Flexibility
- Scalability

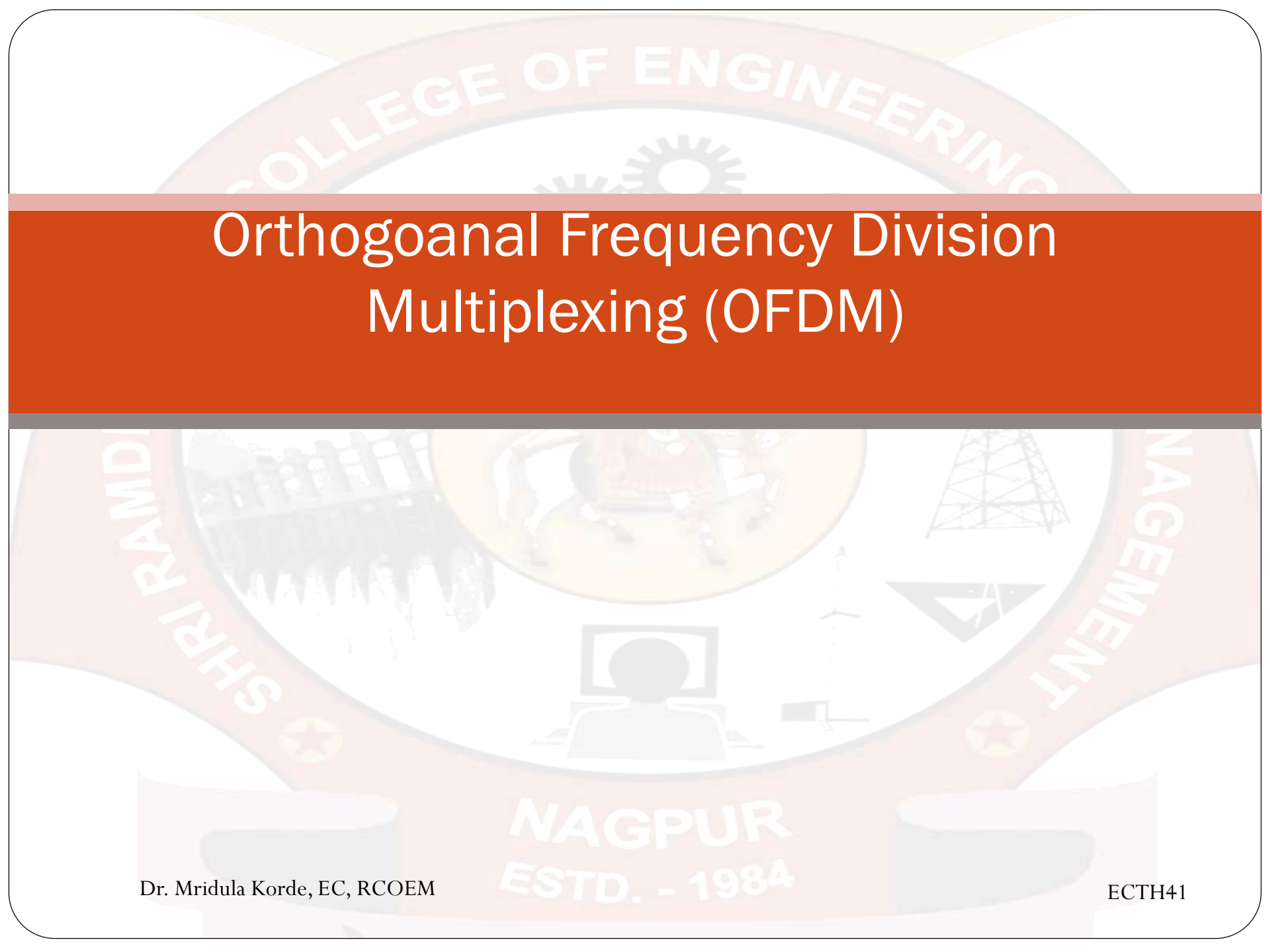


**SHRI RAMDEOBABA COLLEGE OF ENGINEERING AND
MANAGEMENT,
KATOL ROAD, NAGPUR, MAHARASHTRA, INDIA**

Communication System Analysis(ECTH 41)

Dr. (Mrs.) Mridula Korde

Associate Professor, Department of Electronics and
Communication Engineering, RCOEM, Nagpur



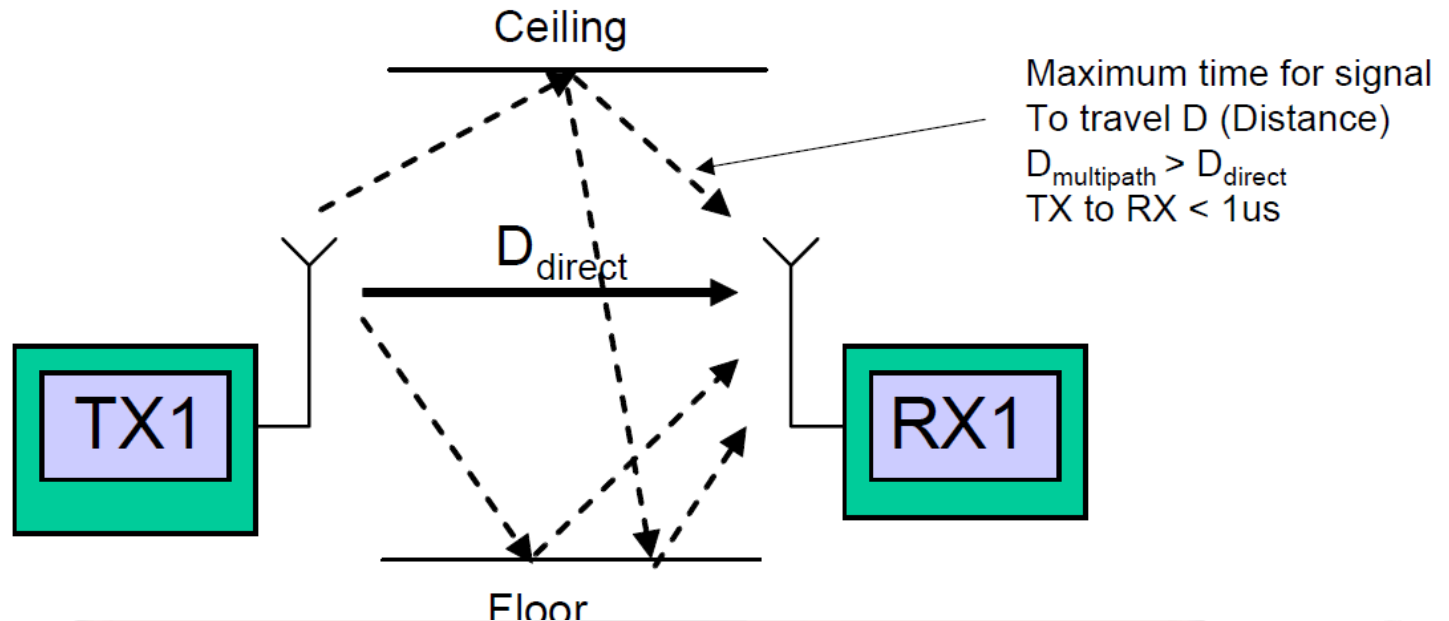
Orthogonal Frequency Division Multiplexing (OFDM)

The Multi-Path Problem

Example: Bluetooth Transmitter & Receiver

Symbol Rate = 1MSymbols/s
Symbol Duration = $1/1\text{E}6 = 1\mu\text{s}$

Maximum Symbol Delay < $1\mu\text{s}$



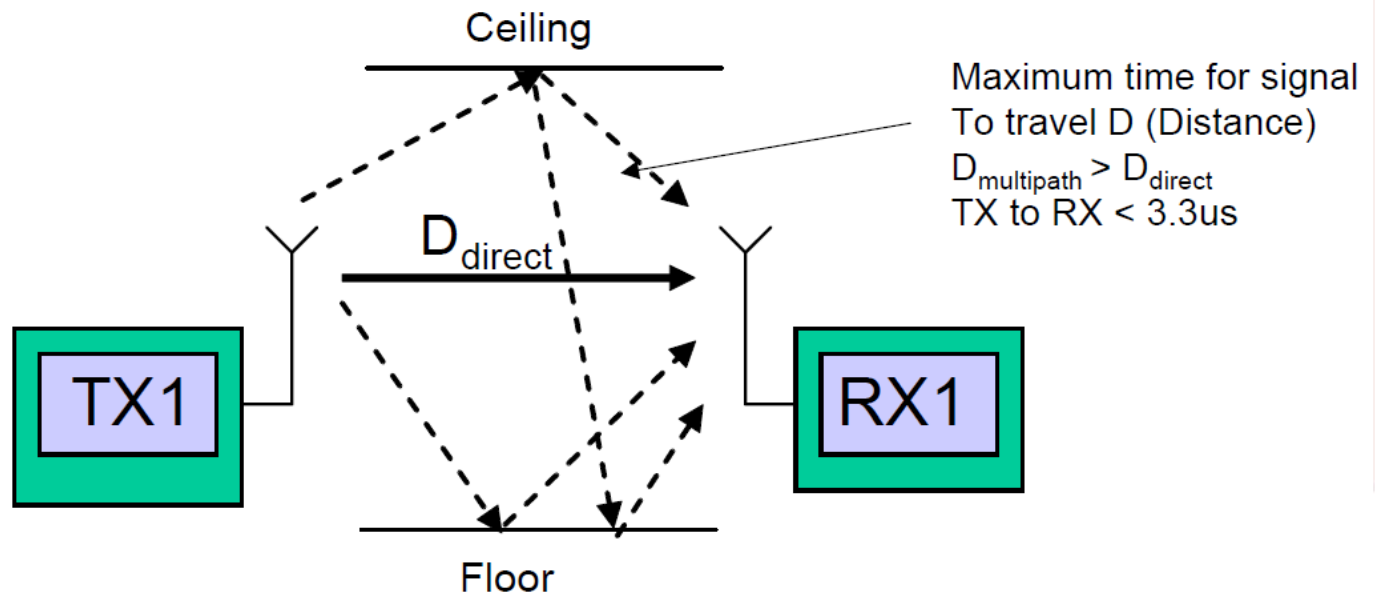
Single Carrier –Single Symbol

- Bluetooth, GSM, CDMA and other communications standards use a single carrier to transmit a single symbol at a time.
- Data throughput is achieved by using a very fast symbol rate.
- W-CDMA -3.14 Msymbols/sec
- Bluetooth —1 Msymbols/sec
- A primary disadvantage is that fast symbol rates are more susceptible to Multi-path distortion.

Slow the symbol rate.....Reduce the previous examples symbol rate by a third

Symbol Rate = 300kSymbols/s
Symbol Duration = $1/300 = 3.3\mu\text{s}$

Maximum Symbol Delay < $3.3\mu\text{s}$



But the data throughput is reduced!

Improve the throughput -use more than one carrier!

Low symbol rate per carrier * multiple carriers

= high data rate

250 kbps symbol rate * 48 sub-carriers * 6 coded bits /sub-carrier * $\frac{3}{4}$ coding rate = 54 Mbps

Introduction to OFDM

Basic idea

- Using a large number of parallel narrow-band sub-carriers instead of a single wide-band carrier to transport information.

Advantages

- Very easy and efficient in dealing with multi-path
- Robust against narrow-band interference

Disadvantages

- Sensitive to frequency offset and phase noise
- Peak-to-average problem reduces the power efficiency of RF amplifier at the transmitter

- **Modulation** - a mapping of the information on changes in the carrier phase, frequency or amplitude or combination.
- **Multiplexing** - method of sharing a bandwidth with other independent data channels.
- OFDM is a combination of modulation and multiplexing.
- Multiplexing generally refers to independent signals, those produced by different sources.
- In OFDM the question of multiplexing is applied to independent signals but these independent signals are a sub-set of the one main signal.
- In OFDM the signal itself is first split into independent channels, modulated by data and then re-multiplexed to create the OFDM carrier.
- OFDM is a special case of Frequency Division Multiplex (FDM).
- As an analogy, a FDM channel is like water flow out of a faucet, in contrast the OFDM signal is like a shower. In a faucet all water comes in one big stream and cannot be sub-divided. OFDM shower is made up of a lot of little streams.

Basic Concept of OFDM

- The basic concept of OFDM was first proposed by R. W. Chang, recognizing that band limited orthogonal signals can be combined with *significant overlap* while avoiding inter channel interference.
- These subcarriers must be orthogonal functions.
- The precise mathematical definition for orthogonality between two functions is that the integral of their product over the designated time interval is zero.

- **OFDM** - Orthogonal Frequency Division Multiplexing is a form of signal waveform or modulation that provides some significant advantages for data links.
- It is used for many of the latest wide bandwidth and high data rate wireless systems including Wi-Fi, cellular telecommunications and many more.
- The fact that OFDM uses a large number of carriers, each carrying low bit rate data, means that it is very resilient to selective fading, interference, and multipath effects, as well providing a high degree of spectral efficiency.

Where OFDM is Used

Wireless Standards

- **Cellular (Mobile) telecommunications** : LTE and LTE-A
- **Wi-Fi standards** : 802.11a, 802.11g, 802.11n, 802.11ac and HIPERLAN/2
- **Mobile broadband wireless access (MBWA) standard:** IEEE 802.20.
- **Broadcast standards:** DAB Digital Radio, DVB and Digital Radio Mondiale.
- **Mobile TV standard** as DVB-H, T-DMB, ISDB-T and MediaFLO forward link.

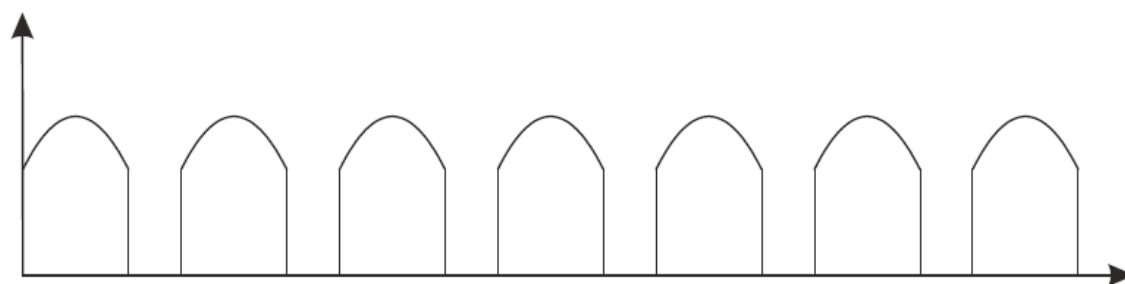
Cable Standards

- ADSL and VDSL broadband utilising copper wiring.
- Cable TV, with DVB-C2.
- Power Line Communications
- Data over cable service interface specification (DOCSIS)

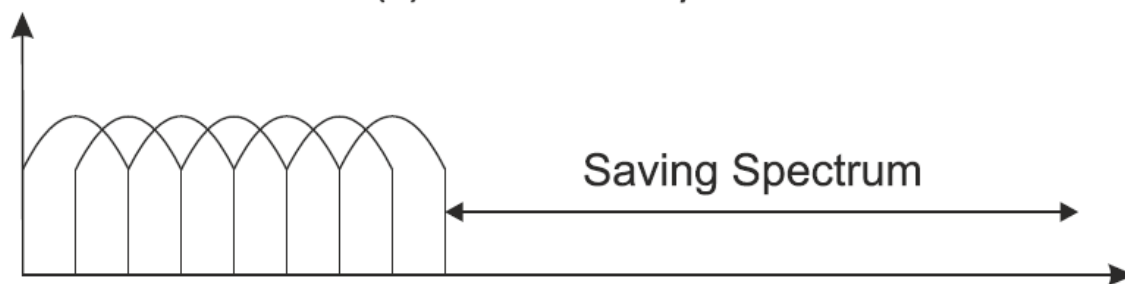
Overview of OFDM

Introduction

- Orthogonal Frequency Division Multiplex (OFDM).
- Special case of multicarrier transmission → **overlapping** and **orthogonal** narrowband subchannels.

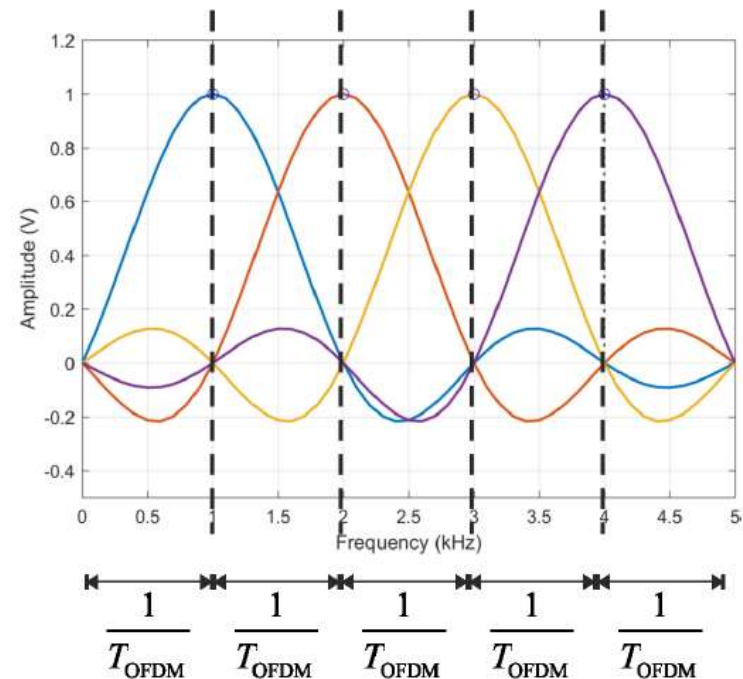
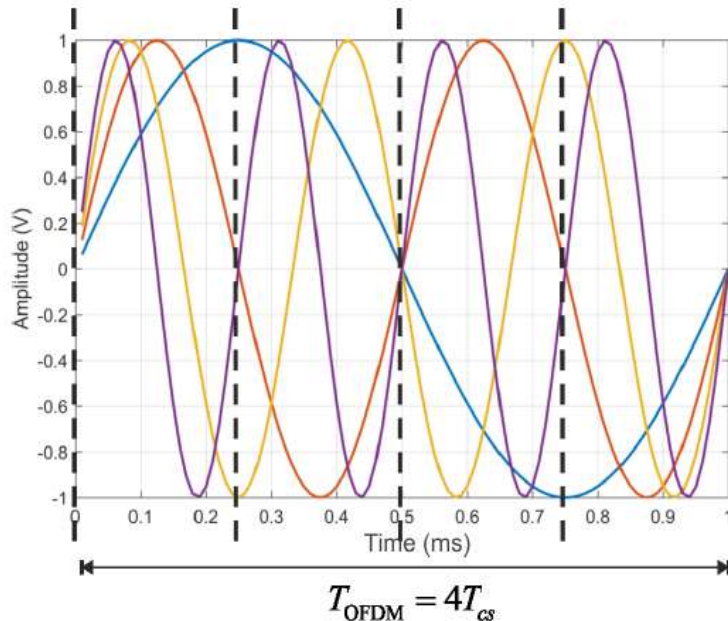


(a) Multicarrier system



(b) OFDM system

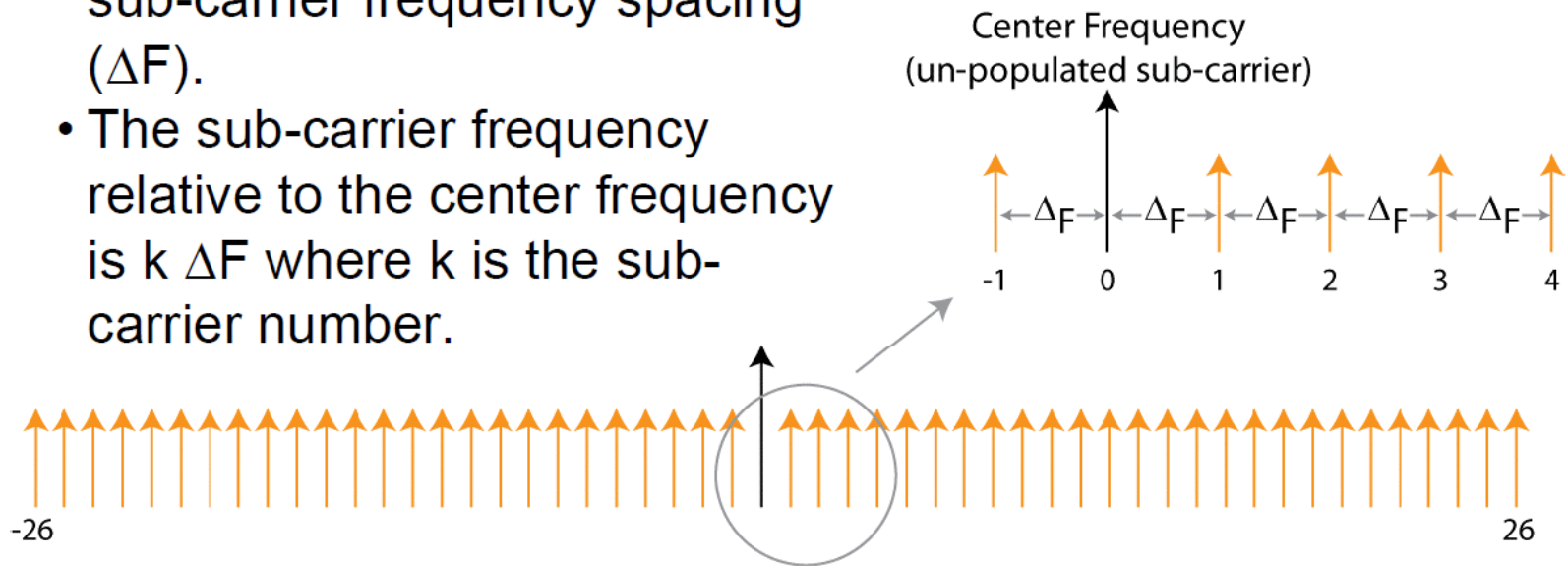
OFDM Subcarriers



- The chosen subcarriers are harmonics and hence orthogonal to each other.
- No guard band between subcarriers are needed.
- The peak of a subcarrier is precisely at the zero-crossings of its adjacent subcarriers.

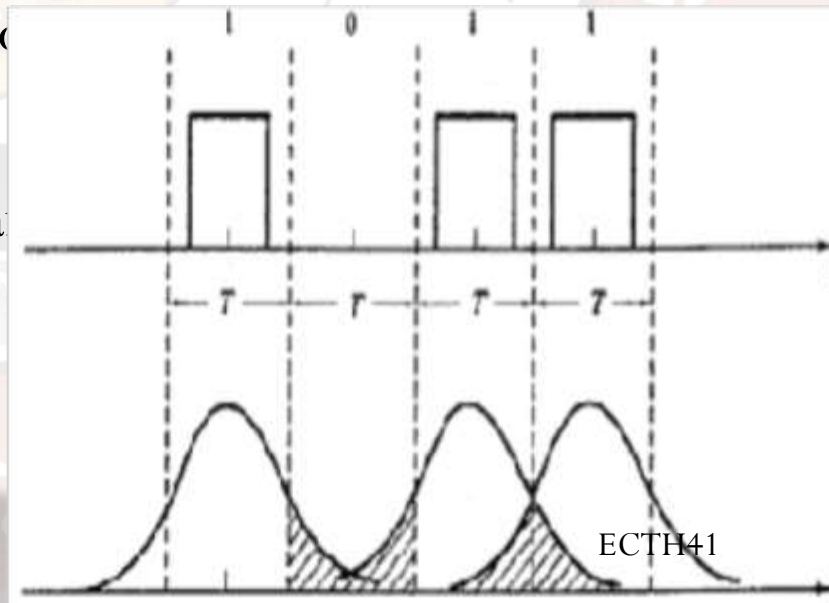
Sub Carrier Spacing

- The sub-carriers are spaced at regular intervals called the sub-carrier frequency spacing (ΔF).
- The sub-carrier frequency relative to the center frequency is $k \Delta F$ where k is the sub-carrier number.



ISI - Inter Symbol Interference

- ISI caused by a delayed version of one symbol or more from an indirect path that interferes with a different symbol.
- High data rate requires smaller symbol period.
- If symbol period is less than delay spread then we have ISI.
- Possible solutions to resolve ISI in single carrier-Expensive channel equalizers or low data rates.
- Solution is to increase the symbol period while using multiple carriers.
- The period of a symbol on each sub-carrier
- is more than the delay spread.

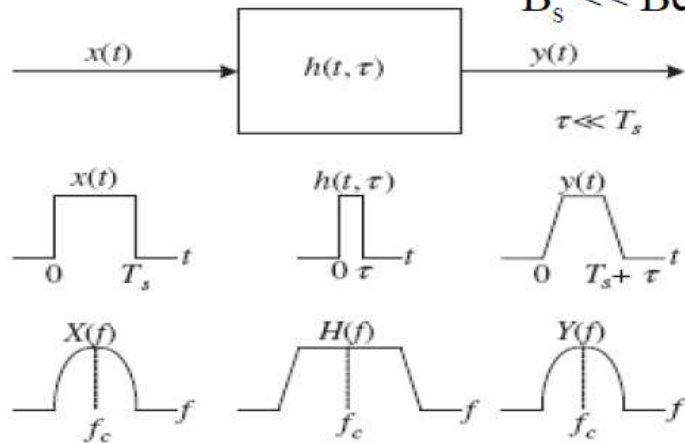


Need of OFDM [7]

Flat Fading

$$T_s \gg \sigma\tau$$

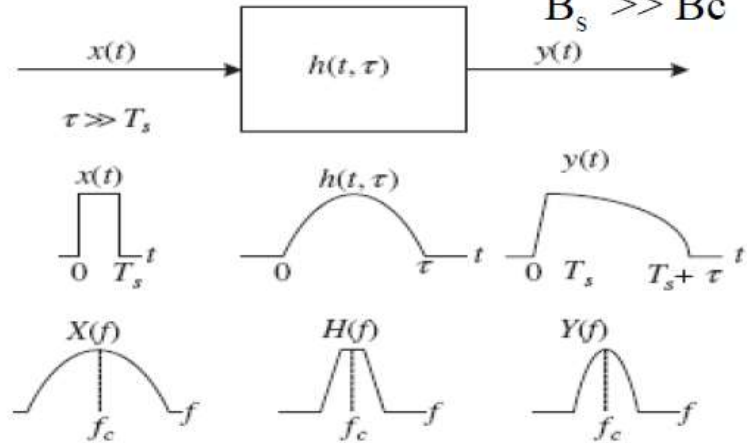
$$B_s \ll B_c$$



Freq. Selective Fading

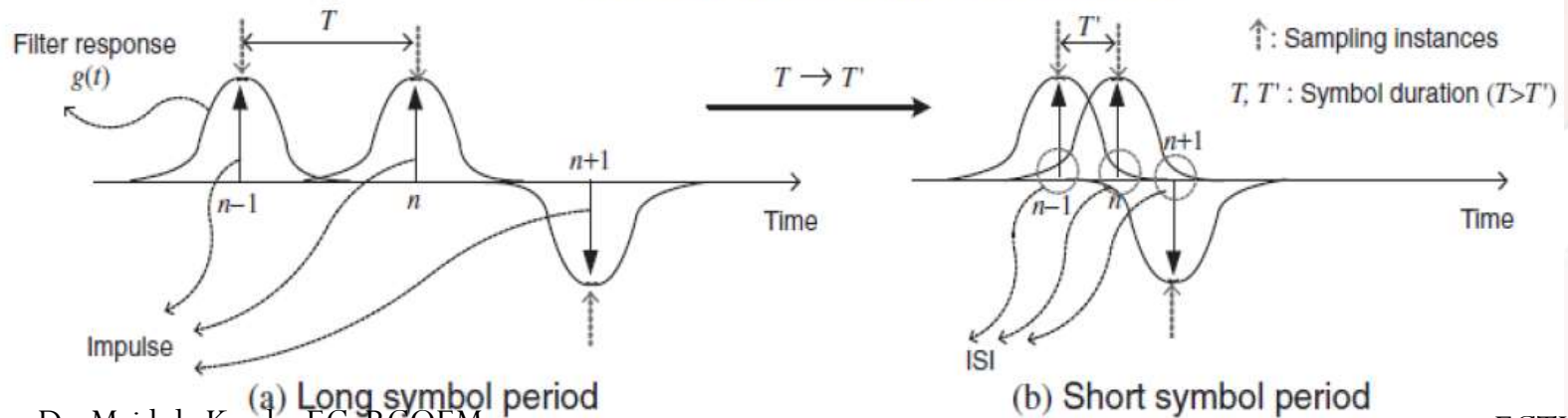
$$T_s \ll \sigma\tau$$

$$B_s \gg B_c$$



Single Carrier Transmission : Problem of Frequency Selective Fading

Problem of ISI as data rate increases



Principle of Orthogonal Multiplexing

- A high-rate data stream typically faces a problem in having a symbol period T_s much smaller than the channel delay spread T_d if it is transmitted serially.
- This generates Intersymbol Interference (ISI) which can only be undone by means of a complex equalization procedure.
- In OFDM, the high-rate stream of data symbols is first serial-to-parallel converted for modulation onto M parallel subcarriers as shown in Figure 5.2.
- This increases the symbol duration on each subcarrier by a factor of approximately M , such that it becomes significantly longer than the channel delay spread.

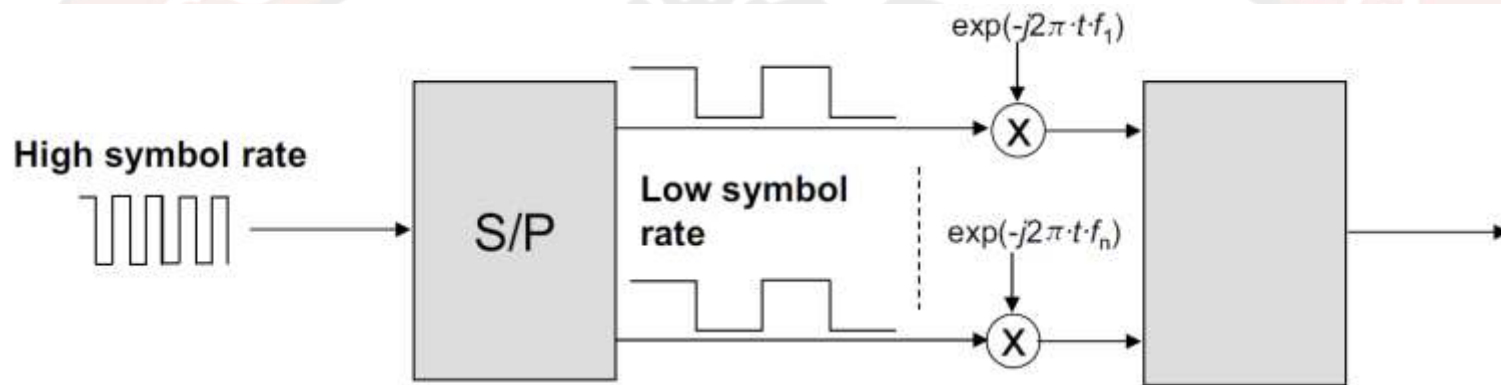


Figure 5.2 Serial-to-parallel conversion operation for OFDM.

Figure 5.3 shows how the resulting long symbol duration is virtually unaffected by ISI compared to the short symbol duration, which is highly corrupted.

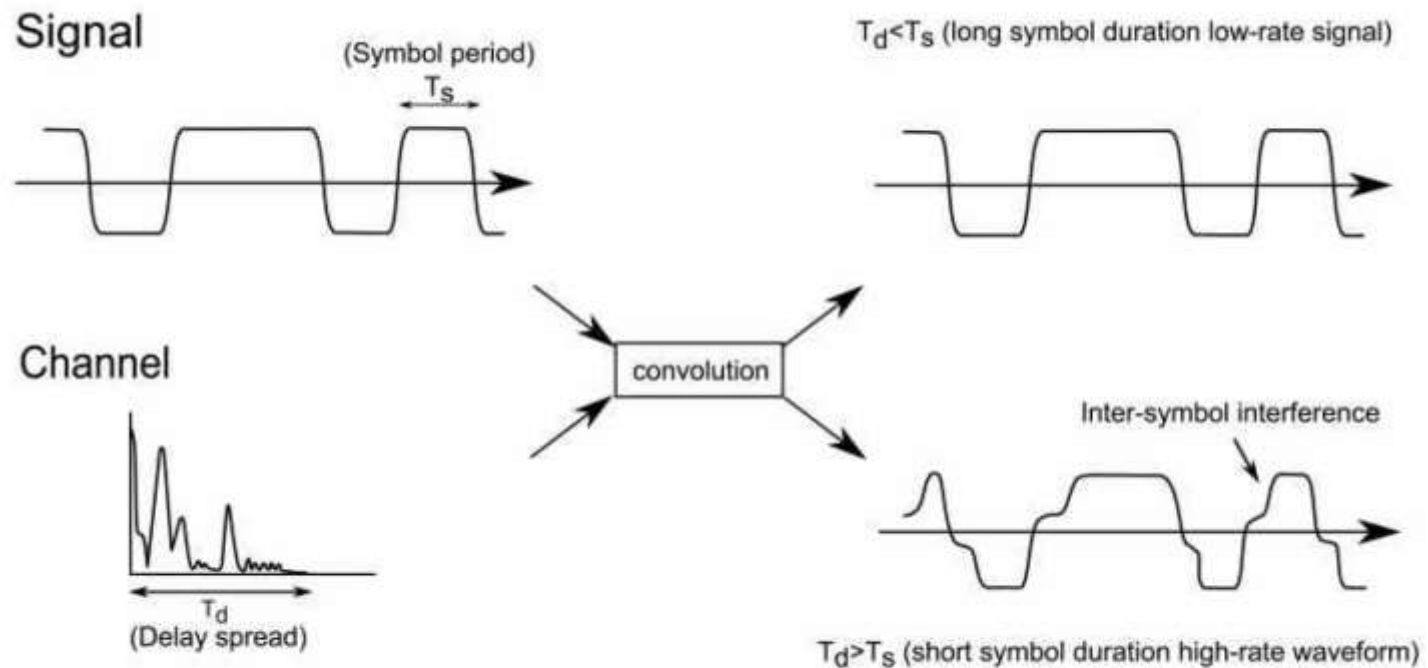
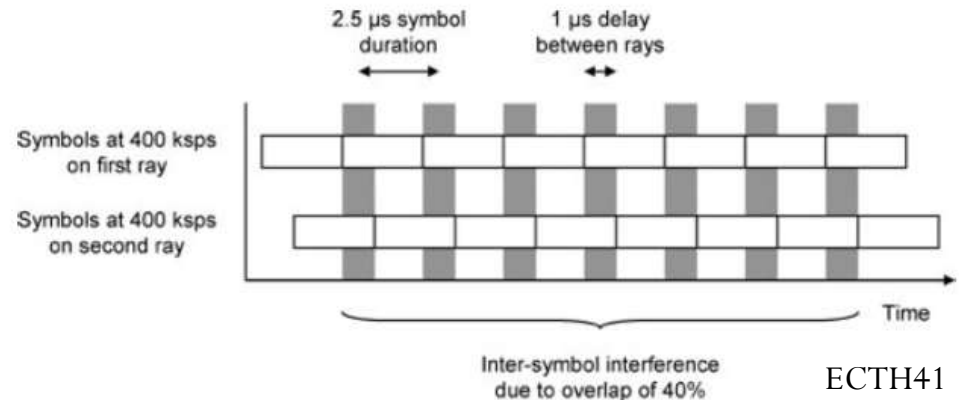


Figure 5.3 Effect of channel on signals with short and long symbol duration.

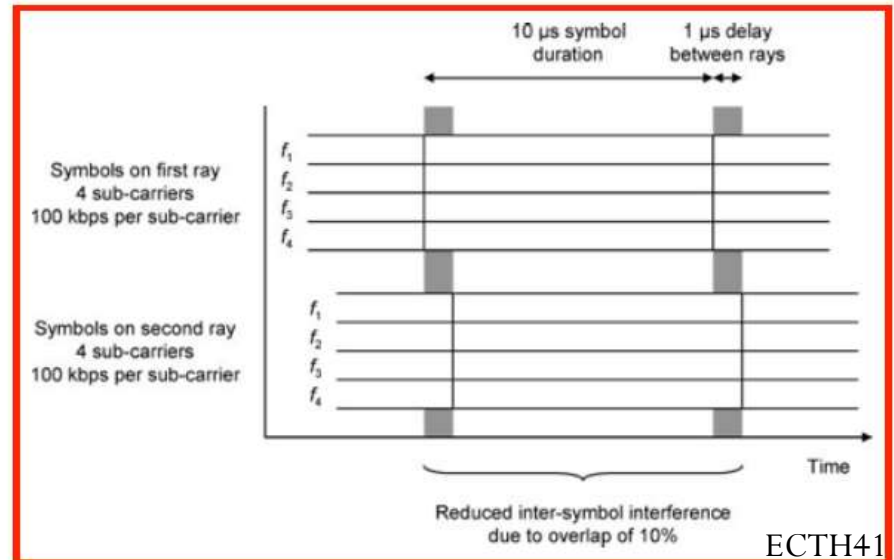
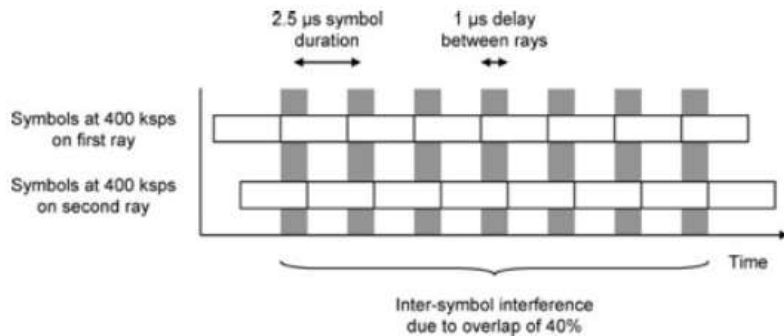
1.1 Reduction of Inter-Symbol Interference using OFDM

- High data rate transmission in a multipath environment leads to Inter Symbol Interference (ISI)
 - ✓ Example, the delay spread was $1\mu s$ and the data rate was 400 kbps, so the symbols overlapped at the receiver by 40%
 - ✓ This led to interference and bit errors at the receiver
 - ✓ OFDM is a powerful way to solve the problem



- OFDM transmitter

- ✓ Divides the information into several parallel sub-streams
- ✓ Sends each sub-stream on a different frequency (sub-carrier)
 - ▶ If the total data rate stays the same, then
 - The data rate on each sub-carrier is less than before, so the symbol duration is longer
 - ▶ This reduces the amount of ISI, and reduces the error rate



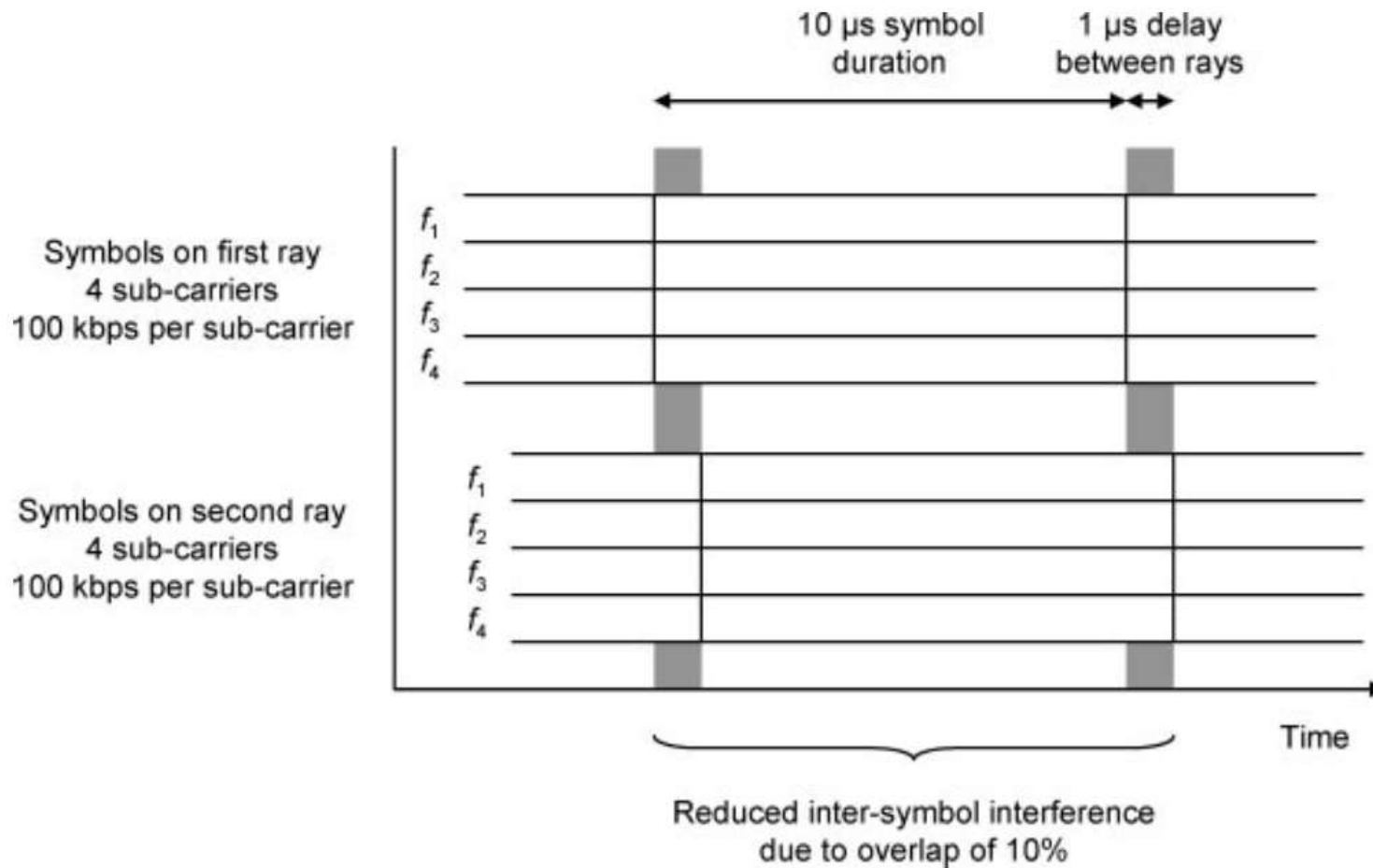
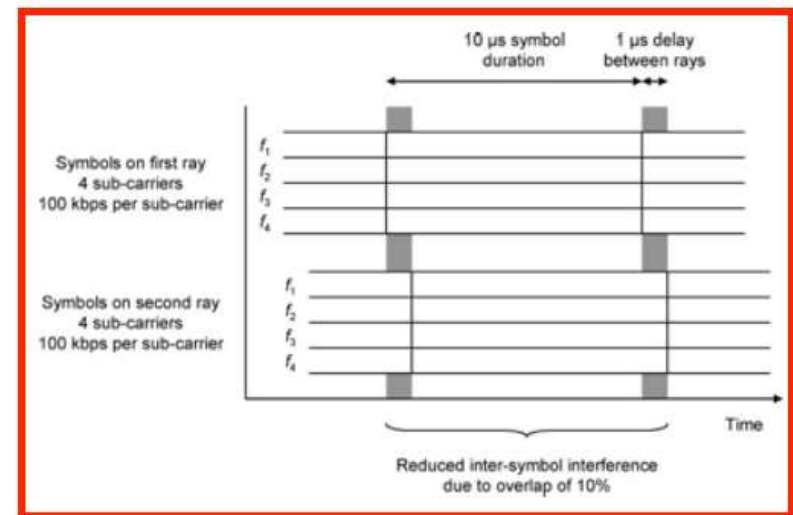
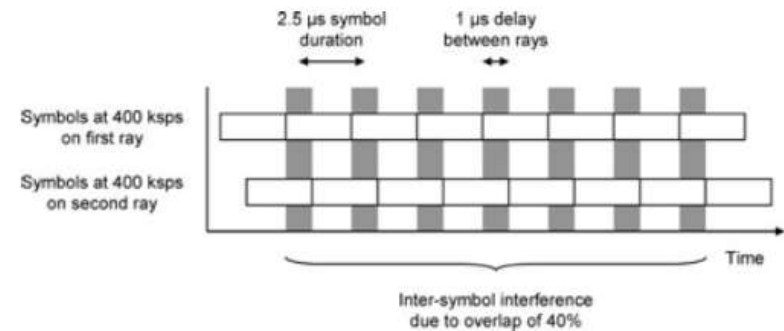
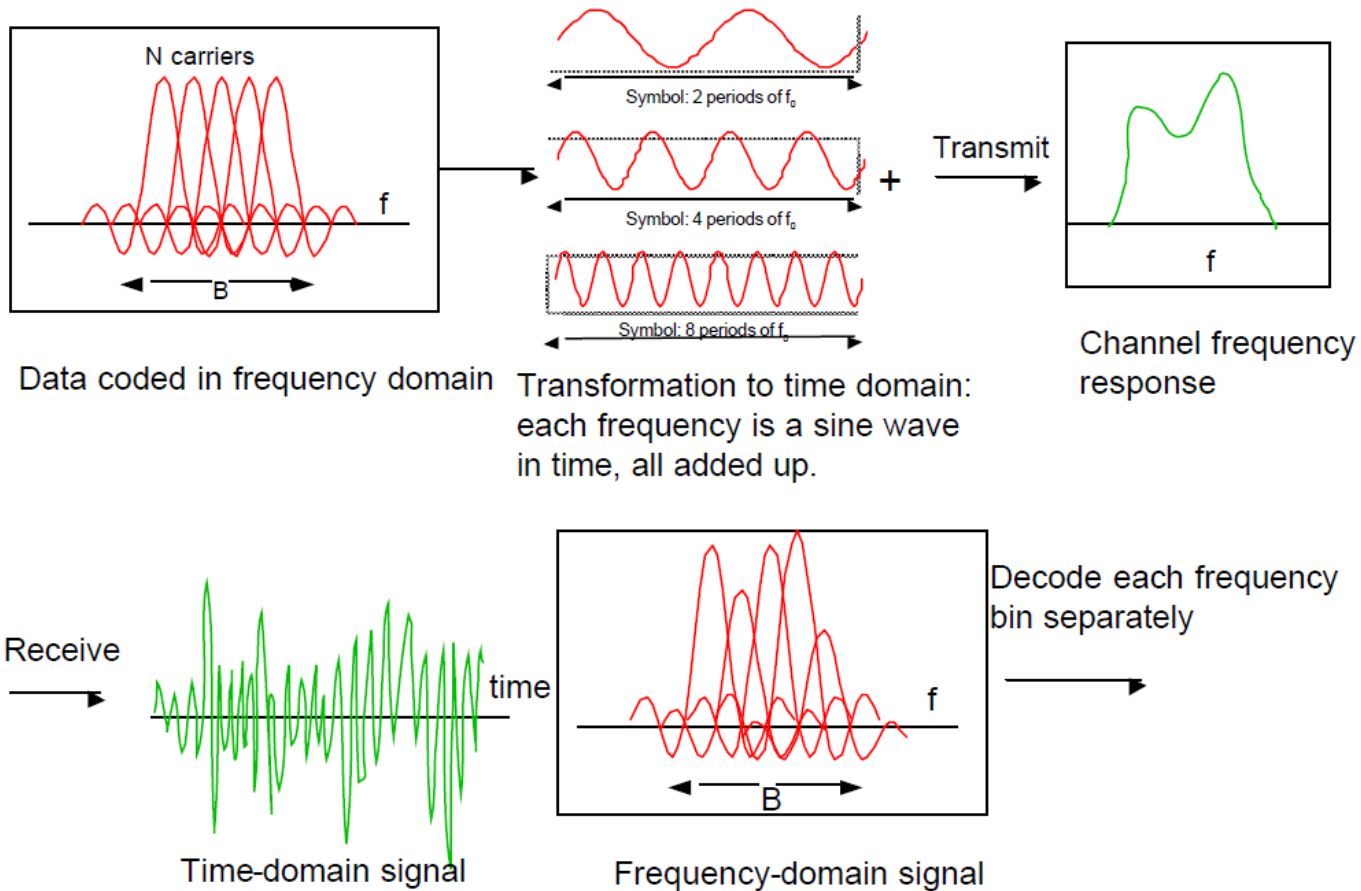


Figure 4.1 Reduction of inter-symbol interference by transmission on multiple sub-carriers.

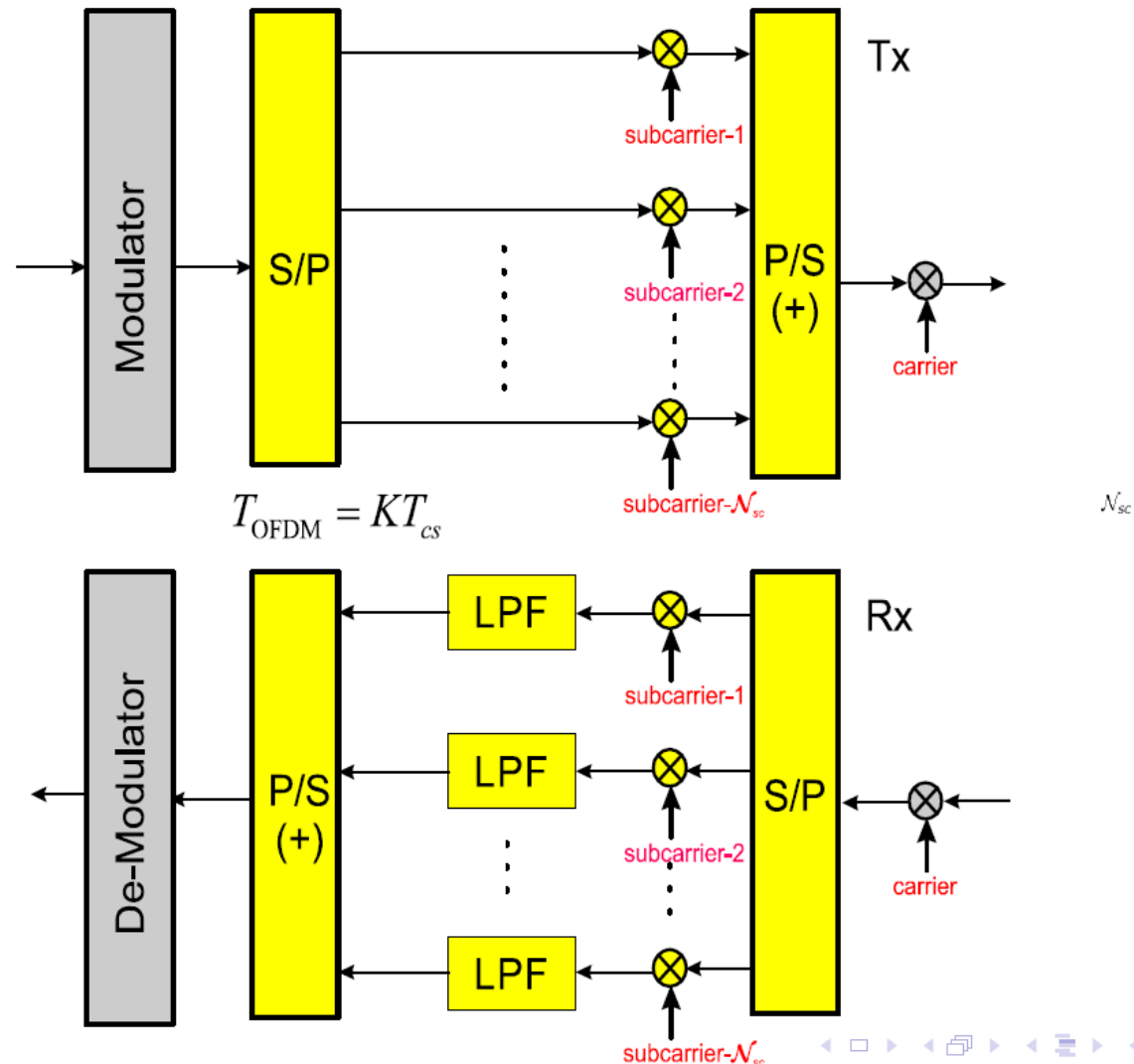
- Figure 4.1
 - ✓ Divid the original data stream amongst four sub-carriers with frequencies f_1 to f_4
 - ✓ The data rate on each sub-carrier is now 100 kbps, so the symbol duration has increased to $10\mu s$
 - ✓ If the delay spread remains at $1\mu s$, then the symbols only overlap by 10%
 - ✓ This reduces the amount of ISI to one quarter of what it was before and reduces the number of errors in the receiver
- LTE can use a very large number of sub-carriers, up to a maximum of 1200 in Release 8, which reduces the amount of ISI to negligible levels



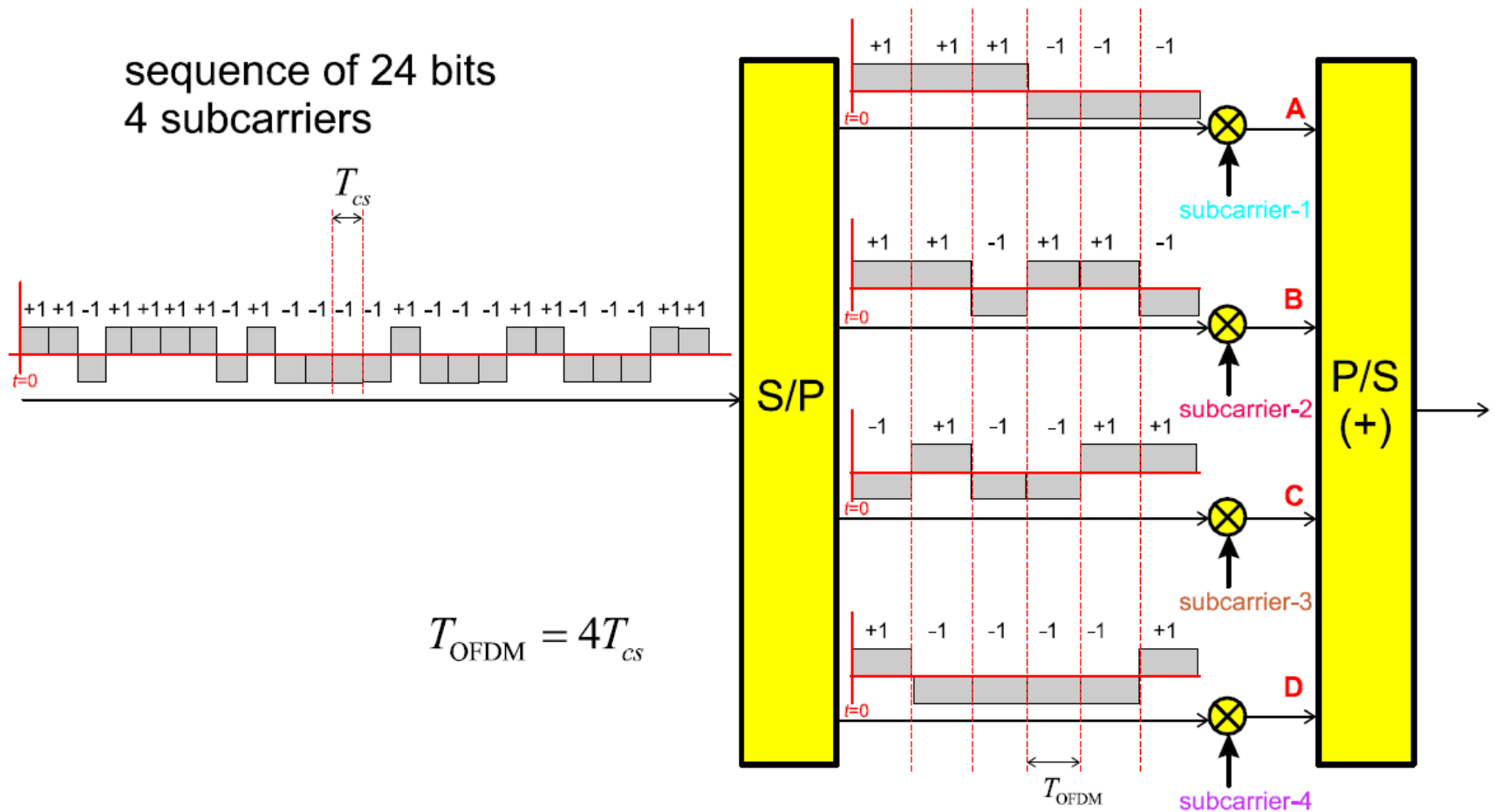
Orthogonal Frequency Division Modulation



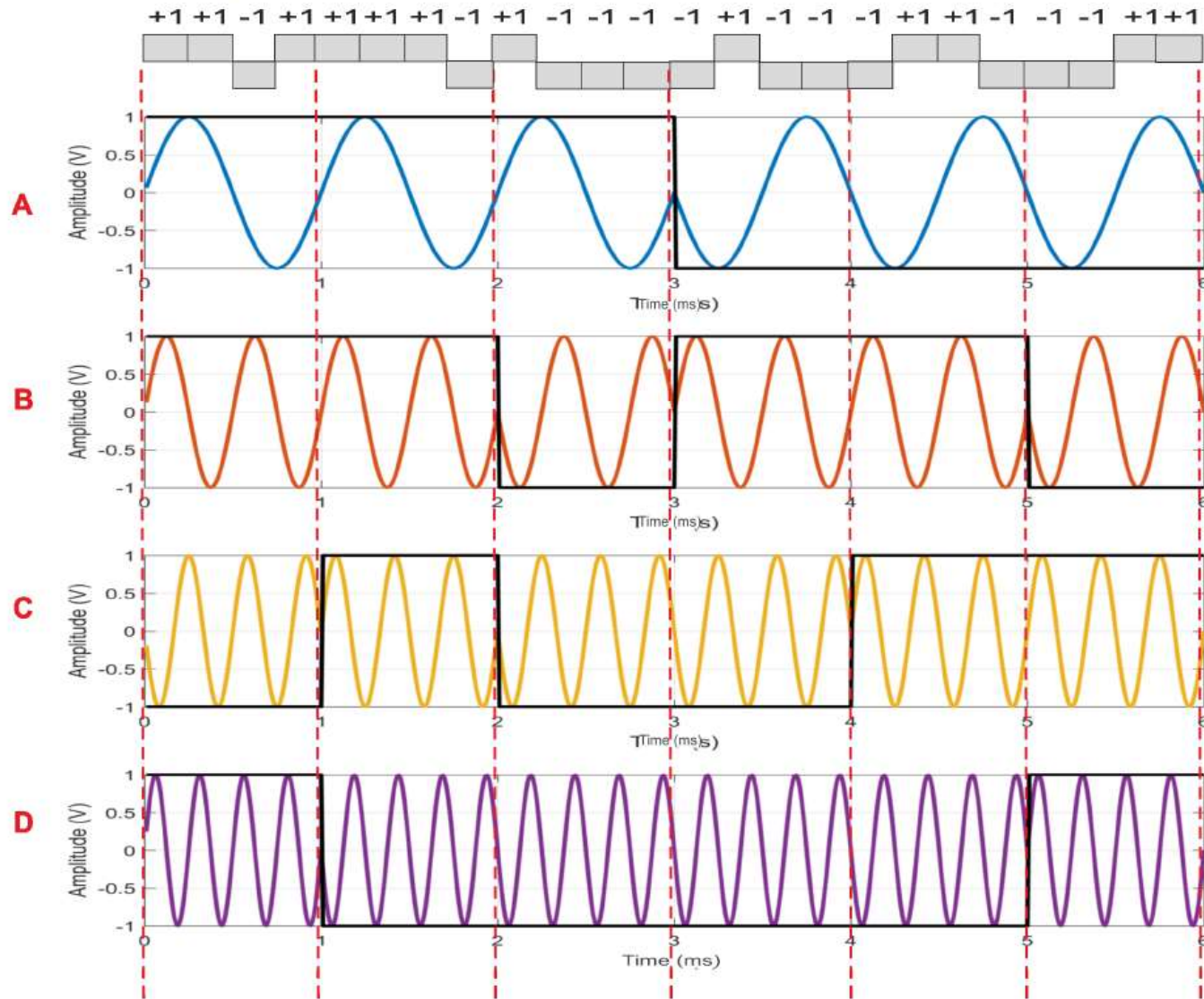
OFDM - Analogue Block Diagram

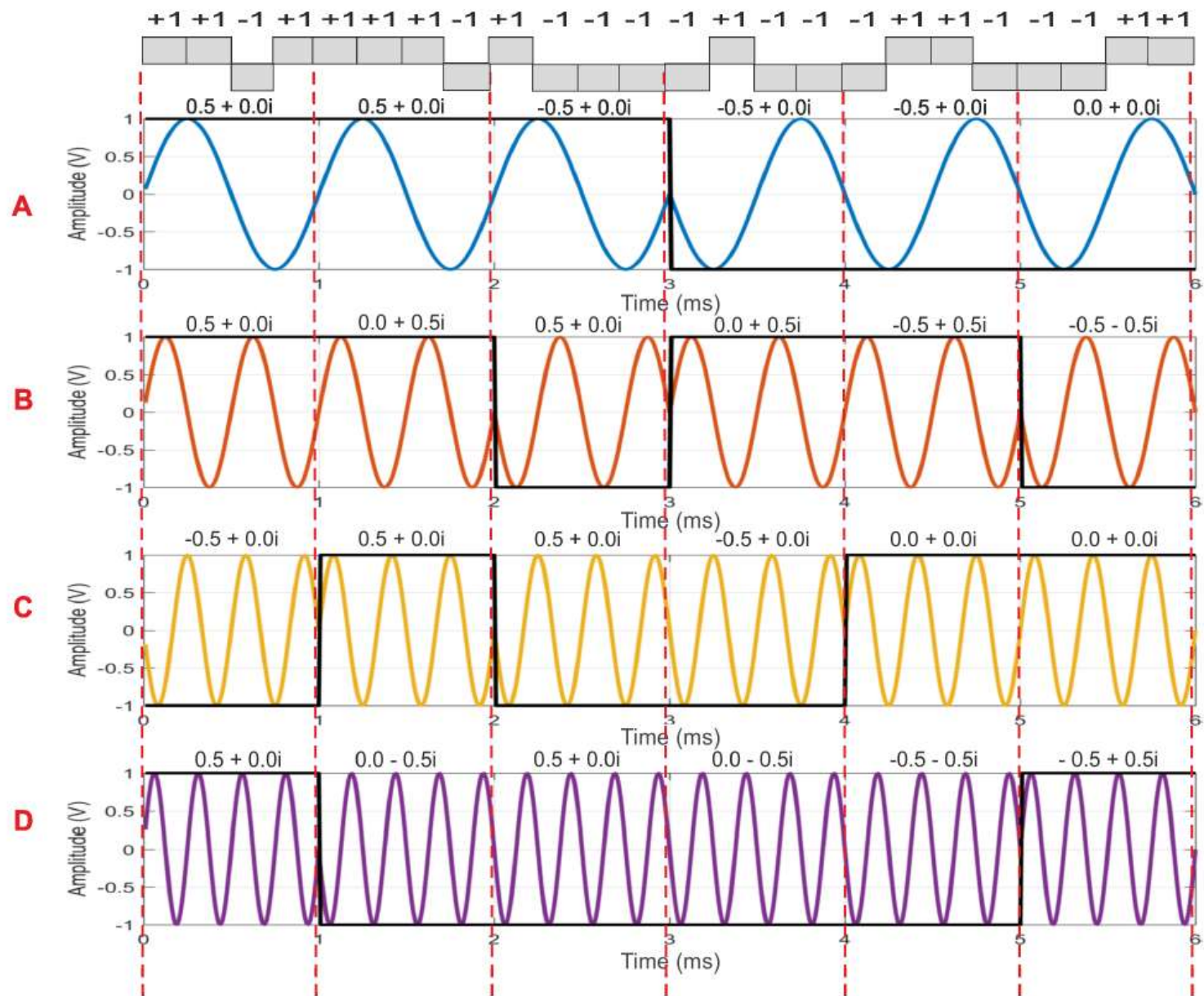


OFDM TX: Example of 4 Subcarriers



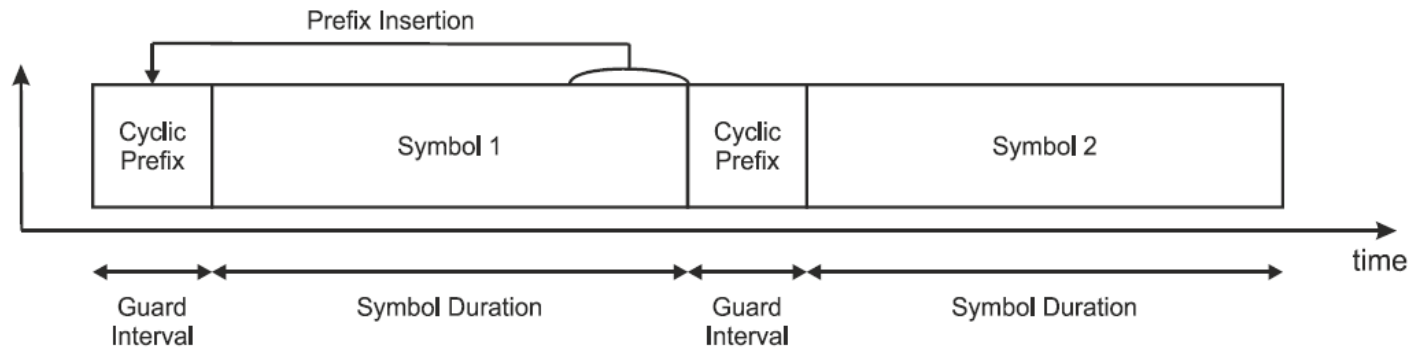
OFDM Tx Output





Cyclic Prefix, CP

- The **cyclic prefix** is a **guard interval (GI)** placed to protect OFDM signals from intersymbol interference in the presence of multipath.



- The last section of a symbol is used as a prefix in the front of the symbol.
- The duration of the guard interval T_{GI} should be

$$T_{GI} > T_{\text{spread}}, \quad (1)$$

to minimise **inter-symbol interference (ISI)**.

- This avoids the need to separate the carriers by means of guard-bands, and therefore makes OFDM highly spectrally efficient.
- The spacing between the subchannels in OFDM is such they can be perfectly separated at the receiver.
- This allows for a low-complexity receiver implementation, which makes OFDM attractive for high-rate mobile data transmission such as the LTE downlink.

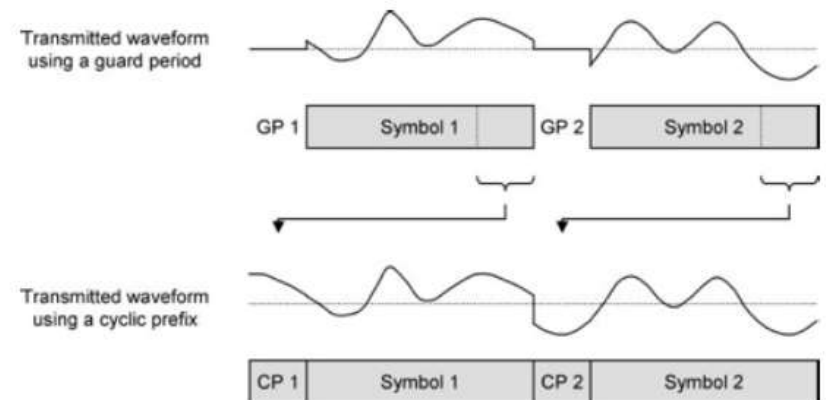
2.4 Cyclic Prefix Insertion

- A technique used to get rid of inter symbol interference (ISI)
- The basic idea is to insert a guard period (GP) before each symbol, in which nothing is transmitted
- If the guard period is longer than the delay spread, then the receiver can be confident of reading information from just one symbol at a time, without any overlap with the symbols that precede or follow
- The symbol reaches the receiver at different times on different rays and some extra processing is required to tidy up the confusion

- LTE uses a cyclic prefix (CP) insertion

✓ The transmitter

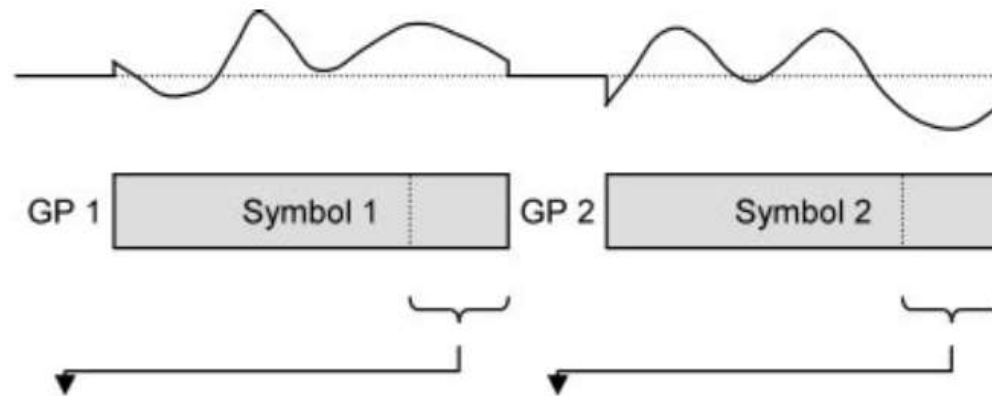
- ▶ Starts by inserting a guard period before each symbol, as before
- ▶ It then copies data from the end of the symbol following, so as to fill up the guard period



✓ The receiver

- ▶ If the cyclic prefix is longer than the delay spread, then the receiver can still be confident of reading information from just one symbol at a time

Transmitted waveform
using a guard period



Transmitted waveform
using a cyclic prefix

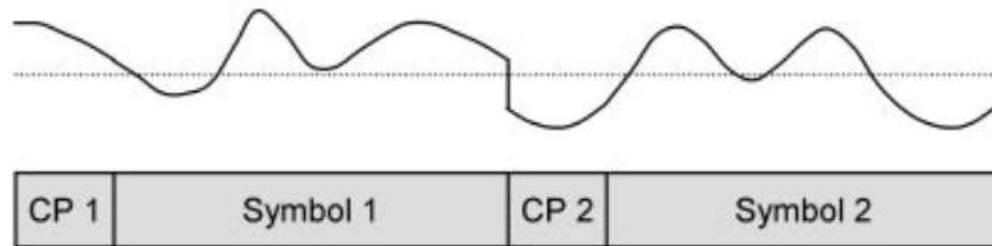
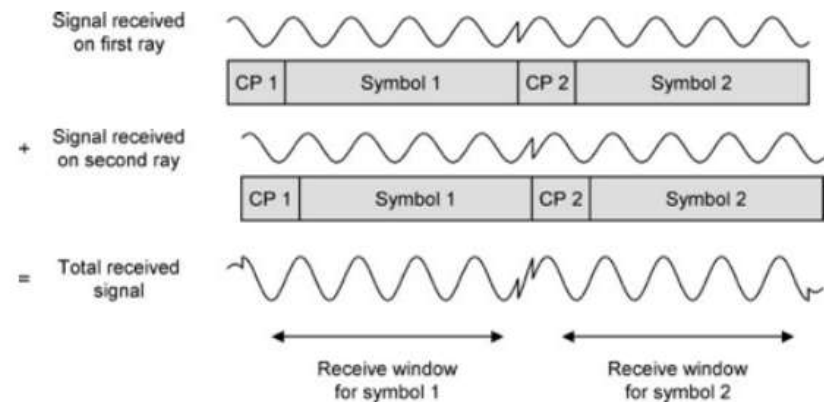


Figure 4.8 Operation of cyclic prefix insertion.

- Figure 4.9 shows how cyclic prefix insertion works

✓ The transmitter

- ▶ The transmitted signal is a sine wave, whose amplitude and phase change from one symbol to the next
- ▶ As noted earlier, each symbol contains an exact number of cycles of the sine wave, so the amplitude and phase at the start of each symbol equal the amplitude and phase at the end
- ▶ Because of this, the transmitted signal changes smoothly as we move from each cyclic prefix to the symbol following



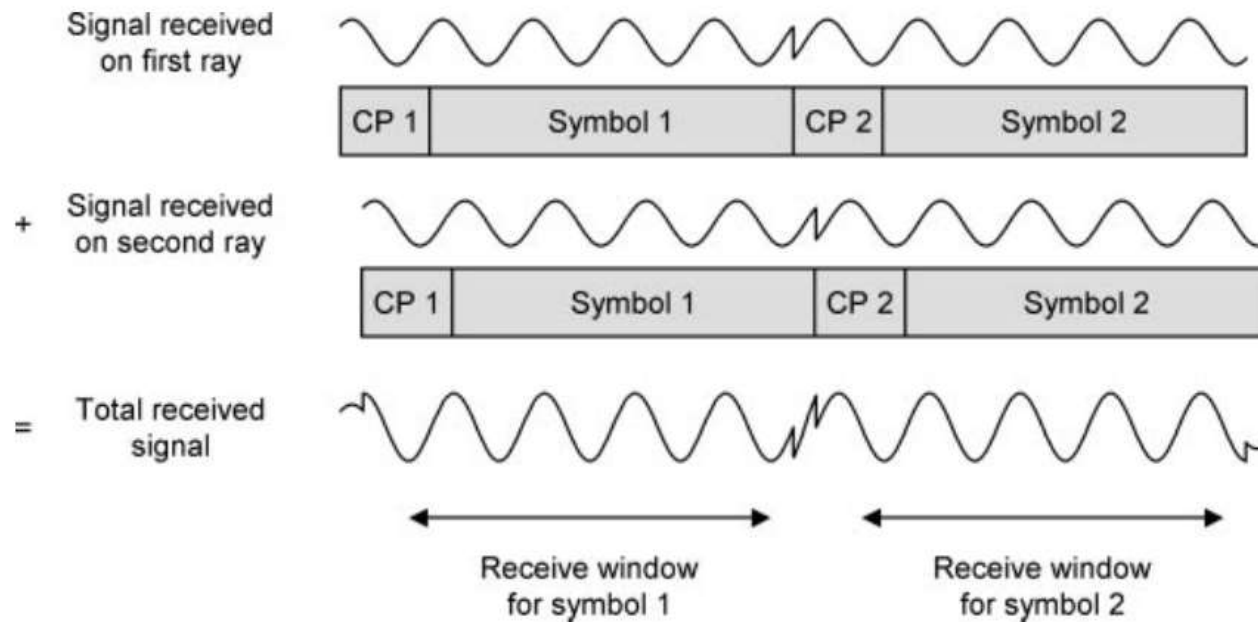


Figure 4.9 Operation of the cyclic prefix on a single sub-carrier.

✓ The receiver

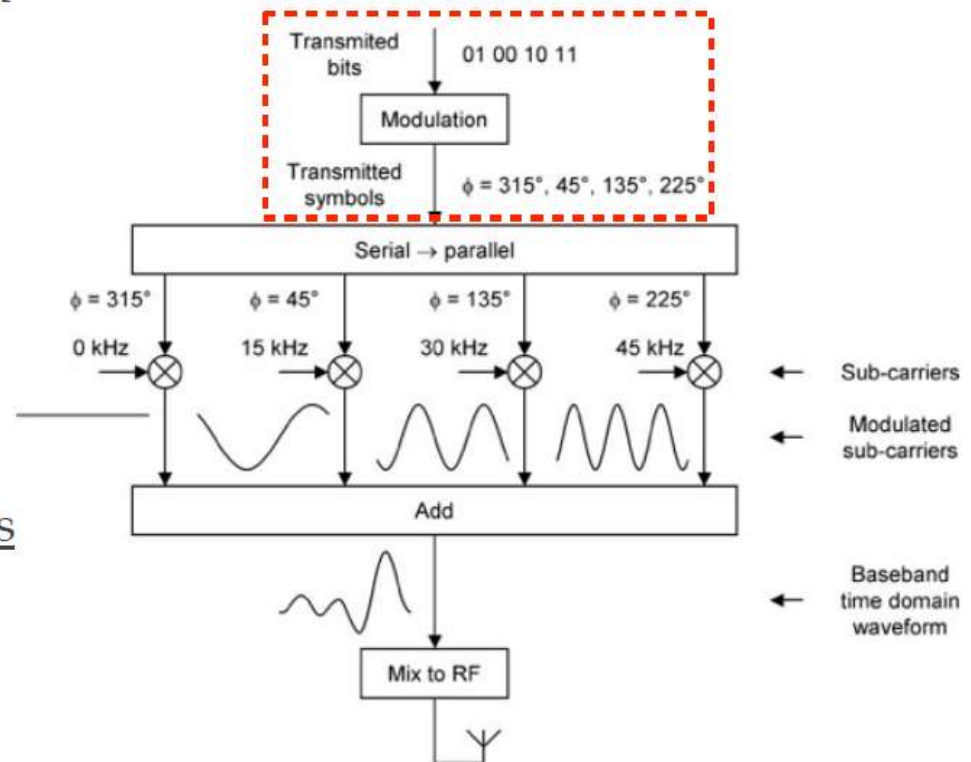
- ▶ In a multipath environment, the receiver picks up multiple copies of the transmitted signal with multiple arrival times
- ▶ These add together at the receive antenna, giving a sine wave with the same frequency but a different amplitude and phase
- ▶ The received signal still changes smoothly at the transition from a cyclic prefix to the symbol that follows
- ▶ There are a few glitches, but these are only at the start of the cyclic prefix and the end of the symbol, where the preceding and following symbols start to interfere

- ▶ The receiver processes the received signal within a window whose length equals the symbol duration, and discards the remainder
- ▶ If the window is correctly placed, then the received signal is exactly what was transmitted, without any glitches, and subject only to an amplitude change and a phase shift
- ▶ But the receiver can compensate for these using the channel estimation and equalization techniques described above
- ▶ Admittedly the system uses multiple sub-carriers, not just one
- ▶ The sub-carriers do not interfere with each other and can be treated independently, so the existence of multiple sub-carriers does not affect this argument at all

- LTE uses a cyclic prefix of about $4.7 \mu\text{s}$
 - ✓ This corresponds to a maximum path difference of about 1.4 km between the lengths of the longest and shortest rays, which is enough for all but the very largest and most cluttered cells
 - ✓ The cyclic prefix reduces the data rate by about 7%, but this is a small price to pay for the removal of ISI

1.2 The OFDM Transmitter

- Figure 4.2 is a block diagram of an analogue OFDM transmitter
- Transmitter
 - ✓ Accepts a stream of bits from higher layer protocols, and
 - ✓ Converts them to symbols using the chosen modulation scheme, for example quadrature phase shift keying (QPSK)



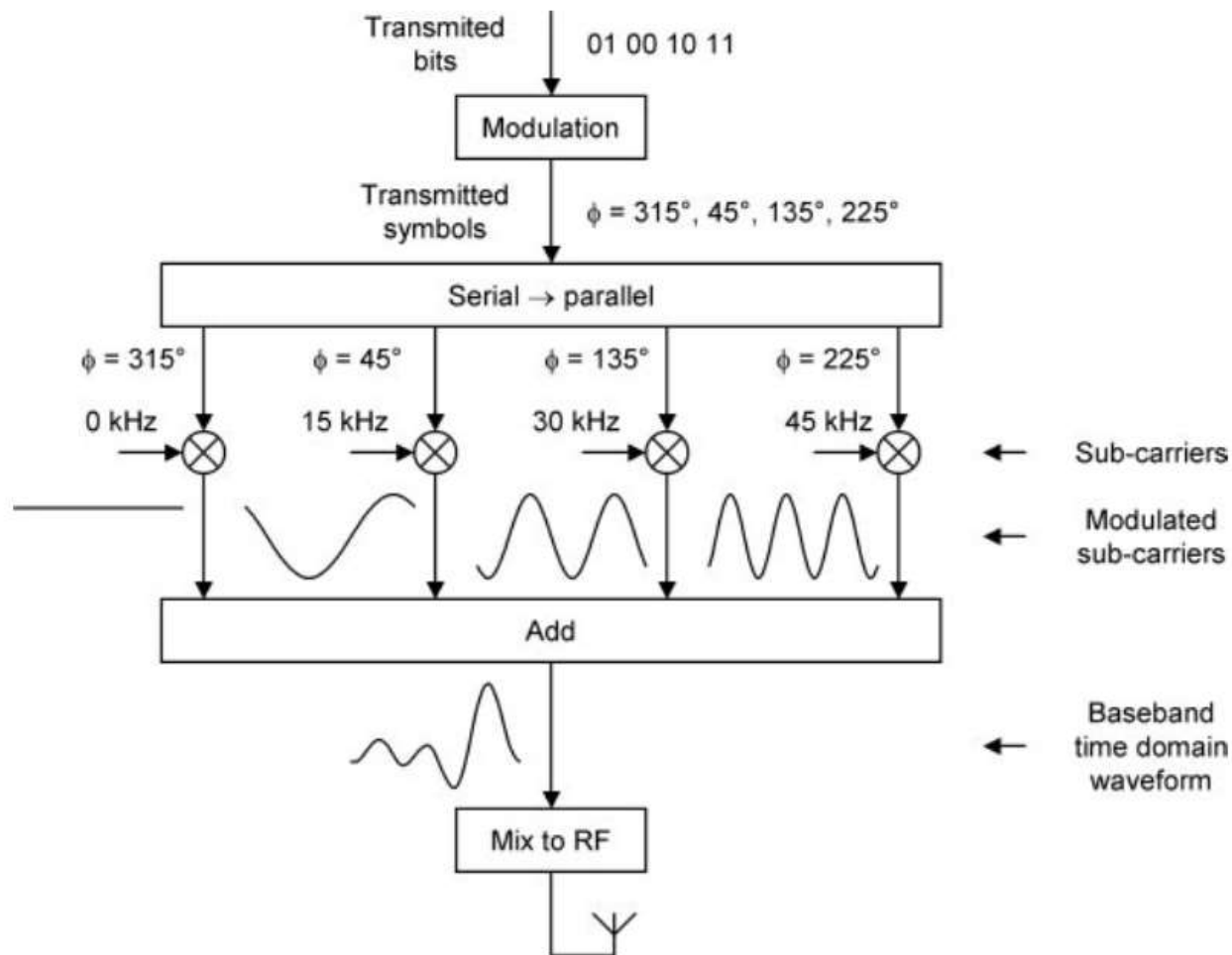
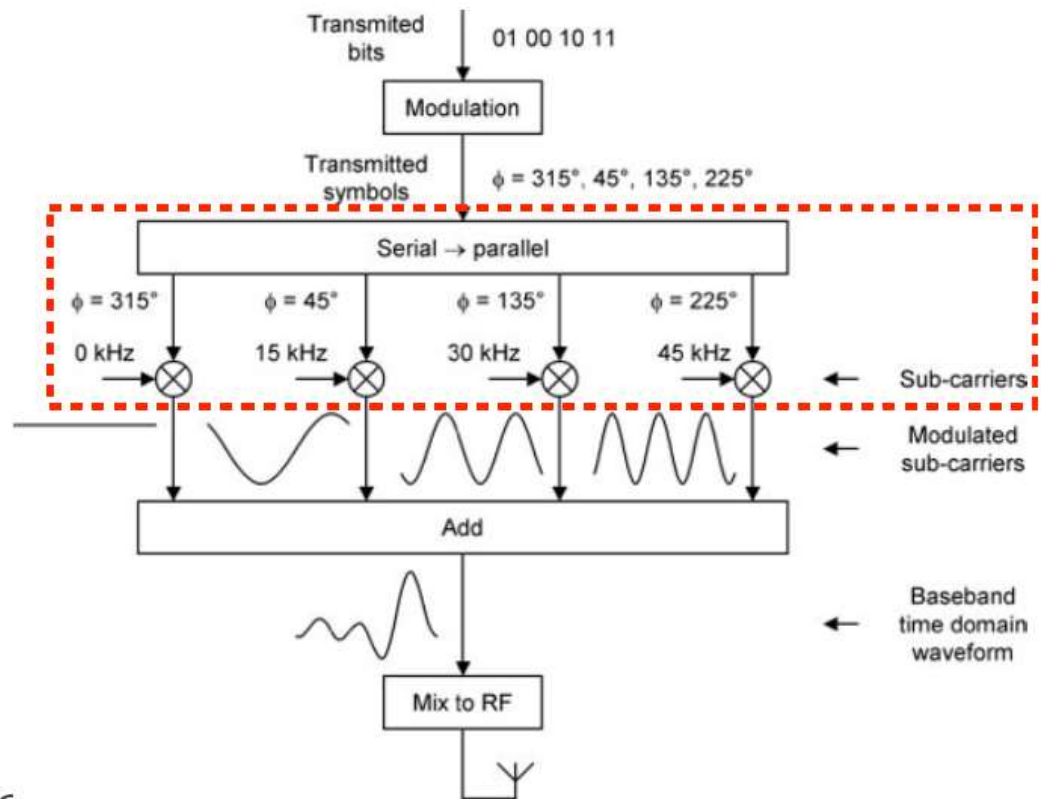
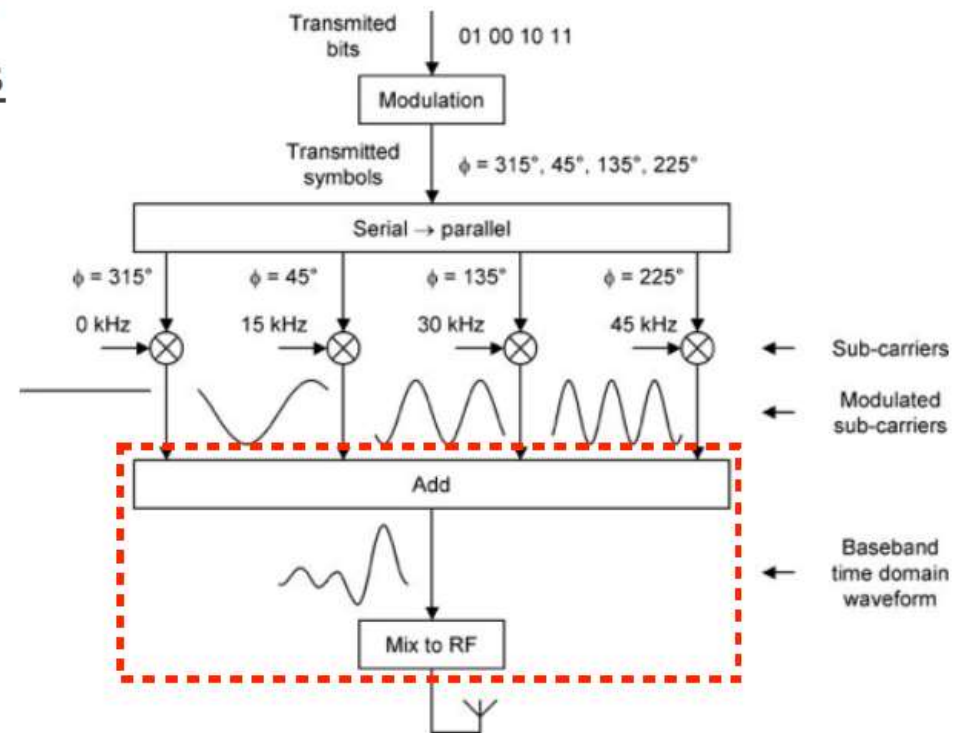


Figure 4.2 Processing steps in a simplified analogue OFDM transmitter.

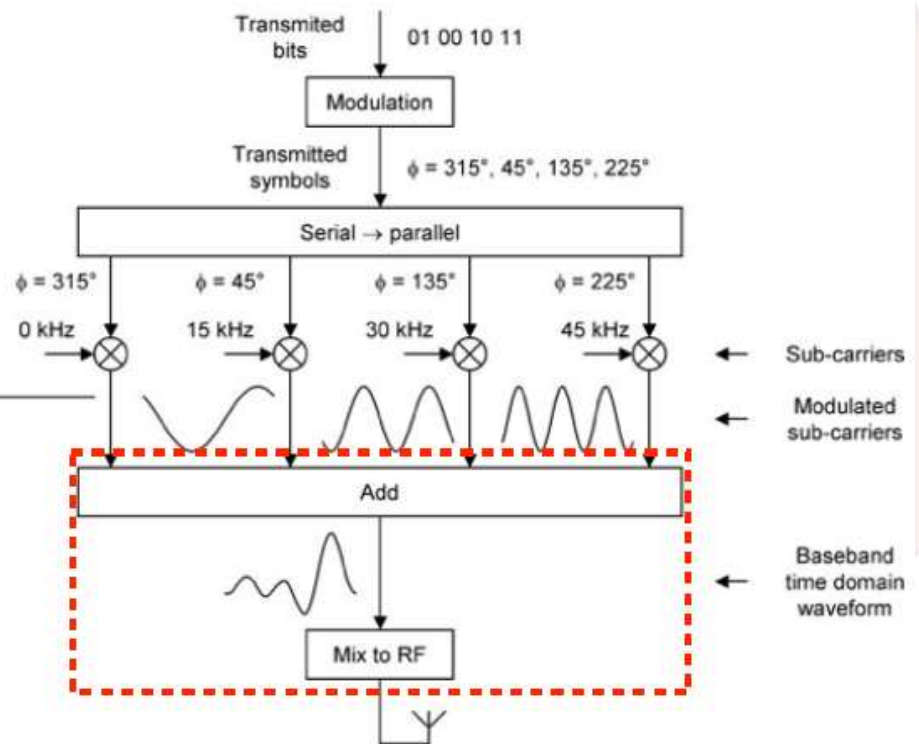
- Serial-to-parallel converter
- ✓ Takes a block of symbols, four in this example
- ✓ Mixes each symbol with one of the sub-carriers by adjusting its amplitude and phase



- LTE uses a fixed sub-carrier spacing of 15 kHz, so the sub-carriers in Figure 4.2 have frequencies of 0, 15, 30 and 45 kHz. (Mix the signals up to radio frequency (RF) at the end)
- The symbol duration is the reciprocal of the sub-carrier spacing, so is about $66.7\mu s$
- ✓ 15 kHz, 30 kHz, and 45 kHz sub-carriers goes through one, two and three cycles during the $66.7\mu s$ symbol duration, respectively



- We now have four sine waves, at frequencies of 0, 15, 30 and 45 kHz, whose amplitudes and phases represent the eight transmitted bits (01 00 10 11)
- By adding these sine waves together, we can generate a single time-domain waveform, which is a low frequency representation of the signal that we need to send
- Finally, mix the waveform up to radio frequency (RF) for transmission

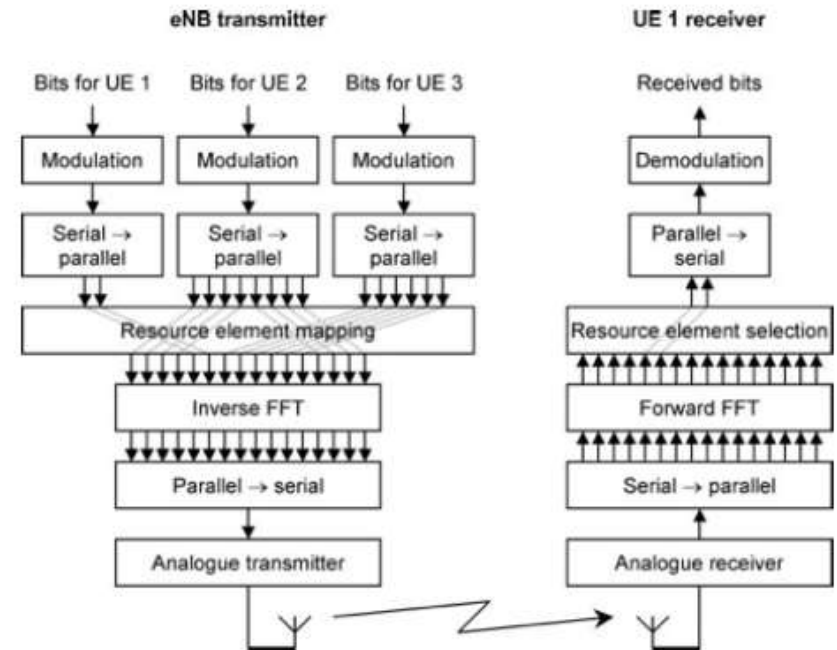


- Two important points in the processing chain
 - ✓ Serial-to-parallel conversion stage
 - ▶ The data represent the amplitude and phase of each sub-carrier, as a function of frequency
 - ✓ After the addition stage towards the end
 - ▶ The data represent the in-phase and quadrature components of the transmitted signal, as a function of time
 - ✓ The mixing and addition steps have simply converted the data from a function of frequency to a function of time

- This conversion is called the inverse discrete Fourier transform (DFT)
 - ✓ The Fourier transform converts data from the time domain to the frequency domain, so the transmitter requires an inverse transform, which carries out the reverse process
- In turn, the discrete Fourier transform can be implemented extremely quickly using fast Fourier transform (FFT)
 - ✓ This limits the computational load on the transmitter and receiver, and allows the two devices to be implemented in a computationally efficient way
 - ✓ However, there is one important restriction
 - ▶ For the FFT to work efficiently, the number of data points should be either an exact power of two or a product of small prime numbers alone

1.3 Initial Block Diagram

- Figure 4.4 is a block diagram of an OFDM transmitter and receiver
- Assume that the system is operating on the downlink, so that the transmitter is in the BS and the receiver is in the mobile



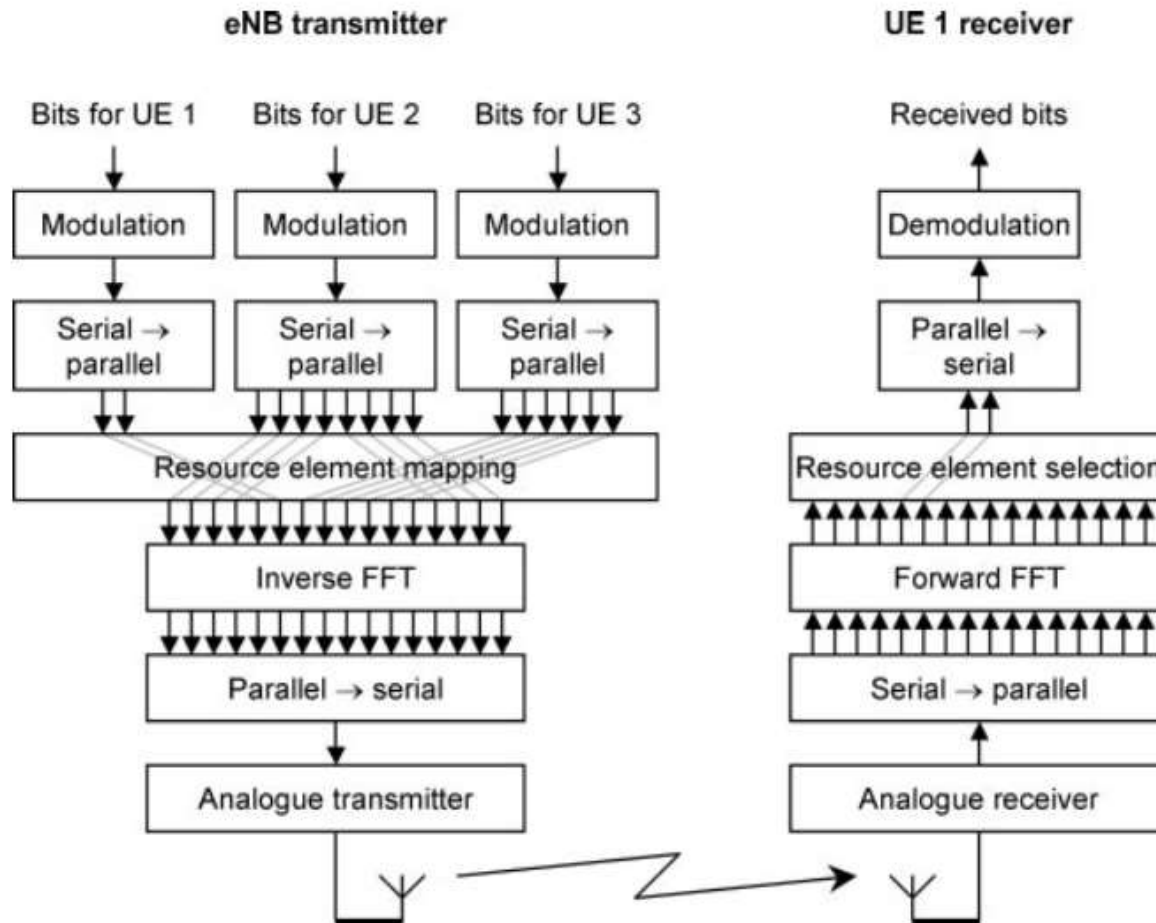
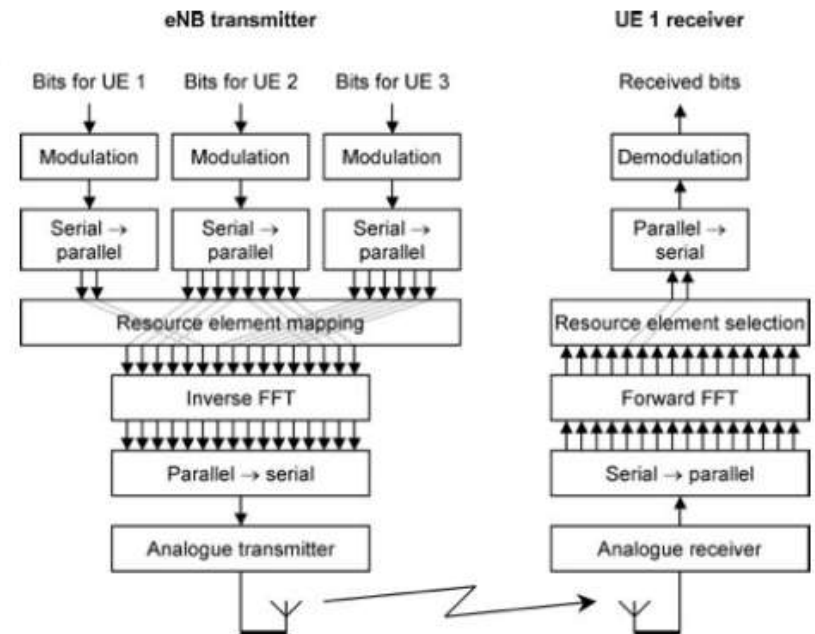
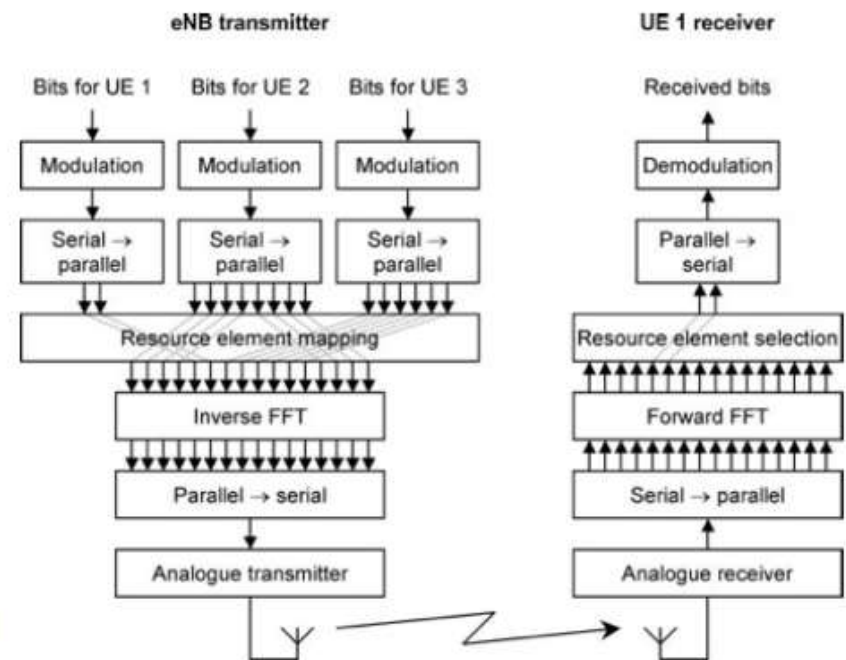


Figure 4.4 Initial block diagram of an OFDM transmitter and receiver.

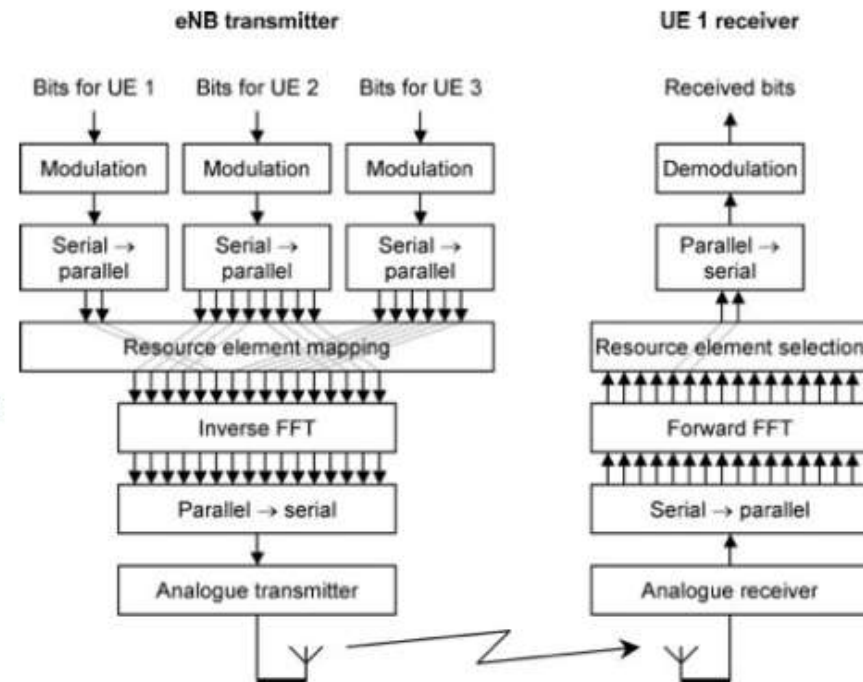
- In the diagram, the BS is sending streams of bits to three different mobiles
- It modulates each bit stream independently, possibly using a different modulation scheme for each one
- It then passes each symbol stream through a serial-to-parallel converter to divide it into sub-streams
- The num of sub-streams per mobile depends on the data rate
 - ✓ A voice application might only use a few sub-streams
 - ✓ A video application might use many more

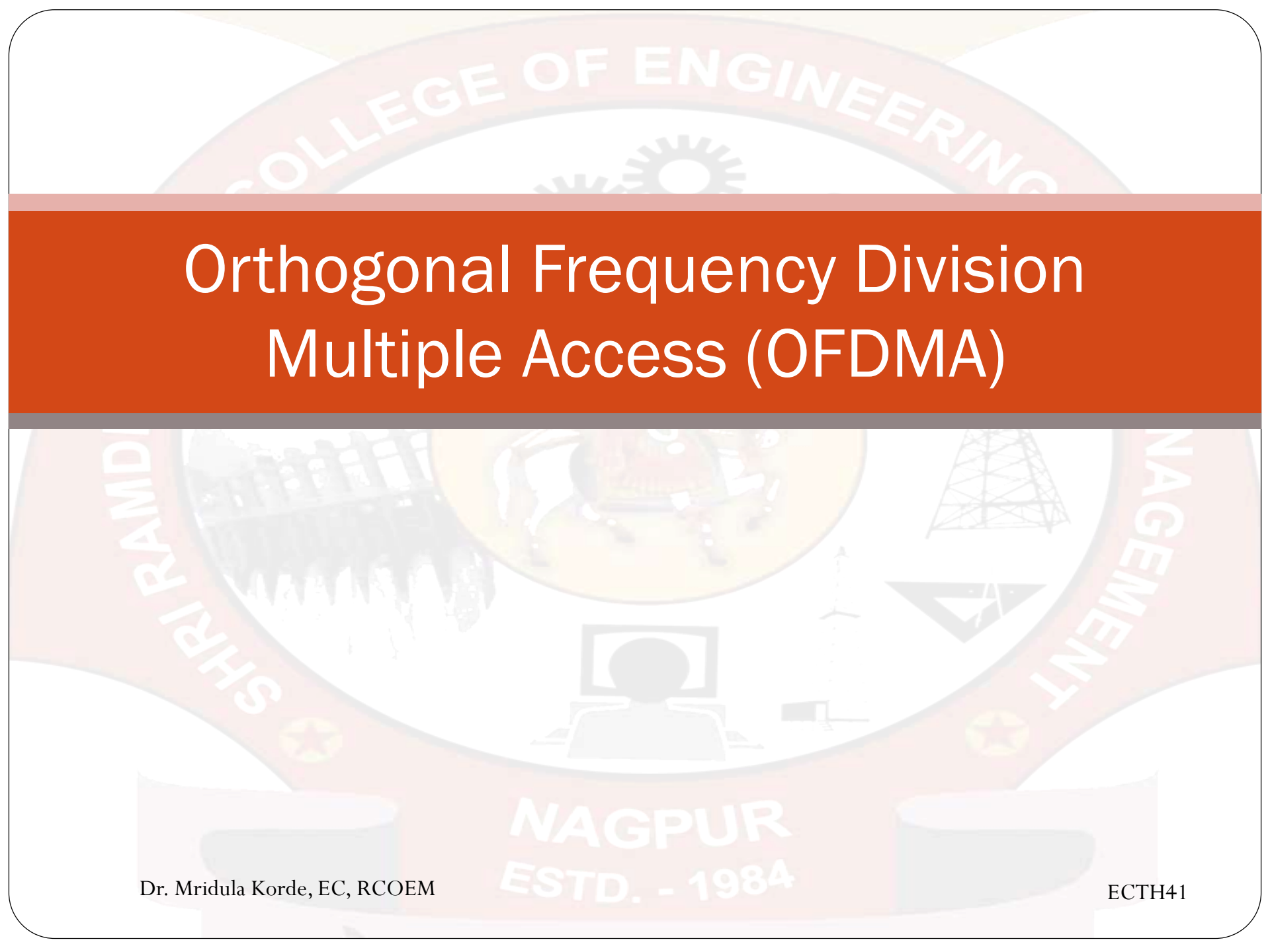


- The resource element mapper takes the individual sub-streams and chooses the subcarriers on which to transmit them
- A mobile's sub-carriers may lie in one contiguous block (as in the case of mobiles 1 and 3), or they may be divided (as for mobile 2)
- The resulting information is the amplitude and phase of each sub-carrier as a function of frequency
- By passing it through an inverse FFT, we can compute the in-phase and quadrature components of the corresponding time-domain waveform
- This can then be digitized, filtered and mixed up to radio frequency for transmission



- The mobile reverses the process
 - ✓ Starts by sampling the incoming signal, filtering it, and converting it down to baseband
 - ✓ Passes the data through a forward FFT, to recover the amplitude and phase of each sub-carrier
 - ✓ Assume that the BS has already told the mobile which sub-carriers to use
 - ▶ Using this knowledge, the mobile
 - Selects the required sub-carriers
 - Recovers the transmitted information
 - Discarding the remainder





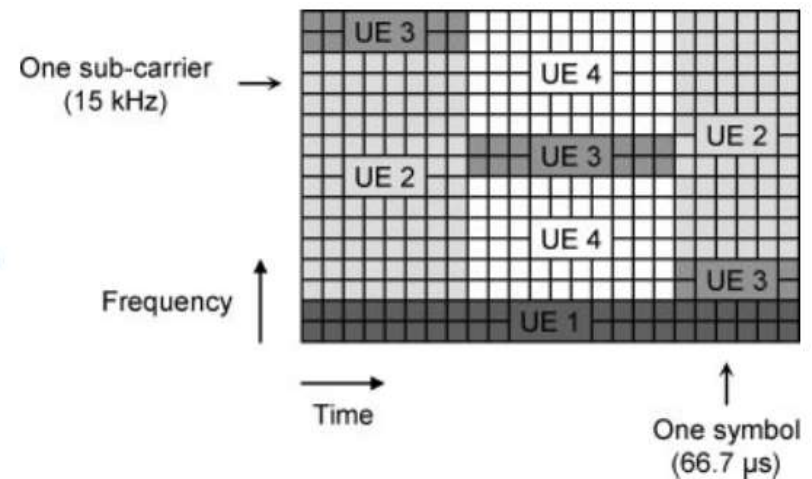
Orthogonal Frequency Division Multiple Access (OFDMA)

- Orthogonal Frequency Division Multiple Access (OFDMA) is an extension of OFDM to the implementation of a multiuser communication system.
- OFDMA distributes subcarriers to different users at the same time, so that multiple users can be scheduled to receive data simultaneously.
- OFDMA can also be used in combination with Time Division Multiple Access (TDMA), such that the resources are partitioned in the time-frequency plane – i.e. groups of subcarriers for a specific time duration.
- In LTE, such time-frequency blocks are known as Resource Blocks (RBs).

2.1 Multiple Access

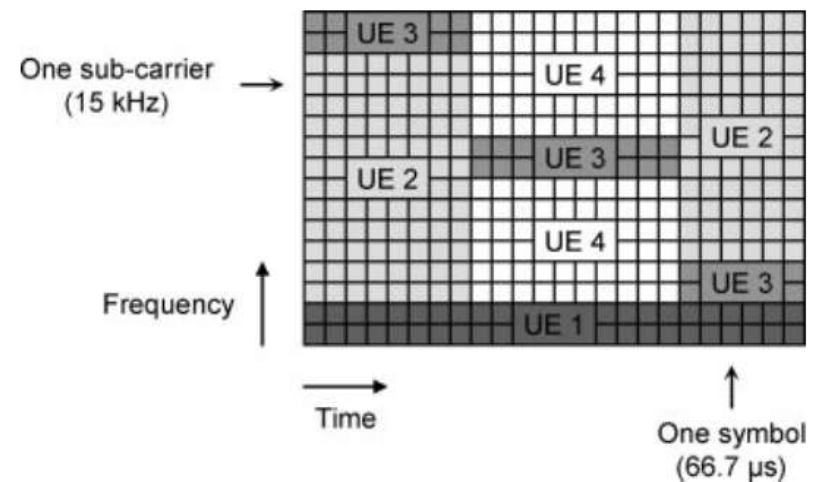
- Orthogonal Frequency Division Multiple Access (OFDMA)
- ✓ The BS shares its resources by transmitting to the mobiles at different times and frequencies, so as to meet the requirements of the individual applications

(Figure 4.5)

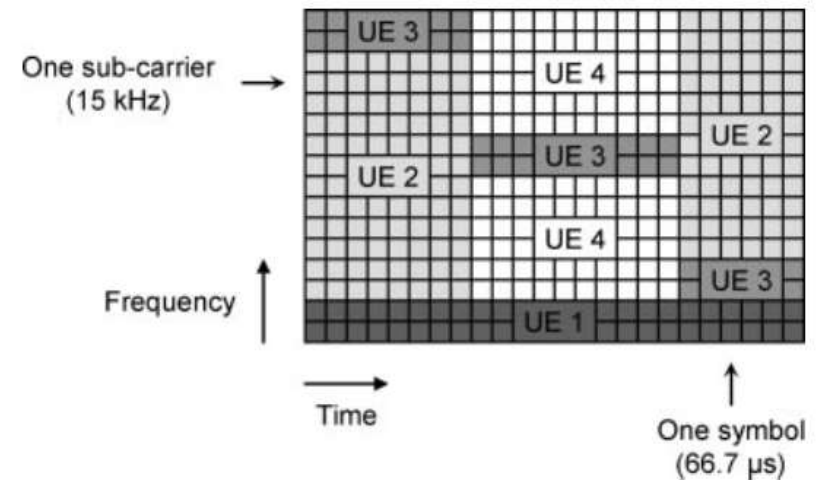


- Example

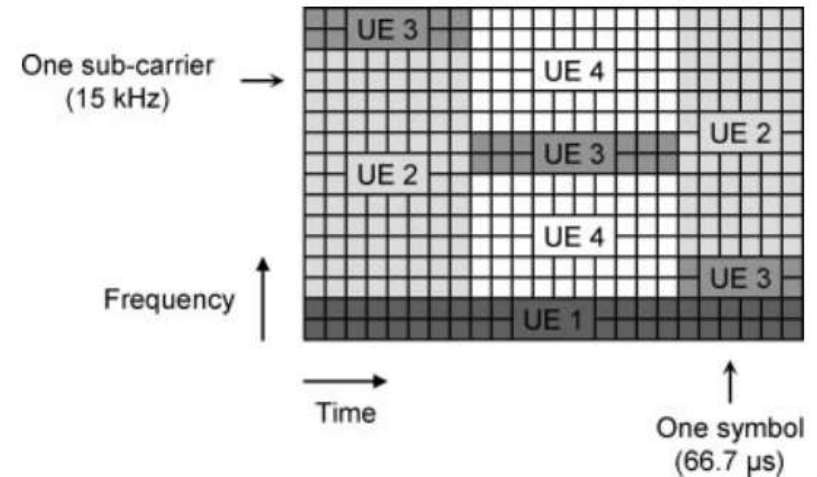
- ✓ Mobile 1 is receiving a voice over IP stream, so the data rate, and hence the num of sub-carriers, is low but constant
- ✓ Mobile 2 is receiving a stream of non real time packet data
- ✓ The average data rate is higher, but the data come in bursts, so the num of sub-carriers can vary



- The BS can also respond to frequency dependent fading, by allocating subcarriers on which the mobile is receiving a strong signal
- ✓ In the figure, mobile 3 is receiving a VoIP stream, but it is also affected by frequency dependent fading
- ✓ In response, the BS allocates subcarriers on which the mobile is receiving a strong signal, and changes this allocation as the fading pattern changes
- ✓ In a similar way, it can transmit to mobile 4 using two separate blocks of sub-carriers, which are separated by a fade



- By allocating sub-carriers in response to changes in the fading patterns, an OFDMA transmitter can greatly reduce the impact of time- and frequency-dependent fading
- The process requires feedback from the mobile



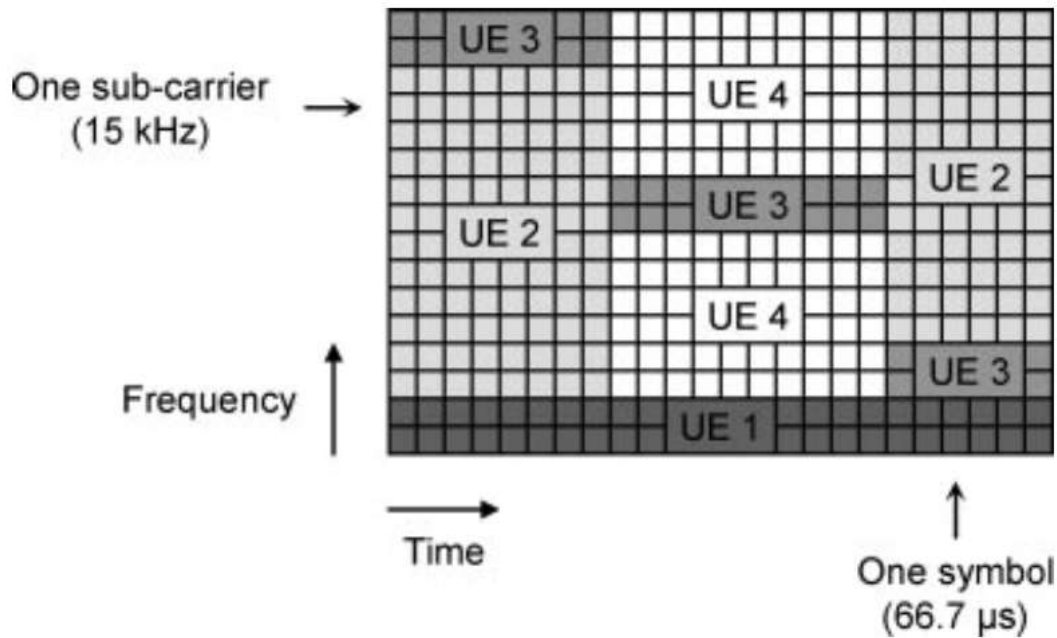


Figure 4.5 Implementation of time and frequency division multiple access when using OFDMA.

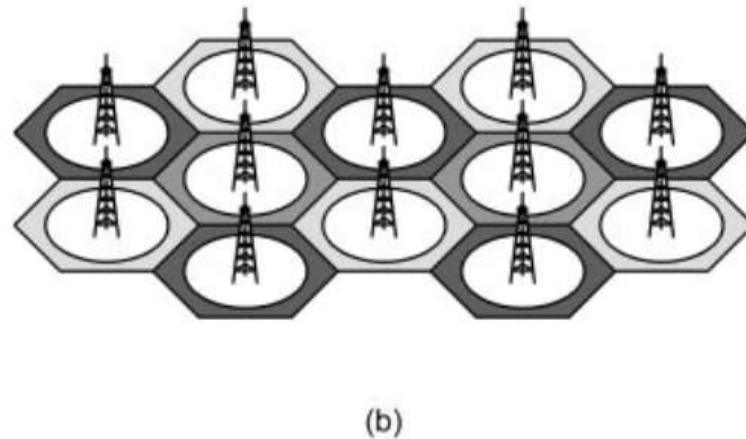
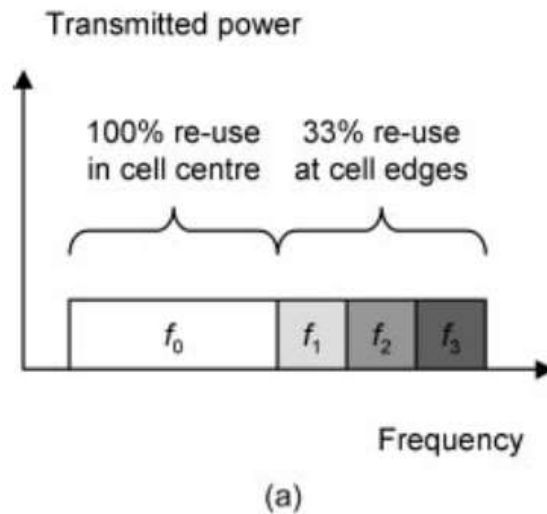
2.2 Fractional Frequency Re-Use

- In a mobile communication system
 - ✓ One BS can send information to a large number of mobiles
 - ✓ Every mobile has to receive a signal from one BS in the presence of interference from the others
- Need to minimize the interference, so that mobile can receive information successfully

- Previous systems have used two different techniques
 - ✓ GSM
 - ▶ Nearby cells transmit using different carrier frequencies
 - ▶ Typically, each cell might use a quarter of the total bandwidth, with a re-use factor of 25%
 - ▶ This technique reduces the interference between nearby cells, but it means that the frequency band is used inefficiently
 - ✓ UMTS
 - ▶ Each cell has the same carrier frequency, with a re-use factor of 100%
 - ▶ This technique uses the frequency band more efficiently than before, at the expense of increasing the interference in the system

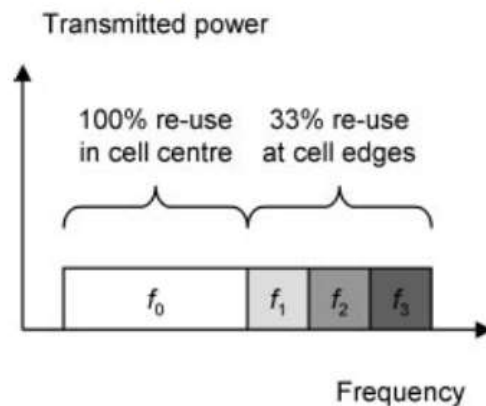
- LTE

- ✓ Every BS can transmit in the same frequency band
- ✓ It can allocate the sub-carriers within that band in a flexible way using fractional frequency re-use

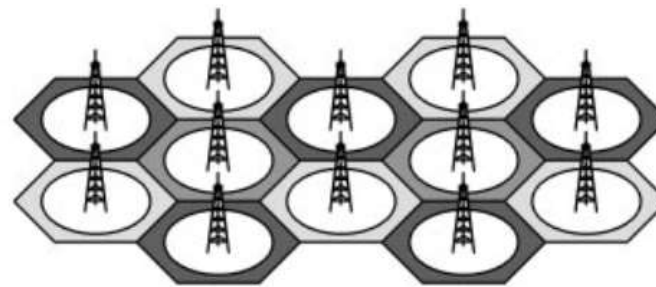


- Figure 4.6

- ✓ Every BS is controlling one cell and every cell is sharing the same frequency band
- ✓ Within that band, each cell transmits to nearby mobiles using the same set of sub-carriers, denoted f_0
- ✓ The mobiles are close to their respective BSs, so the received signals are strong enough to overwhelm any interference
- ✓ Distant mobiles receive much weaker signals, which are easily damaged by interference
 - ▶ To avoid this, neighboring cells can transmit to those mobiles using different sets of sub-carriers
 - ▶ Half the frequency band is reserved for nearby mobiles, while the remainder is divided into three sets, denoted f_1, f_2 and f_3 , for use by distant mobiles
 - ▶ Resulting re-use factor: 67%



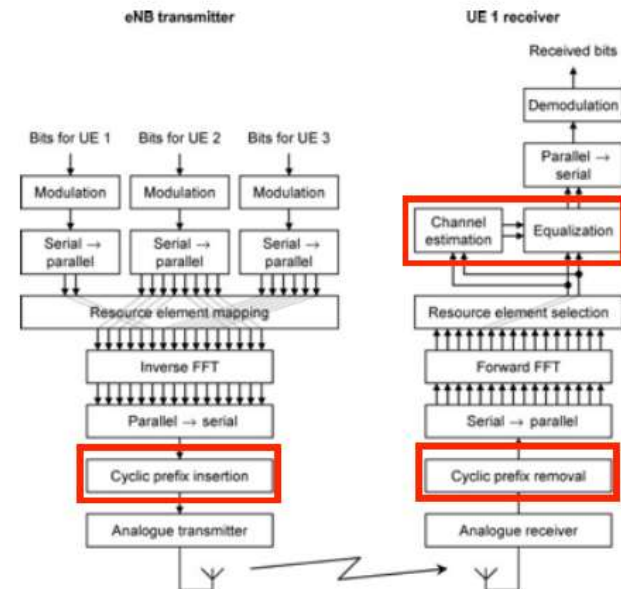
(a)



(b)

2.3 Channel Estimation

- Figure 4.7
 - ✓ A detailed block diagram of OFDMA
 - ✓ Two extra processes
 - ▶ The receiver contains the extra steps of channel estimation and equalization
 - ▶ The transmitter inserts a cyclic prefix into the data stream, which is then removed in the receiver



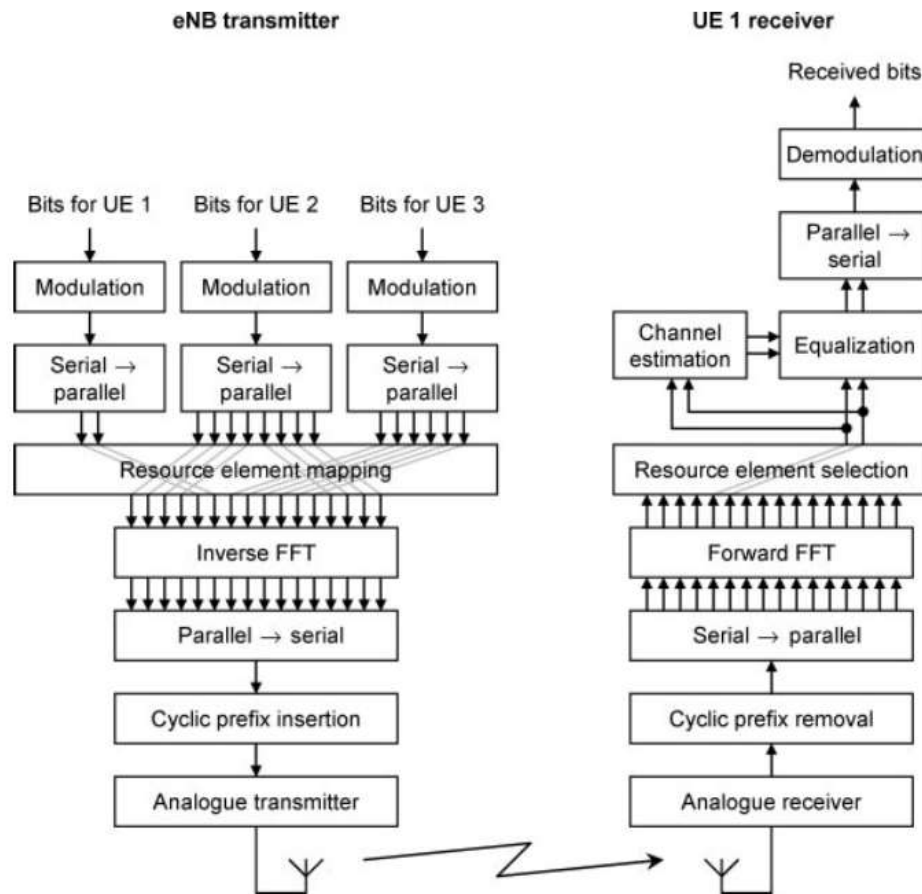
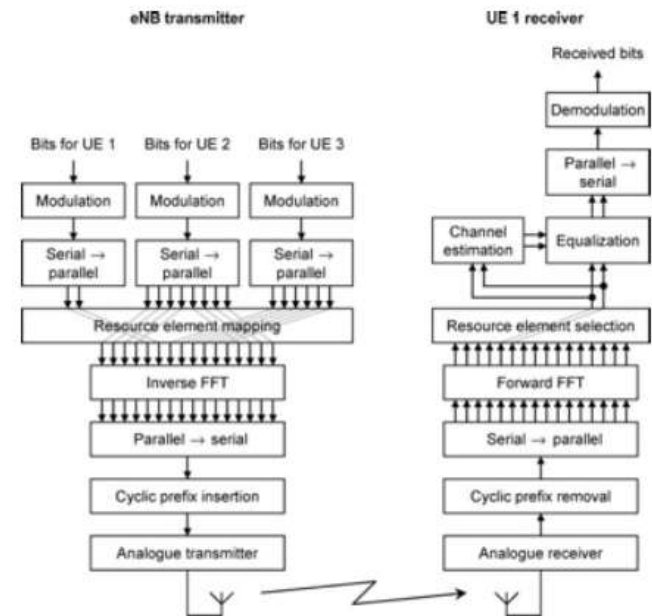


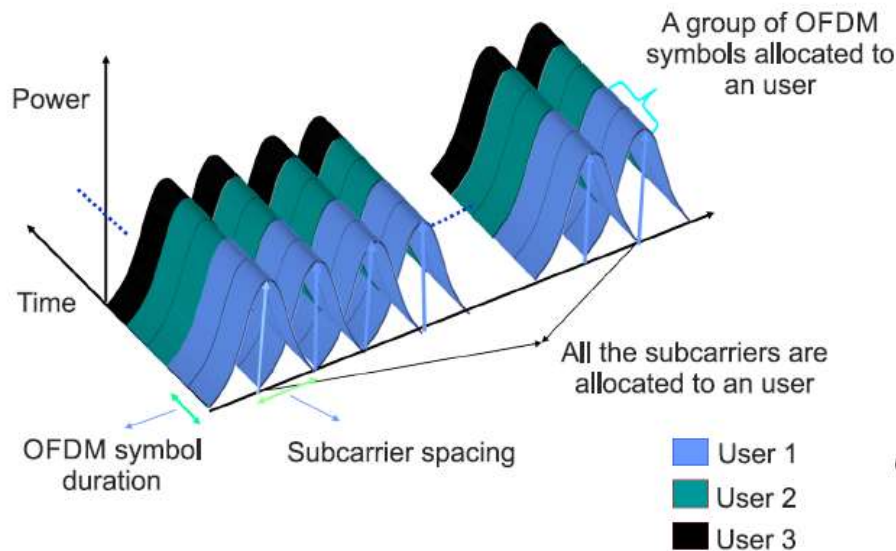
Figure 4.7 Complete block diagram of an OFDMA transmitter and receiver.

- Channel estimation
 - ✓ Each sub-carrier can reach the receiver with a completely arbitrary amplitude and phase
 - ✓ The OFDMA transmitter
 - ▶ Injects reference symbols into the transmitted data stream
 - ✓ The receiver
 - ▶ Measures the incoming reference symbols
 - ▶ Compares them with the ones transmitted
 - ▶ Uses the result to remove the amplitude changes and phase shifts from the incoming signal

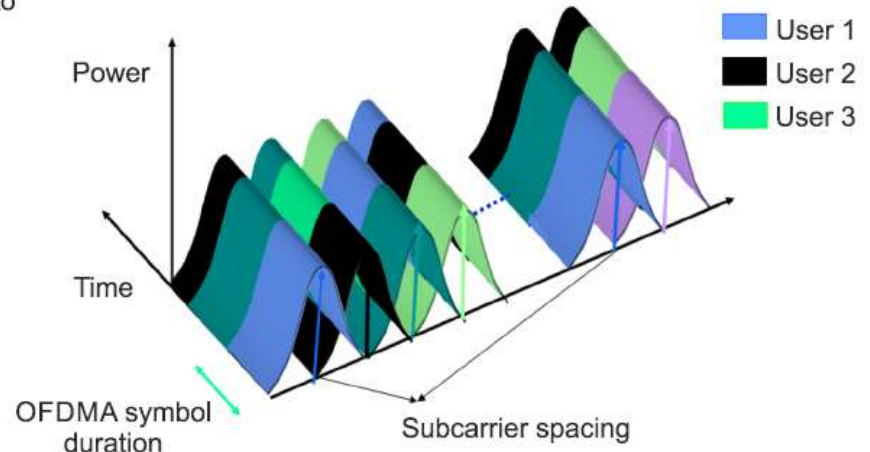


OFDM vs OFDMA

- **OFDMA** is the “**access**” version of OFDM, that is a scheme where the access to resources is shared between users.



(a) OFDM-TDMA Signal



(B) OFDMA Signal

- In **OFDM-TDMA**, time-slots are attributed to each users and the whole OFDM symbol is dedicated to one user.
- In **OFDMA**, both time and frequency are distributed among users.

OFDMA Compared to CDMA

	OFDMA	CDMA
Channel Equalisation	OFDMA chops its large bandwidth into subchannels that are easy to equalise.	CDMA requires a large bandwidth (more difficult to equalise than OFDM).
Mobility	Sensitive to Doppler shift.	Robust to Doppler shift.
Scalability	Easy to aggregate additional spectrum.	Hard to support more users.
Security	May require additional security measures	PN-codes are an additional level of security.

- Orthogonal Frequency Division Multiplexing (OFDM) is a special case of multicarrier transmission which is highly attractive for implementation.
- In OFDM, the non-frequency selective narrowband sub channels into which the frequency-selective wideband channel is divided are overlapping but orthogonal, as shown in Figure 5.1(b).

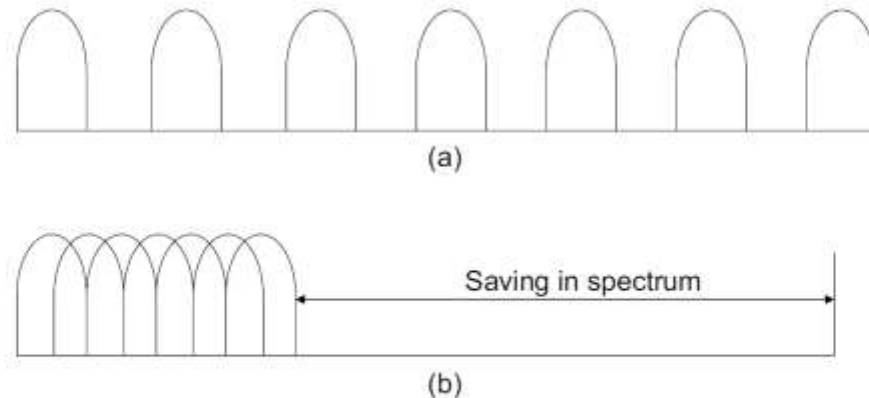
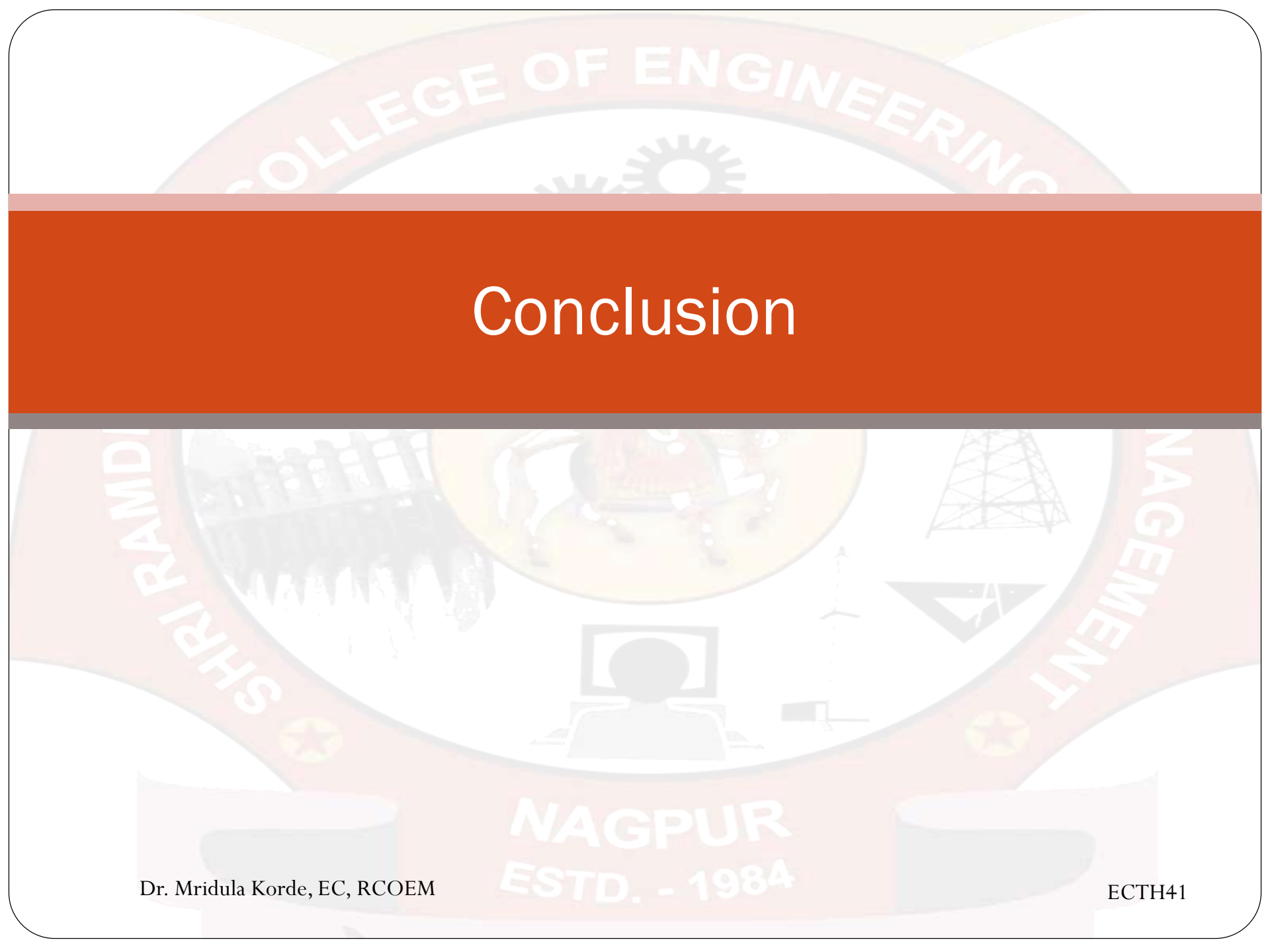


Figure 5.1 Spectral efficiency of OFDM compared to classical multicarrier modulation:
(a) classical multicarrier system spectrum; (b) OFDM system spectrum.



Conclusion

- Certain key parameters determine the performance of OFDM and OFDMA systems.
- As already mentioned, the CP should be longer than the channel impulse response in order to ensure robustness against ISI.
- OFDM is a mature technology.
- It achieves high transmission rates of broadband transmission, with low receiver complexity.
- It makes use of a CP to avoid ISI, enabling block-wise processing.
- It exploits orthogonal subcarriers to avoid the spectrum wastage associated with inter subcarrier guard-bands.
- It can be extended to a multiple-access scheme, OFDMA, in a straightforward manner.
- These factors together have made OFDMA the technology of choice for the LTE downlink.

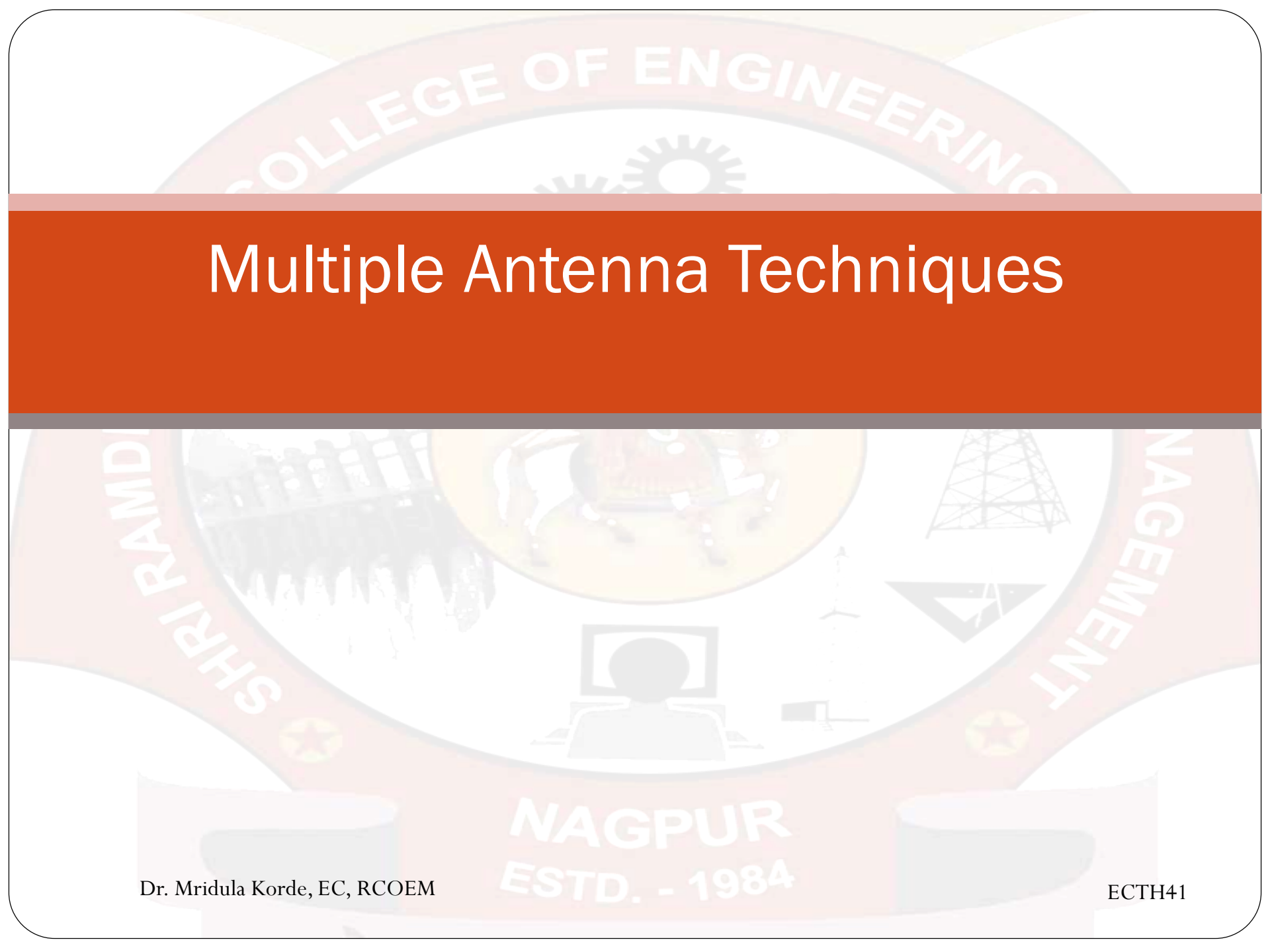


**SHRI RAMDEOBABA COLLEGE OF ENGINEERING AND
MANAGEMENT,
KATOL ROAD, NAGPUR, MAHARASHTRA, INDIA**

Communication System Analysis(ECTH 41)

Dr. (Mrs.) Mridula Korde

Associate Professor, Department of Electronics and
Communication Engineering, RCOEM, Nagpur

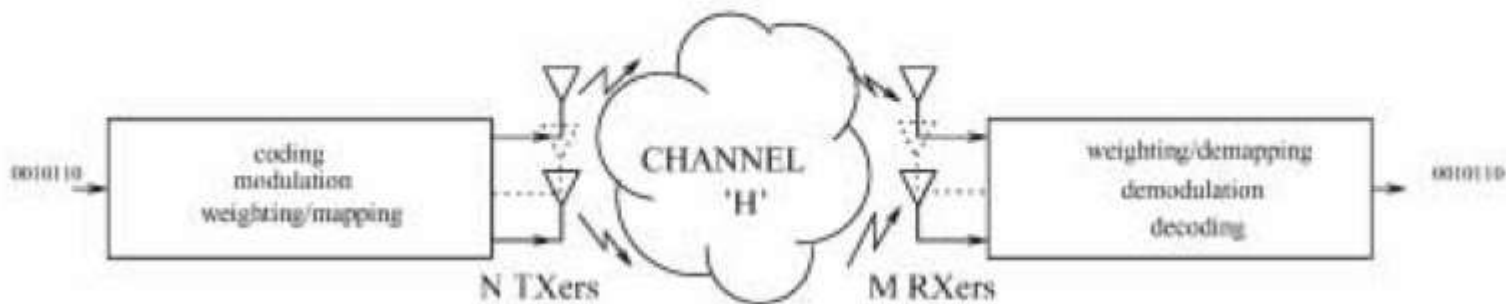


Multiple Antenna Techniques

Multiple-Input Multiple-Output (MIMO) Wireless Systems

1.1 What are MIMO systems ?

- A MIMO system consists of several antenna elements, plus adaptive signal processing, at both transmitter and receiver
- First introduced at Stanford University (1994) and Lucent (1996)
- Exploit multipath instead of mitigating it



1.2 Wireless channels limitations

Wireless transmission introduces:

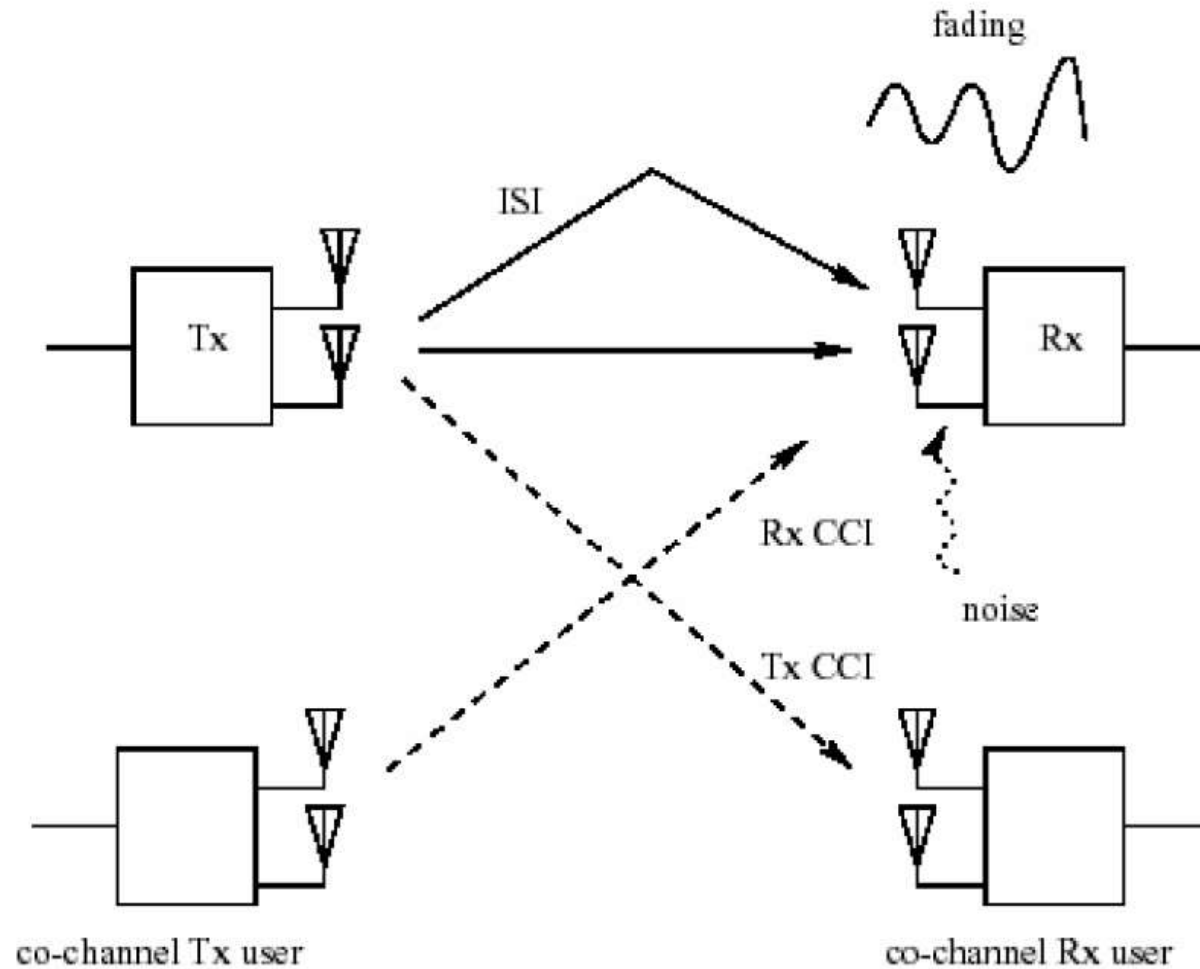
Fading: multiple paths with different phases add up at the receiver, giving a random (Rayleigh/Ricean) amplitude signal.

ISI: multiple paths come with various delays, causing intersymbol interference.

CCI: Co-channel users create interference to the target user

Noise: electronics suffer from thermal noise, limiting the SNR.

Wireless channels limitations : summary



Introduction

- Wireless networks are experiencing a huge increase in the delivered amount of data due to emerging applications: M2M, video, etc.
- Key problems: High data rates require extra spectrum and energy which are very scarce, scalability of devices and signal processing algorithms.
- Future networks will require techniques that can substantially increase the capacity (bits/Hz) whilst not requiring extra spectrum or extra energy.
- Massive MIMO is a potential solution to these problems:
 - Very large arrays with an order of magnitude higher number of sensors.
 - Deployment of devices (access points, mobile phones and tables) with a large number of antenna elements.
 - Huge multiplexing gains allowing an order of magnitude higher data rates.

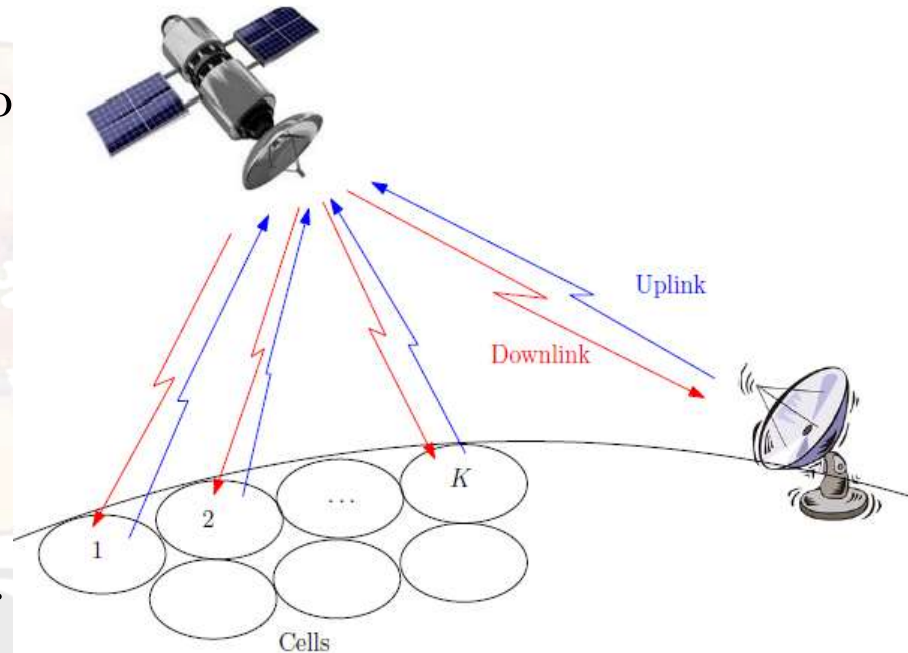
Introduction

Massive MIMO networks will be structure around the following elements:

- Antennas:
 - Reduction of RF chains and costs,
 - Compact antennas and mitigation of coupling effects.
- Electronic components:
 - Low-cost components such as power amplifiers and RF components.
 - Flexibility for different air interfaces and replacement of coaxial cables.
- Network architectures:
 - Heterogeneous networks, small cells and network MIMO,
 - Cloud radio access networks to help different devices.
- Protocols:
 - Scheduling and medium-access protocols for numerous heterogeneous users.
- Signal processing:
 - Transmit and receive processing,
 - Scalability and hardware implementation,
 - Integration between signal processing and RF devices to deal with impairments.

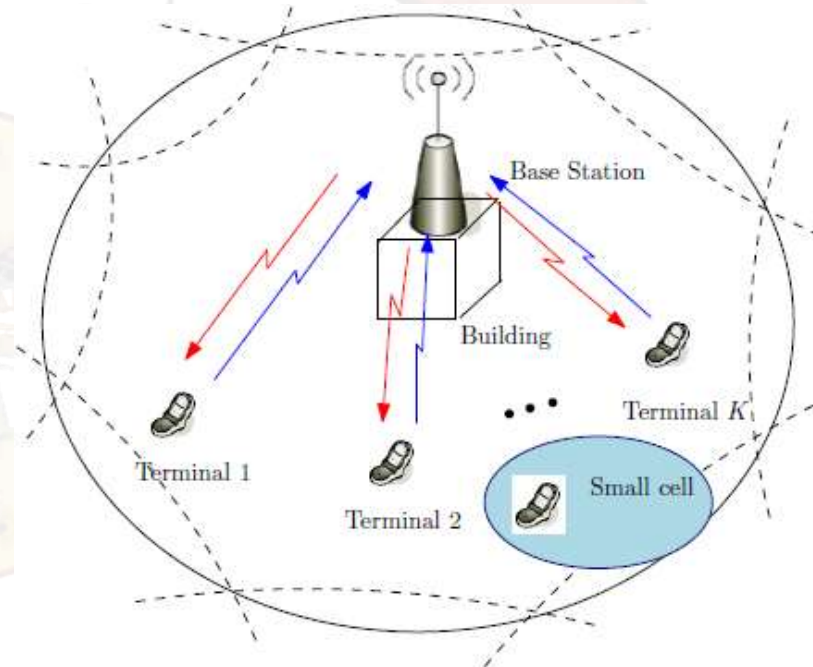
Application Scenarios: Satellite Networks

- Multi-beam satellite systems:
 - Clear and well-defined scenario for massive MIMO.
 - Coverage region is served by multiple spot beams intended for the users .
 - Beams are shaped by the antenna feeds forming part of the payload.
- Research problem:
 - interference caused by multiple adjacent spot beams that share the same frequency band.



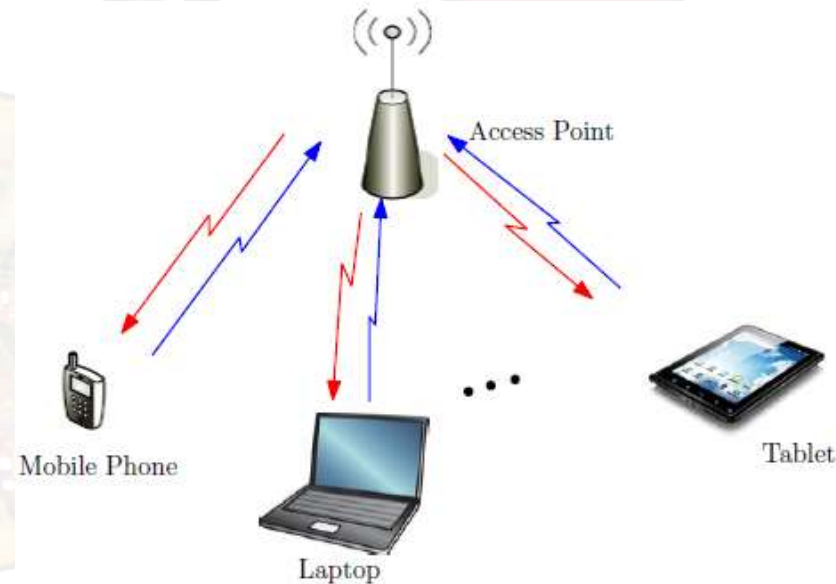
Application Scenarios: Mobile Cellular Networks

- 5G mobile cellular networks:
 - Base stations: very large arrays placed on rooftops and façades.
 - User terminals: phones, tablets with a significant number of antenna elements.
 - Compact antennas: mutual coupling
 - Coordination between cells.
 - Operation in TDD Mode.
- Research problems:
 - Uplink channel estimation: use of non-orthogonal pilots, the existence of adjacent cells and the coherence time of the channel require the system to reuse the pilots.
- Topics for investigation:
 - Design of channel estimation strategies that avoid pilot contamination.
 - Design of precoding and detection algorithms for Network MIMO with large arrays.



Application Scenarios: Local Area Networks

- **Future wireless local area networks:**
 - Tremendous increase in the last years with the proliferation of access points (APs) in hot spots and home users.
 - Massive MIMO: compact antennas, planar array geometries, etc.
- **Research problems:**
 - Coupling effects and I/Q imbalances.
 - Physical dimensions of APs and user devices.
- **Topics for investigation:**
 - Design of efficient precoding and detection algorithms.
 - Design of decoding algorithms with reduced delay.



MIMO: an Overview



- MIMO (Multiple-Input Multiple-Output)
 - Transmitter/receiver can have multiple antennas
 - A modern wireless communication technology
 - * Theory: late 1980's
 - * Standards and products: after 2000's
 - * Now: core feature in WLAN (802.11 WiFi) and cellular (3G LTE, WiMax)



MIMO: an Overview

- MIMO network architectures

- Single-user MIMO (in 802.11n-2009, LTE)

- * One TX, one RX. Either TX or RX or both can have multiple antennas



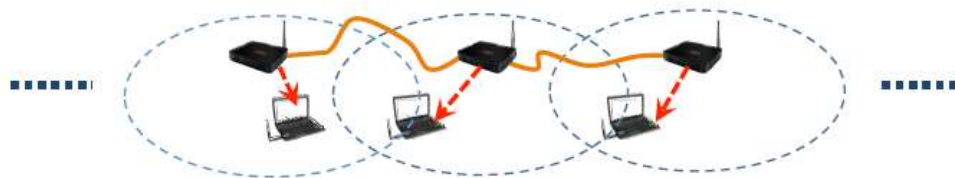
- Multi-user MIMO (in 802.11ac-2014, LTE-Advanced)

- * One TX, multiple RX. Parallel transmissions.



- Network MIMO (expected in near-future)

- * Multiple TX, multiple RX. Parallel transmissions.





SISO



SIMO



MISO



MIMO

Multi-Antenna Techniques

Use multiple antennas at the transmitter and/or receiver in combination with advanced signal processing techniques order to improve system performance:

- **Improved System Capacity:** More users per cell
- **Improved Coverage:** Possibility for larger cells
- **Improved Service Provisioning:** Higher per-user data rates

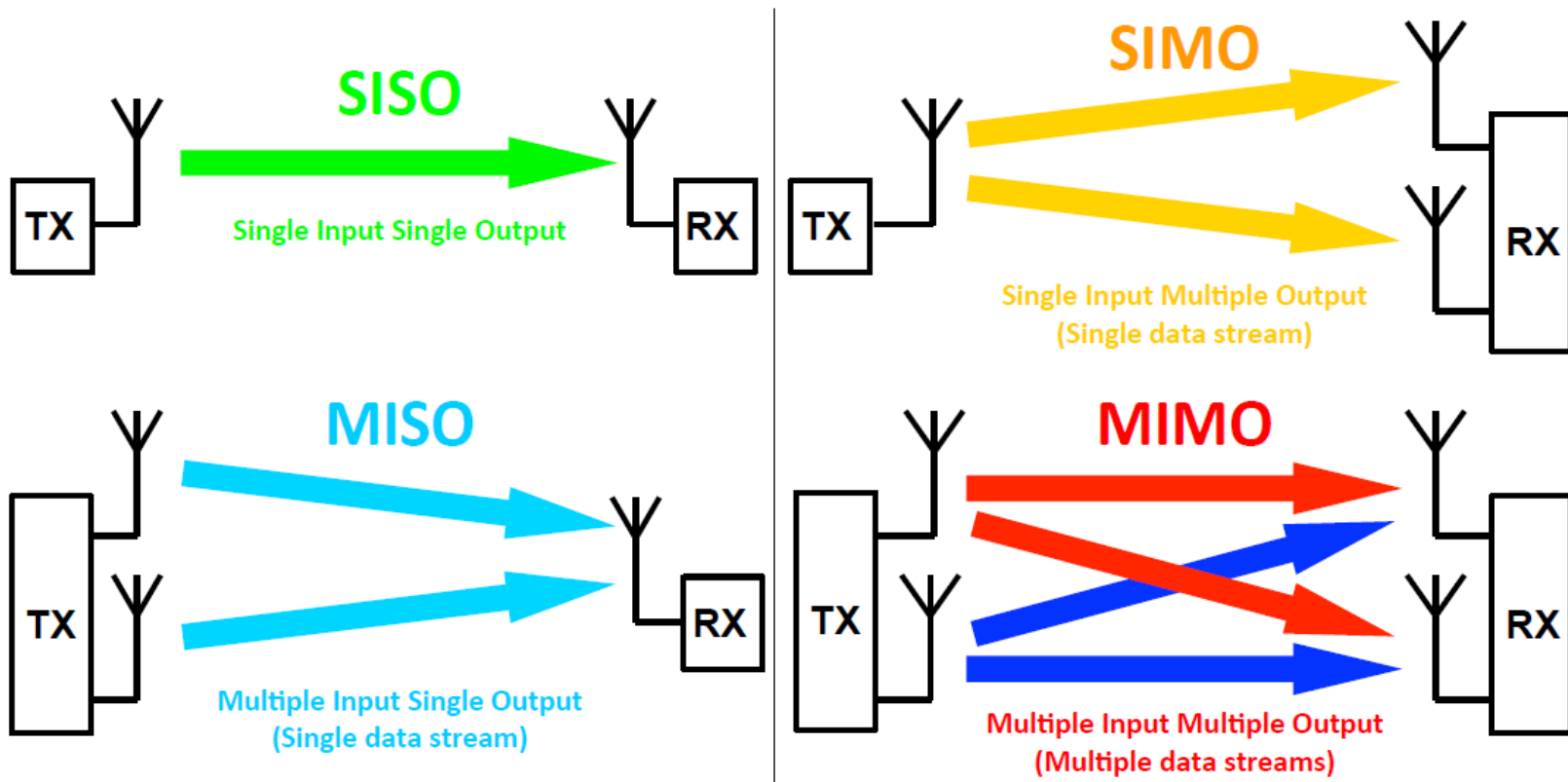
Distance and Correlation Between Antennas

In MIMO systems, distance between antenna elements affect the correlation of the radio-channel fading

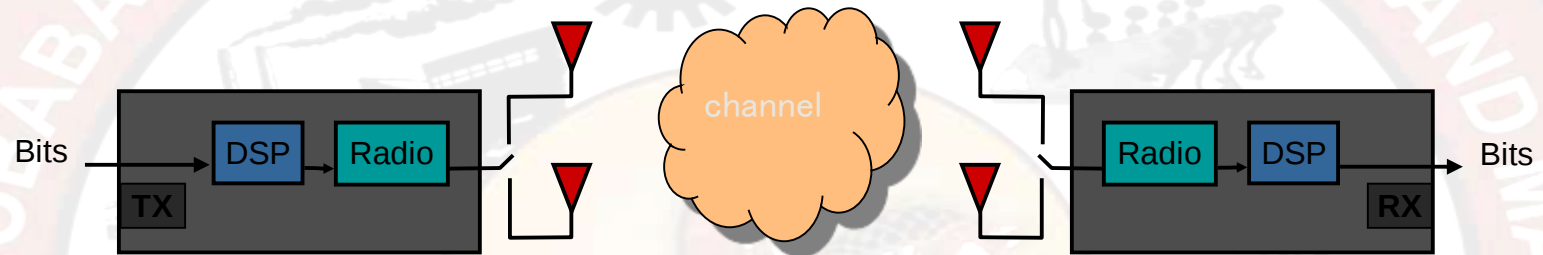
- **Low correlation:** Antennas are physically distanced far apart
- **High correlation:** Antennas are physically close to each other
- Different antenna configurations are suitable for different MIMO configurations: **diversity, beamforming, spatial multiplexing**
- For macro-cellular BS, distance of 10λ is sufficient for low correlation
- For mobile terminals, distance of 0.5λ is needed for low correlation
 - This distance is related to arrival angular spread of multi-path components which is typically much larger for UEs than eNBs

Single-user MIMO

- Basic communication modes



Conventional (SISO) Wireless Systems



Conventional “Single Input Single Output” (SISO) systems were favored for simplicity and low-cost but have some shortcomings:

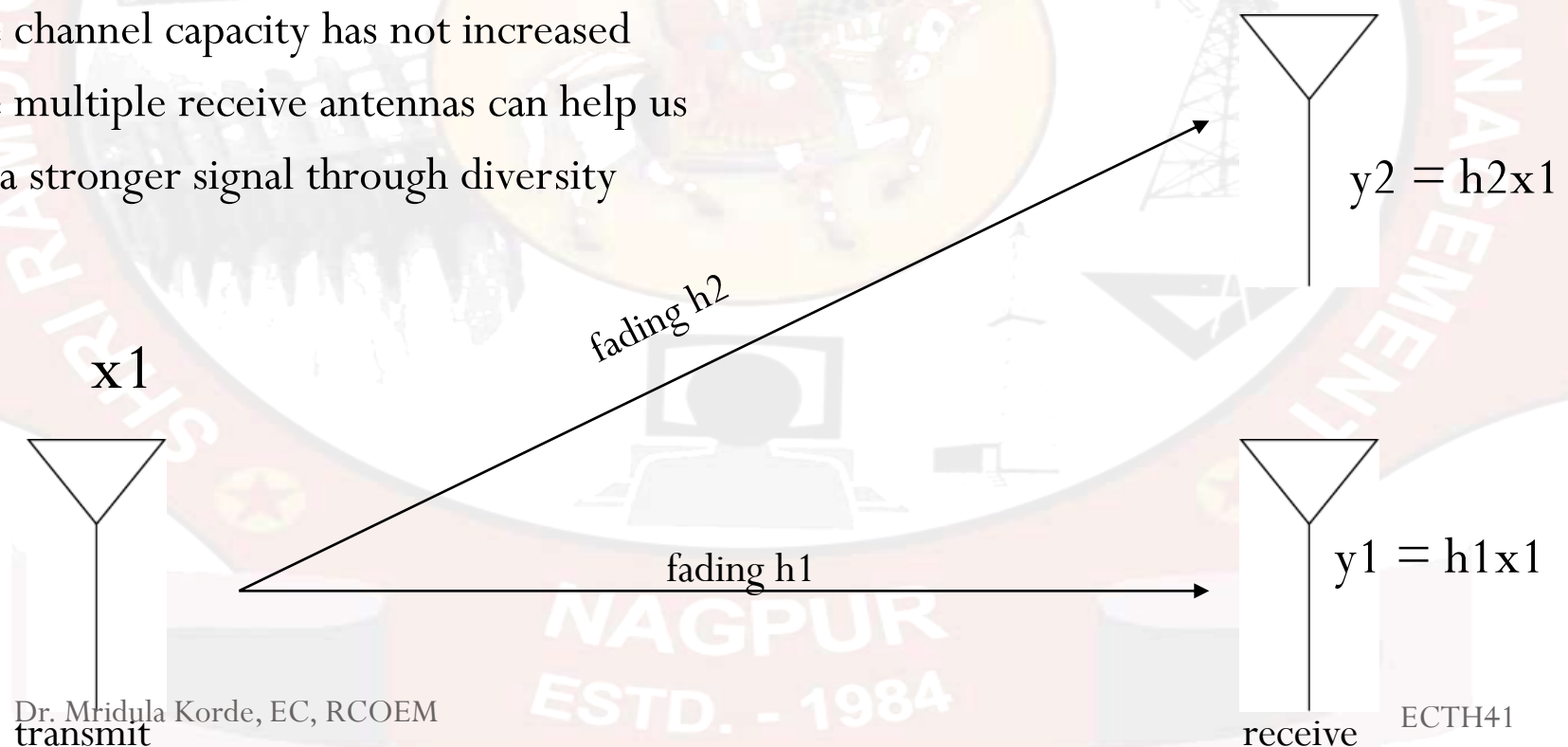
- Outage occurs if antennas fall into null
 - Switching between different antennas can help
- Energy is wasted by sending in all directions
 - Can cause additional interference to others
- Sensitive to interference from all directions
- Output power limited by single power amplifier

SIMO Single Input Multiple Output

Now assume we have two receiving antennas. There will be two received signals y_1 and y_2 with different fading coefficients h_1 and h_2 . The effect upon the signal x for a given path (from a transmit antenna to a receive antenna) is called a **channel**.

The channel capacity has not increased

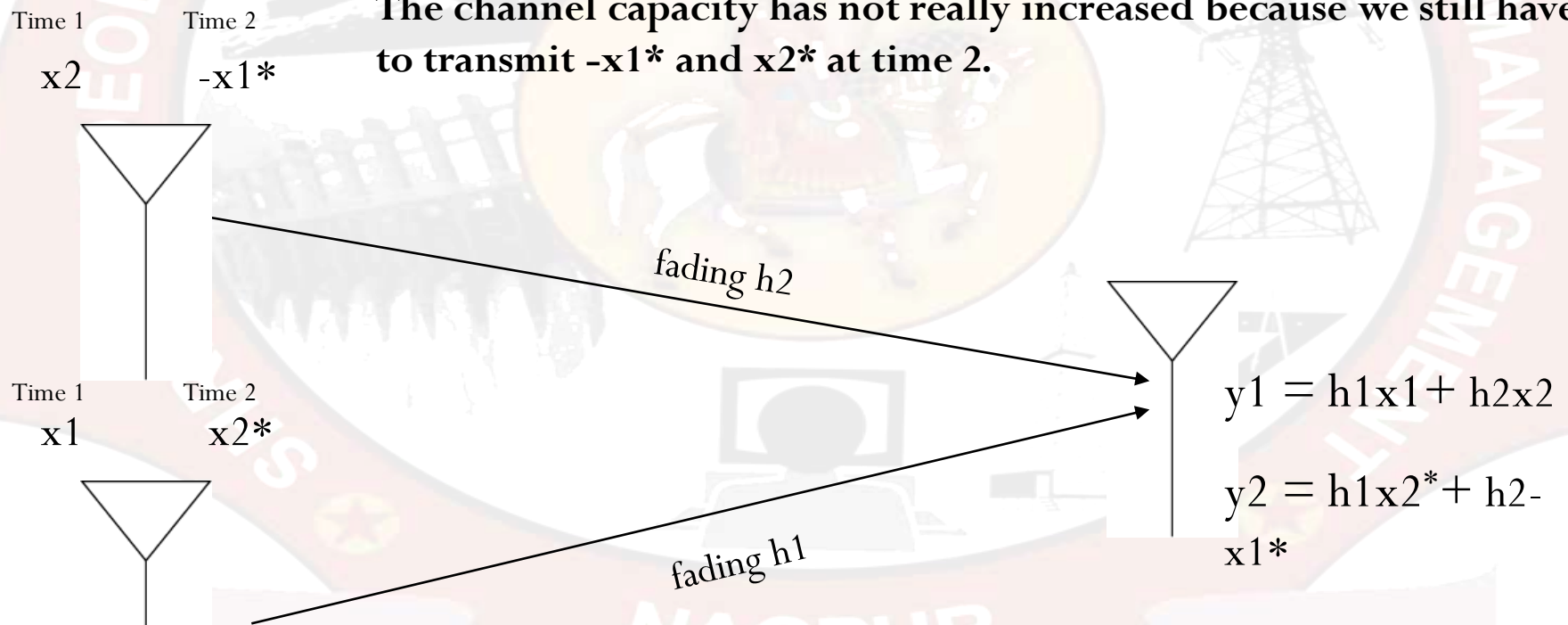
The multiple receive antennas can help us
get a stronger signal through diversity



MISO Multiple Input Single Output

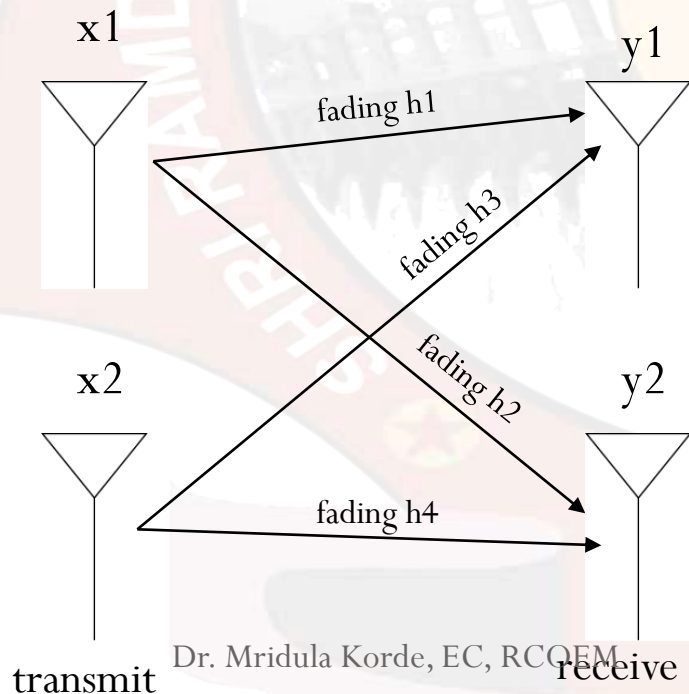
Assume 2 transmitting antennas and 1 receive antenna. There will be one received signal y_1 (sum of $x_1 h_1$ and $x_2 h_2$). In order to separate x_1 and x_2 we will need to also transmit, at a different time, $-x_1^*$ and x_2^* .

The channel capacity has not really increased because we still have to transmit $-x_1^*$ and x_2^* at time 2.



MIMO Multiple Input Multiple Output

With 2 transmitting antennas and 2 receiving antennas, we actually add a degree of freedom! Its quite simple and intuitive. However, in this simple model, we are assuming that the h coefficients of fading are independent, and uncorrelated. If they are correlated, we will have a hard time finding an approximation for the inverse of H . In practical terms, this means that we cannot recover x_1 and x_2 .



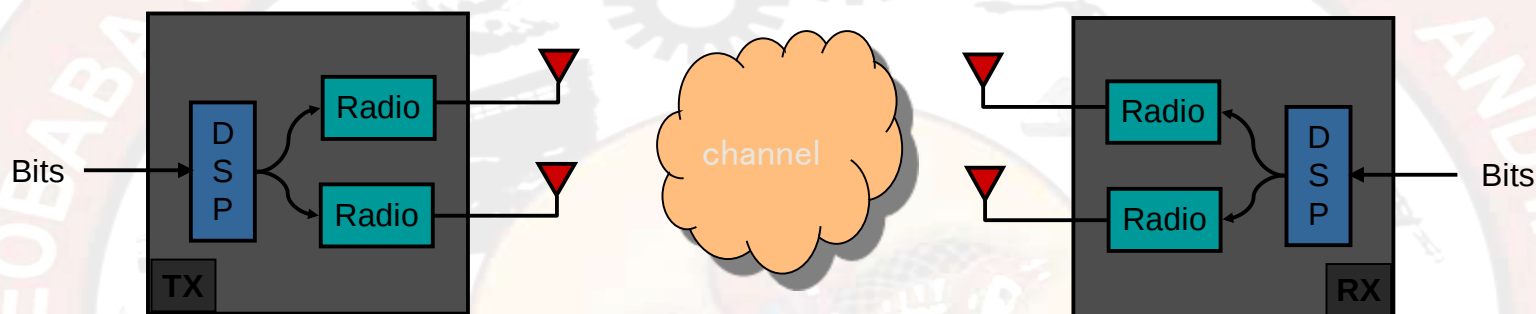
$$\begin{aligned} y_1 &= h_1 x_1 + h_2 x_2 \\ y_2 &= h_3 x_1 + h_4 x_2 \end{aligned} \quad \begin{vmatrix} y_1 \\ y_2 \end{vmatrix} = \begin{vmatrix} h_1 & h_2 \\ h_3 & h_4 \end{vmatrix} \begin{vmatrix} x_1 \\ x_2 \end{vmatrix} + \begin{vmatrix} w_1 \\ w_2 \end{vmatrix}$$

Finally Assume there is some white Gaussian Noise, and we have a set of linear equations

$$y = Hx + w$$

All 2 degrees of freedom are being utilized in the MIMO case, giving us Spatial Multiplexing.

MIMO Wireless Systems



Multiple Input Multiple Output (MIMO) systems with multiple parallel radios improve the following:

- Outages reduced by using information from multiple antennas
- Transmit power can be increased via multiple power amplifiers
- Higher throughputs possible
- Transmit and receive interference limited by some techniques

Multi-Antenna Techniques

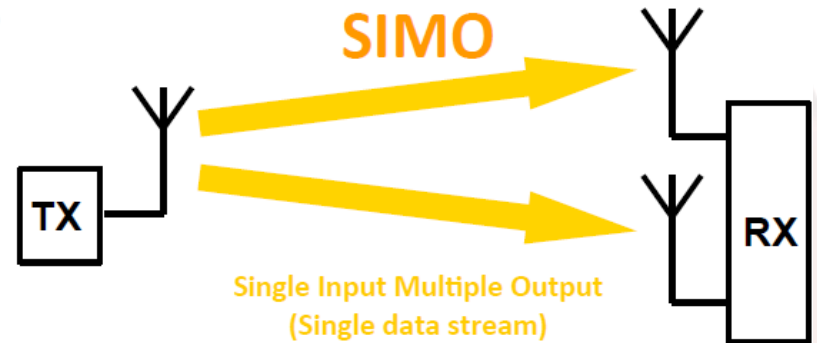
Diversity: Multiple antennas at Tx/Rx side can be used to provide additional diversity against fading. In this case, channel experienced at different antennas should have small mutual correlation, i.e., antennas should be distanced far apart

Beamforming: Multiple antennas can be used to shape the overall antenna beam (both Tx and Rx beam) to obtain directional gain. This can be done based on either high or low channel correlation

Spatial Multiplexing: Multiple antennas can be used to create multiple parallel communication channels. This provides opportunity for very high bandwidth utilization (high data rates in limited band)

Receiver Diversity

- Receiver coherently combines signals received by multiple antennas



- Asymptotic gain: Increasing SNR proportionally to N_r (#of receive antennas)
 - * Intuition: received signal power adds up
- What's the capacity gain?
 - * Logarithmically, according to Shannon's equation: $C = B \log(1 + \text{SNR})$
 - * When SNR is low, $\log(1 + \text{SNR}) \approx \text{SNR}$, so gain is almost linear w.r.t. N_r

Diversity Processing

- *Receive Diversity*
- Receive diversity is most often used in the uplink.
- *Here, the* base station uses two antennas to pick up two copies of the received signal.
- The signals reach the receive antennas with different phase shifts, but these can be removed by antenna-specific channel estimation.
- The base station can then add the signals together in-phase, without any risk of destructive interference between them

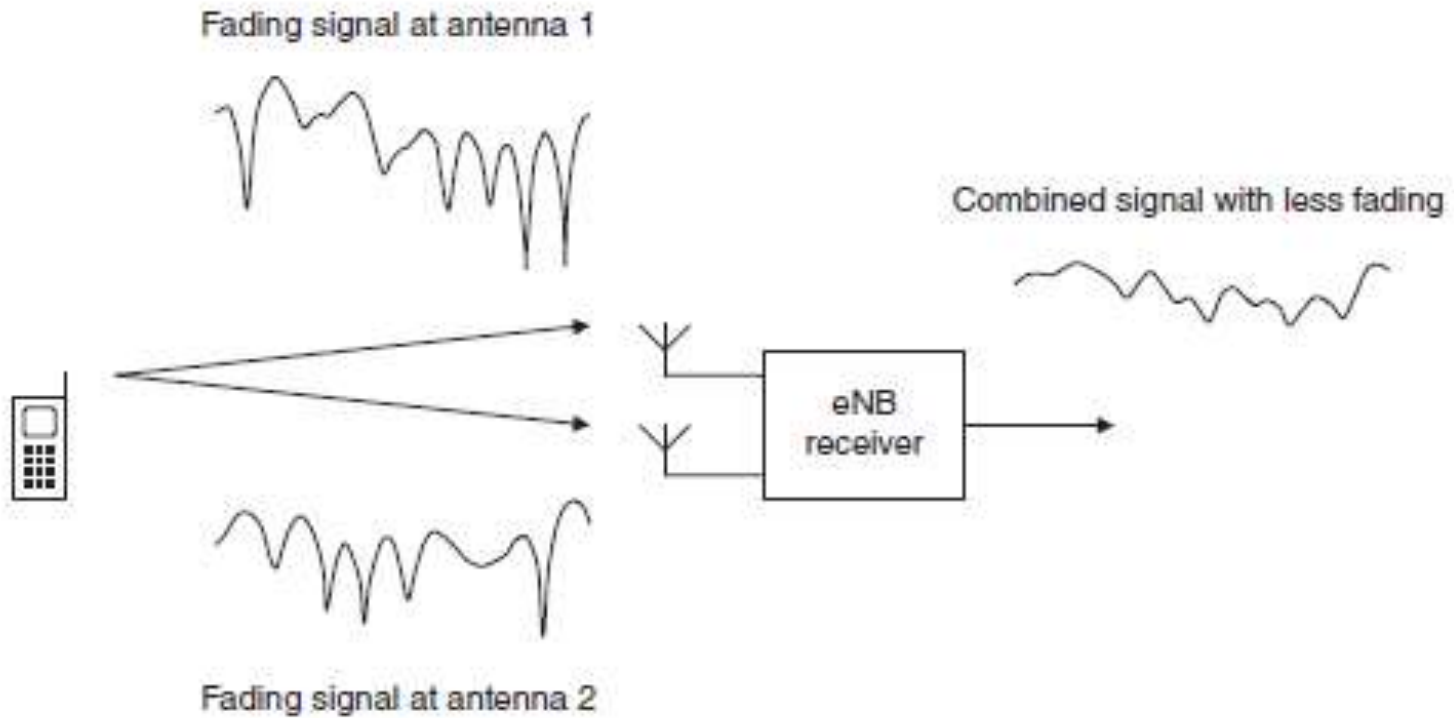
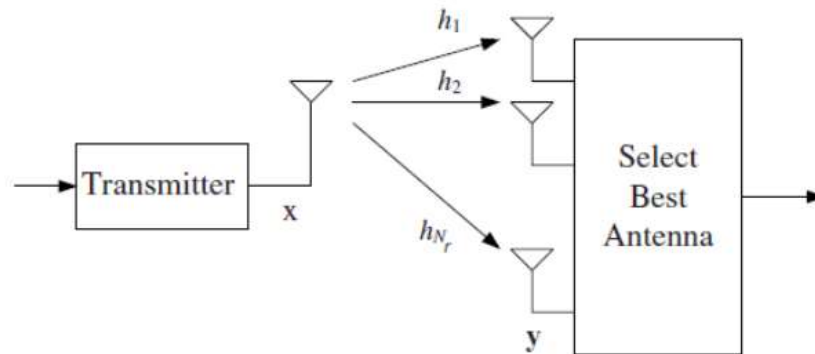


Figure 5.1 Reduction in fading by the use of a diversity receiver

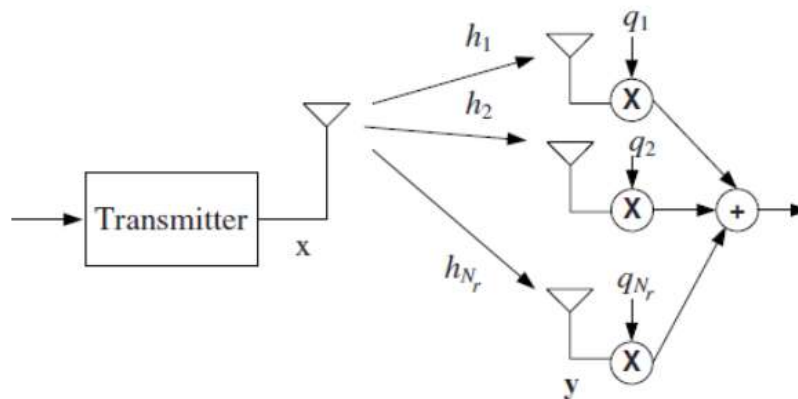
- The signals are both made up from several smaller rays, so they are both subject to fading.
- If the two individual signals undergo fades at the same time, then the power of the combined signal will be low.
- But if the antennas are far enough apart (a few wavelengths of the carrier frequency), then the two sets of fading geometries will be very different, so the signals will be far more likely to undergo fades at completely different times.
- Base stations usually have more than one receive antenna.
- In LTE, the mobile's test specifications assume that the mobile is using two receive antennas, so LTE systems are expected to use receive diversity on the downlink as well as the uplink.

Implementation of Receiver Diversity



* Selection combining

Improves SNR to
$$avgSNR \cdot \left(1 + \frac{1}{2} + \dots + \frac{1}{N_r}\right)$$



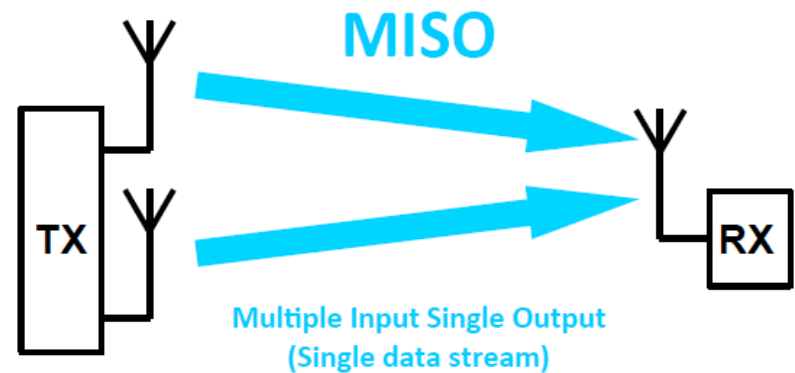
* Maximum Ratio combining

Improves SNR to

$$\sum_{i=1}^{N_r} SNR_i$$

Transmit Diversity

- Transmitter sends multiple versions of the same signal, through multiple antennas



- Two modes of transmit diversity
 - * Open-loop transmit diversity
 - * Closed-loop transmit diversity

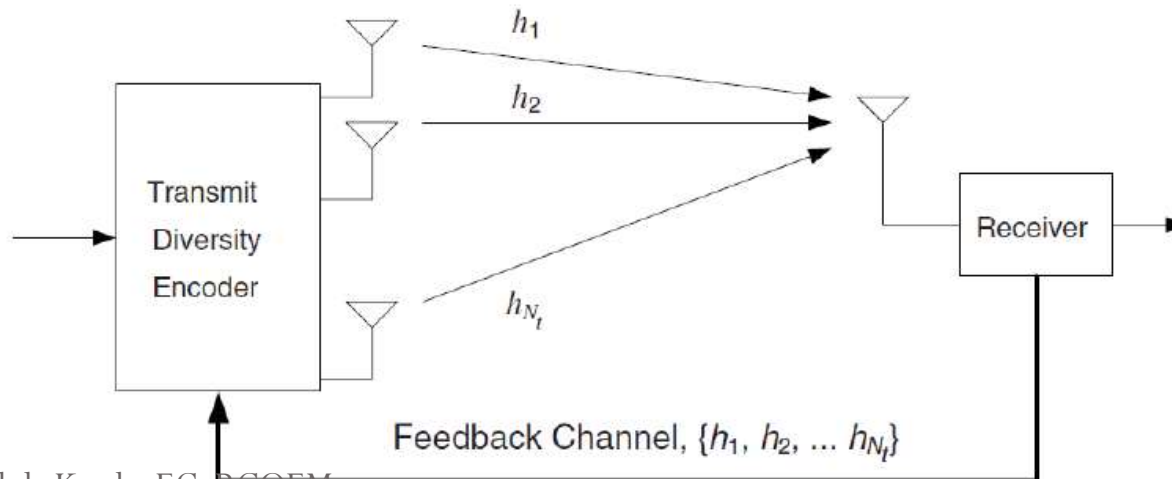
Types of Diversity

- Closed Loop Transmit Diversity
- Open Loop Transmit Diversity

Closed-loop transmit diversity

- Principle

- Send redundant versions of the same signal (symbol), over the same time slot
- Encode the symbols differently for different TX antennas
 - * i.e., weight the symbols on different antennas, following a precoding algorithm
 - * Precoding design requires feedback of channel state information (CSI)



Closed Loop Transmit Diversity

- *Transmit diversity reduces the amount of fading by using two or more antennas at the transmitter.*
- *In closed loop transmit diversity, the transmitter sends two copies of the signal in the expected way, but it also applies a phase shift to one or both signals before transmission.*
- *By doing this, it can ensure that the two signals reach the receiver in-phase, without any risk of destructive interference.*
- *The phase shift is determined by a *pre-coding matrix indicator (PMI)*, which is calculated by the receiver and fed back to the transmitter.*
- *A simple PMI might indicate two options: either transmit both signals without any phase shifts or transmit the second with a phase shift of 180° .*
- *If the first option leads to destructive interference, then the second will automatically work.*

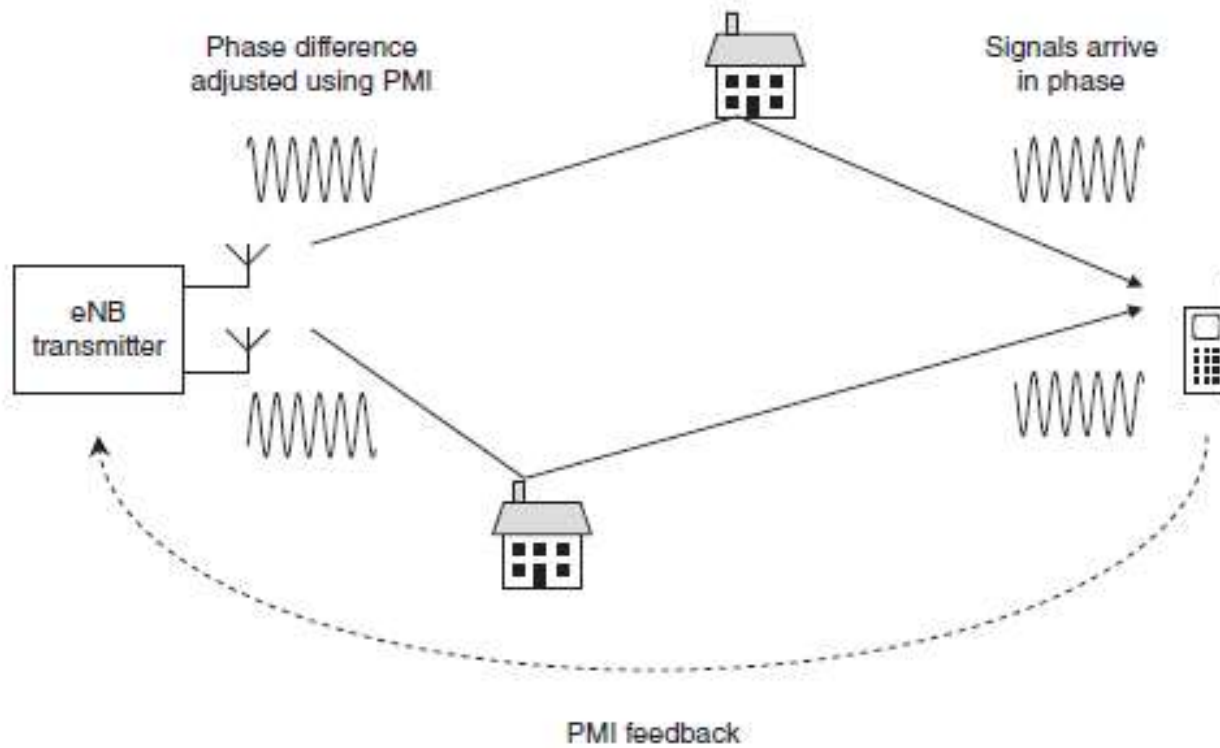


Figure 5.2 Operation of closed loop transmit diversity

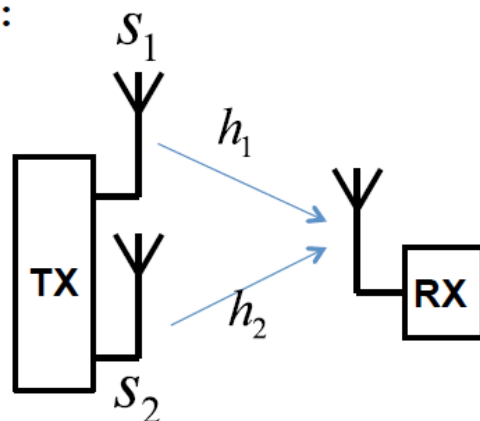
- The phase shifts introduced by the radio channel depend on the wavelength of the carrier signal and hence on its frequency.
- This implies that the best choice of pre coding matrix algorithm (PMI) is a function of frequency as well.
- The best choice of PMI also depends on the position of the mobile, so a fast moving mobile will have a PMI that frequently changes.
- Unfortunately the feedback loop introduces time delays into the system, so in the case of fast moving mobiles, the PMI may be out of date by the time it is used.
- For this reason, closed loop transmit diversity is only suitable for mobiles that are moving sufficiently slowly.
- For fast moving mobiles, it is better to use the open loop technique

Open Loop Transmit Diversity

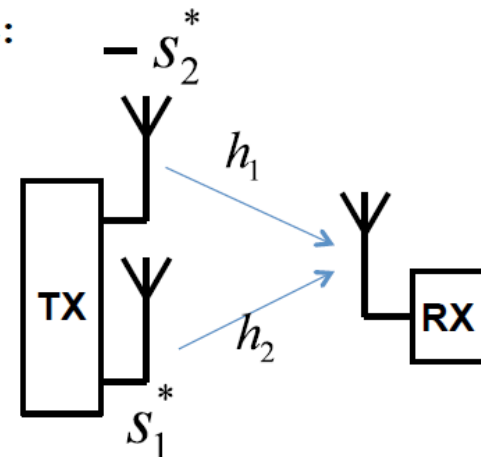
- Example: 2 TX antenna STBC

- Send two data symbols, s_1 and s_2

Time slot 1:



Time slot 2:



- Received signals:

$$r(t_1) = h_1 s_1 + h_2 s_2$$

$$r(t_2) = -h_1 s_2^* + h_2 s_1^*$$

Open Loop Transmit Diversity

- Here, the transmitter uses two antennas to send two symbols, denoted as s_1 and s_2 , in two successive time steps.
- In the first step, the transmitter sends s_1 from the first antenna and s_2 from the second, while in the second step, it sends $-s_2^*$ from the first antenna and s_1^* from the second. (The symbol $*$ indicates that the transmitter should change the sign of the quadrature component, in the process of complex conjugation.)
- The receiver can now make two successive measurements of the received signal, which correspond to two different combinations of s_1 and s_2 .
- It can then solve the resulting equations, so as to recover the two transmitted symbols.
- There are only two requirements: the fading patterns must stay roughly the same between the first time step and the second, and the two signals must not undergo fades at the same time.

Spatial Multiplexing

- ***Principles of Operation***
- *Spatial multiplexing has a different purpose from diversity processing.*
- *If the transmitter and receiver both have multiple antennas, then we can set up multiple parallel data streams between them, so as to increase the data rate.*
- *In a system with NT transmit and NR receive antennas, often known as an $NT \times NR$ spatial multiplexing system, the peak data rate is proportional to $\min(NT, NR)$.*
- *In the transmitter, the antenna mapper takes symbols from the modulator two at a time, and sends one symbol to each antenna.*
- *The antennas transmit the two symbols simultaneously, so as to double the transmitted data rate.*

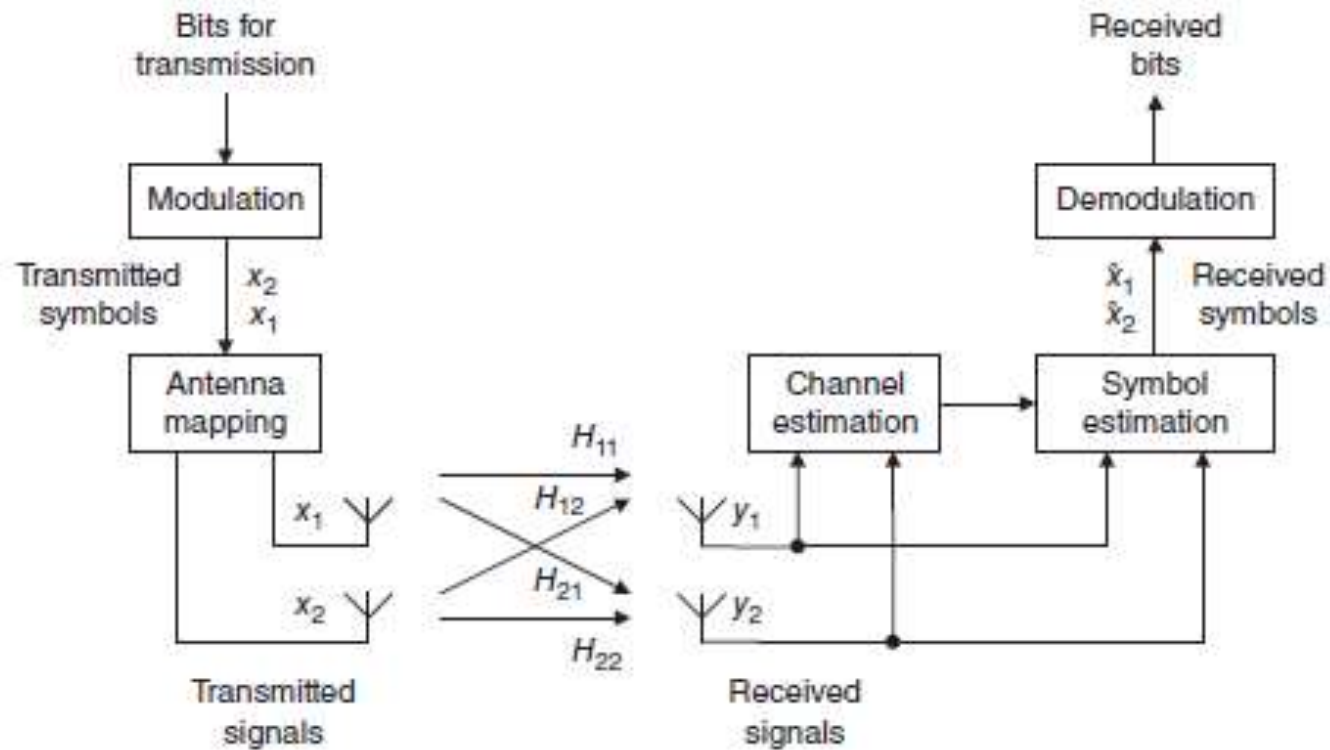


Figure 5.4 Basic principles of a 2×2 spatial multiplexing system

- The symbols travel to the receive antennas by way of four separate radio paths, so the received signals can be written as follows:
- $y1 = H11x1 + H12 x2 + n1$
- $y2 = H21x1 + H22 x2 + n2$
- Here, $x1$ and $x2$ are the signals sent from the two transmit antennas, $y1$ and $y2$ are the signals that arrive at the two receive antennas, and $n1$ and $n2$ represent the received noise and interference.
- H_{ij} expresses the way in which the transmitted symbols are attenuated and phase-shifted, as they travel to receive antenna i from transmit antenna j . (The subscripts i and j may look the wrong way round, but this is for consistency with the usual mathematical notation for matrices.)

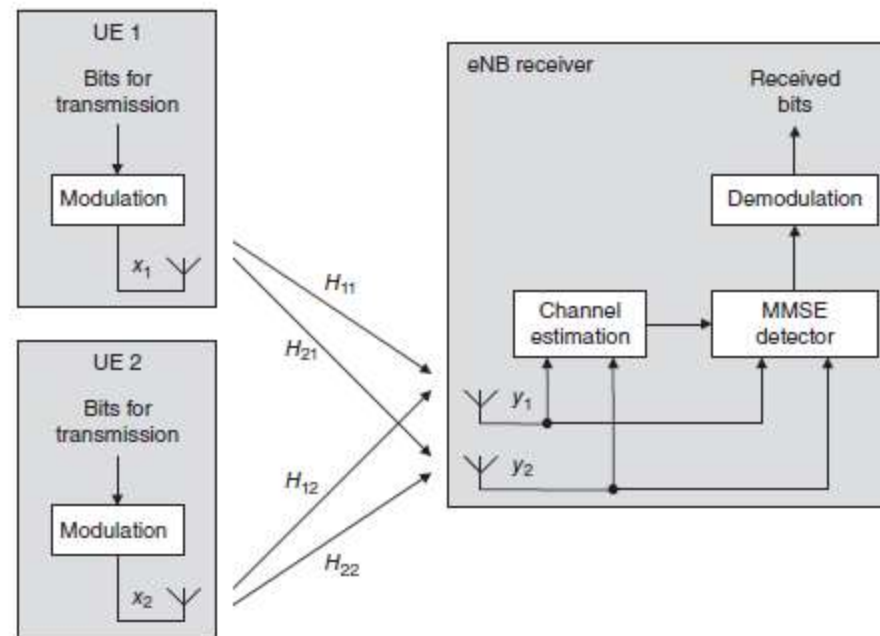
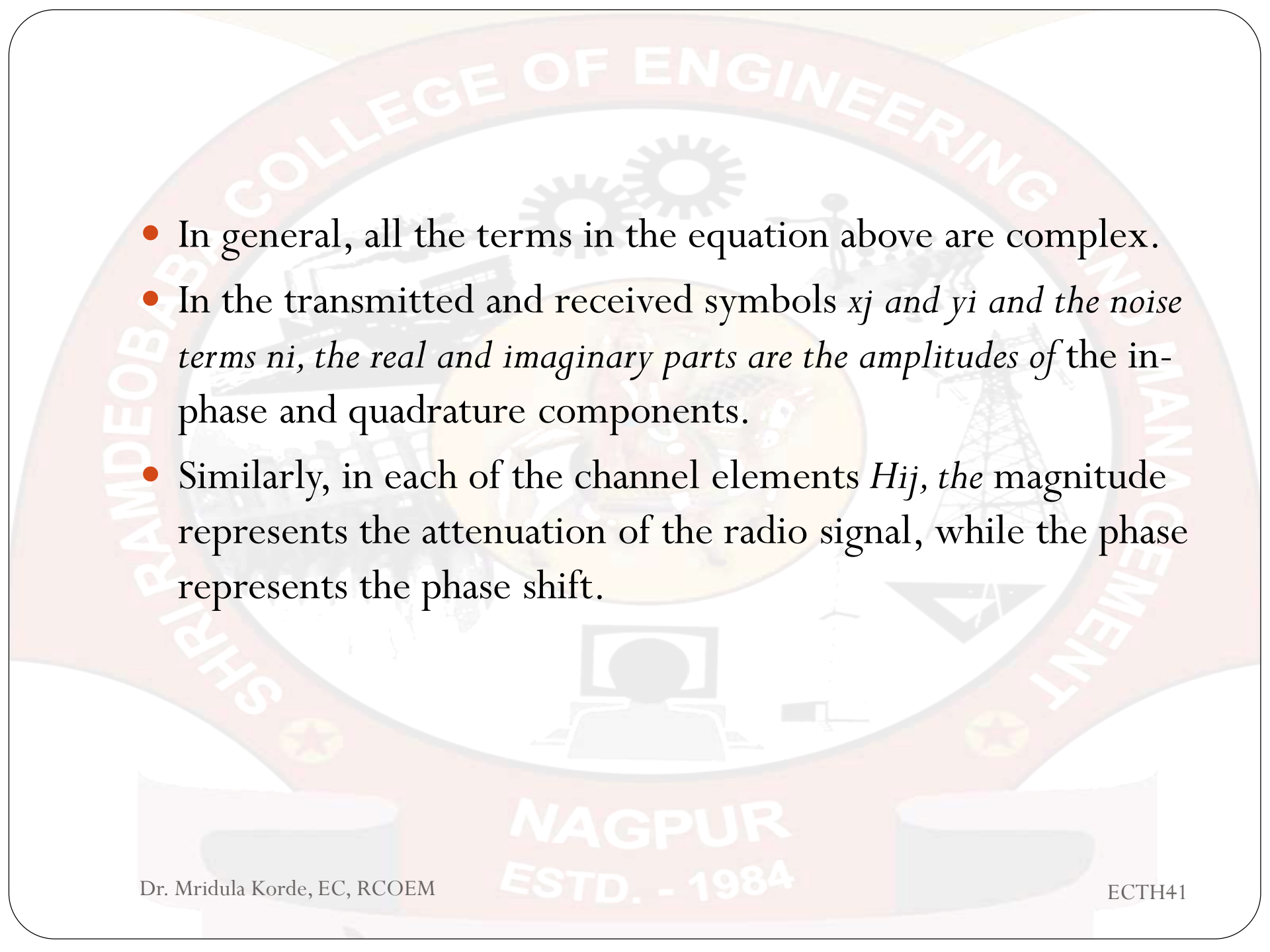
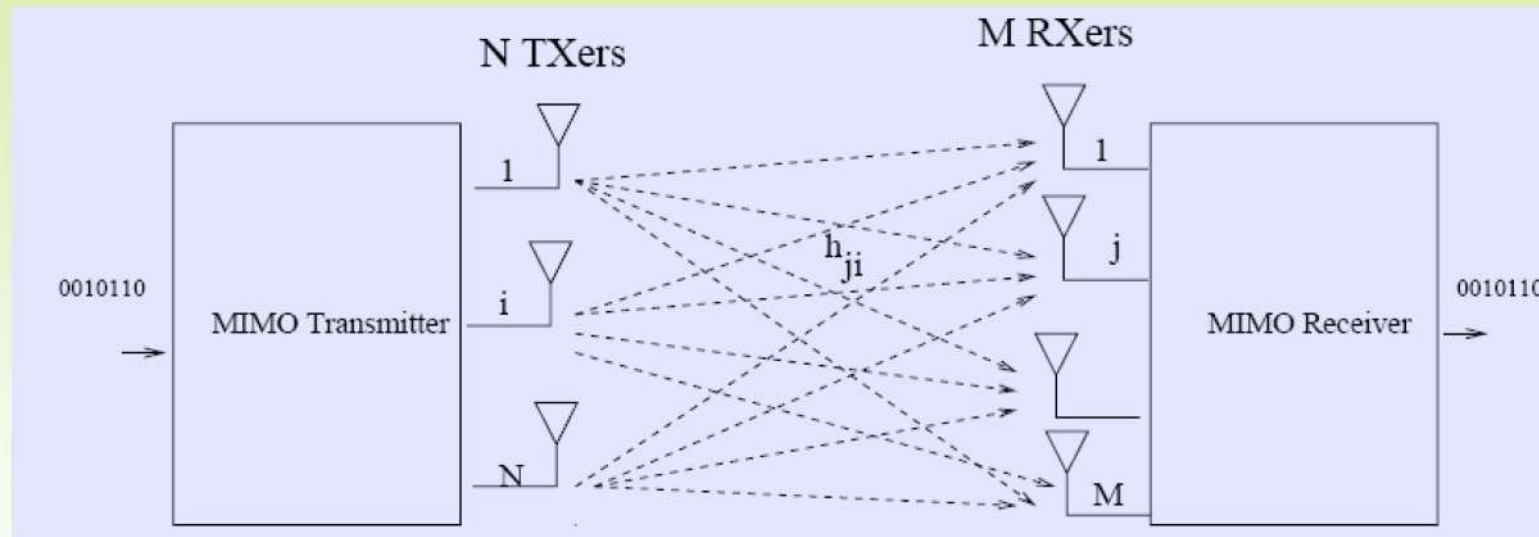


Figure 5.8 Uplink multiple user MIMO

- 
- In general, all the terms in the equation above are complex.
 - In the transmitted and received symbols x_j and y_i and the noise terms n_i , the real and imaginary parts are the amplitudes of the in-phase and quadrature components.
 - Similarly, in each of the channel elements H_{ij} , the magnitude represents the attenuation of the radio signal, while the phase represents the phase shift.

MIMO Channel Matrix

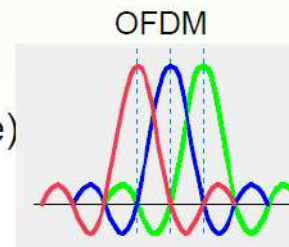


Example for 3 X 4 system:

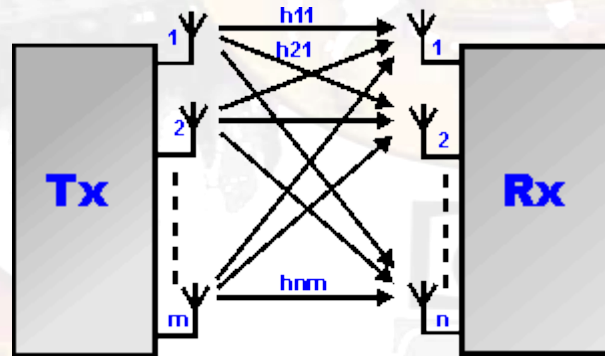
Number of spatial streams equals $\text{rank}(H) \leq \min(M, N)$

$$H = \begin{bmatrix} h_{11} & h_{12} & h_{13} \\ h_{21} & h_{22} & h_{23} \\ h_{31} & h_{32} & h_{33} \\ h_{41} & h_{42} & h_{43} \end{bmatrix}$$

h_{ij} are complex numbers: $a+jb$ (amplitude & phase) and frequency selective



- MIMO is effectively a radio antenna technology as it uses multiple antennas at the transmitter and receiver to enable a variety of signal paths to carry the data, choosing separate paths for each antenna to enable multiple signal paths to be used.



MIMO antenna & MIMO Beamforming Development

For many years antenna technology has been used to improve the performance of systems.

- Directive antennas have been used for very many years to improve signal levels and reduce interference.
- Directive antenna systems have, for example, been used to improve the capacity of cellular telecommunications systems. By splitting a cell site into sector where each antenna illuminates 60° or 120° the capacity can be greatly increased - tripled when using 120° antennas.
- With the development of more adaptive systems and greater levels of processing power, it is possible to utilise antenna beamforming techniques with systems such as MIMO.

MIMO Beamforming Smart Antennas

- Beamforming techniques can be used with any antenna system - not just on MIMO systems.
- They are used to create a certain required antenna directive pattern to give the required performance under the given conditions.
- **Smart antennas** are normally used - these are antennas that can be controlled automatically according the required performance and the prevailing conditions.

Smart antennas can be divided into two groups:

- ***Phased array systems:*** Phased array systems are switched and have a number of pre-defined patterns - the required one being switched according to the direction required.
- ***Adaptive array systems (AAS):*** This type of antenna uses what is termed adaptive beamforming and it has an infinite number of patterns and can be adjusted to the requirements in real time.

- MIMO beamforming using **phased array systems** requires the overall system to determine the direction of arrival of the incoming signal and then switch in the most appropriate beam. This is something of a compromise because the fixed beam is unlikely to exactly match the required direction.
- **Adaptive array systems** are able to direct the beam in the exact direction needed, and also move the beam in real time - this is a particular advantage for moving systems - a factor that often happens with mobile telecommunications. However the cost is the considerable extra complexity required.

Spatial multiplexing: Signal processing

- Example 2x2 MIMO spatial multiplexing

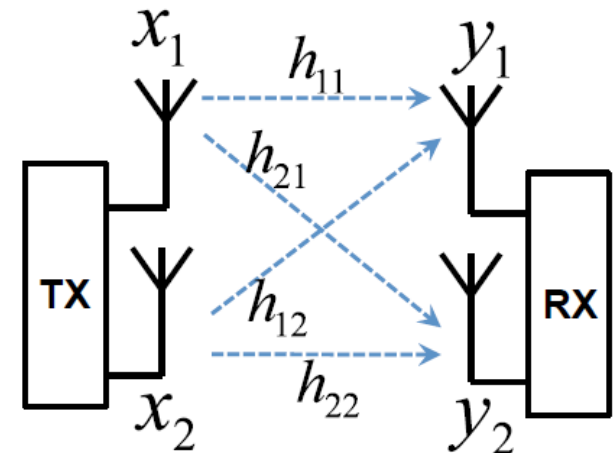
- Data to be sent over two TX antennas:

$$x_1, x_2$$

- Data received on two RX antennas:

$$y_1 = h_{11}x_1 + h_{12}x_2$$

$$y_2 = h_{21}x_1 + h_{22}x_2$$

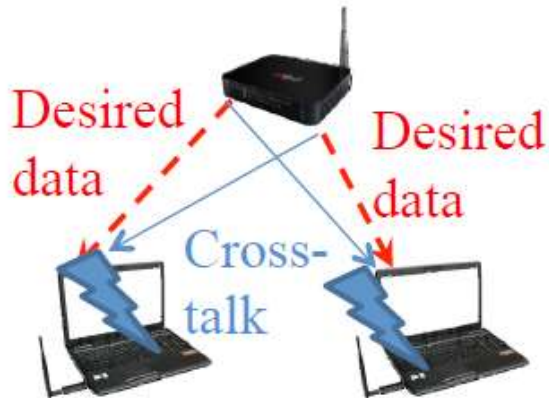


- * Channel distortions: h_{**} can be estimated by the receiver
- * Only two unknowns: x_1, x_2 , easily obtained by solving the equations!

Multi-user MIMO

- Concept of Multi-User MIMO (MU-MIMO)

- * Single-antenna network



- * Multi-user MIMO



- * MU-MIMO enables multiple streams of data to be sent to different users in parallel, without cross-talk interference

MU-MIMO Basics

- MU-MIMO provides a methodology whereby spatial sharing of channels can be achieved. This can be achieved at the cost of additional hardware - filters and antennas - but the incorporation does not come at the expense of additional bandwidth as is the case when technologies such as FDMA, TDMA or CDMA are used.
- When using spatial multiplexing, MU-MIMO, the interference between the different users on the same channel is accommodated by the use of additional antennas, and additional processing when enable the spatial separation of the different users.

Multi-user MIMO or MU-MIMO

- Multi-user MIMO or MU-MIMO is an enhanced form of MIMO technology that is gaining acceptance.
- MU-MIMO, Multi-user MIMO enables multiple independent radio terminals to access a system enhancing the communication capabilities of each individual terminal.
- Accordingly it is often considered as an extension of Space Division Multiple Access, SDMA.
- MU-MIMO exploits the maximum system capacity by scheduling multiple users to be able to simultaneously access the same channel using the spatial degrees of freedom offered by MIMO.

MU-MIMO vs SU-MIMO

- Both Single User-MIMO and Multi-User MIMO systems can be used within wireless and cellular telecommunications systems.
- Each form of MIMO has its advantages and disadvantages.

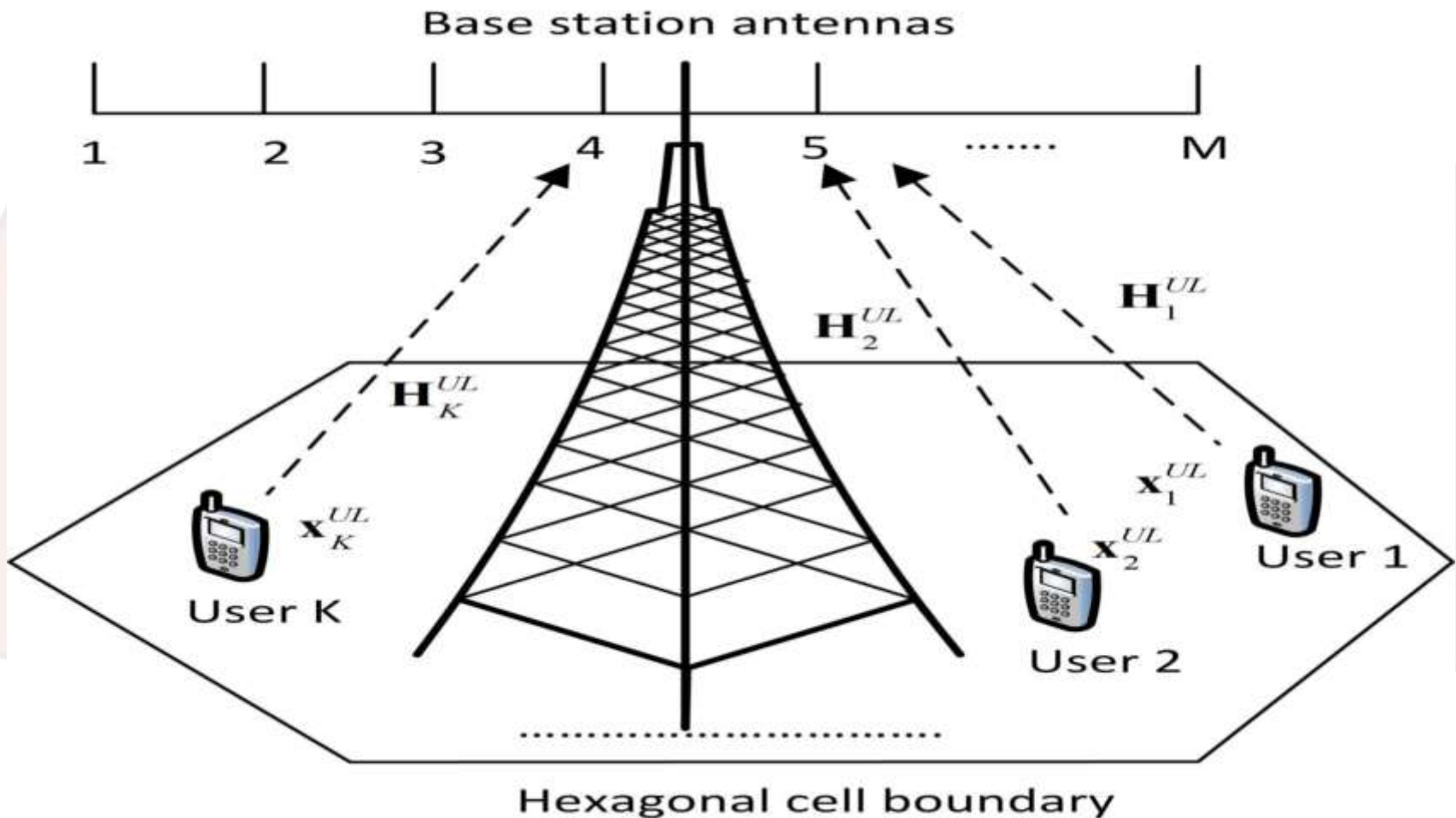
Comparison of Mu-MIMO vs SU-MIMO

Feature	MU-MIMO	SU-MIMO
Main feature	For Mu-MIMO the base station is able to separately communicate with multiple users.	Base station communicates with a single user.
Key aspect	Using MU-MIMO provides capacity gain.	Provides increased data rate for the single user.
Key advantage	Multiplexing gain.	Interference reduction
Data throughput	MU-MIMO provides a higher throughput when the signal to noise ratio is high.	Provides a higher throughput for a low signal to noise ratio.
Channel State Information	Perfect CSI is required.	No CSI needed.

Multiuser MIMO uplink

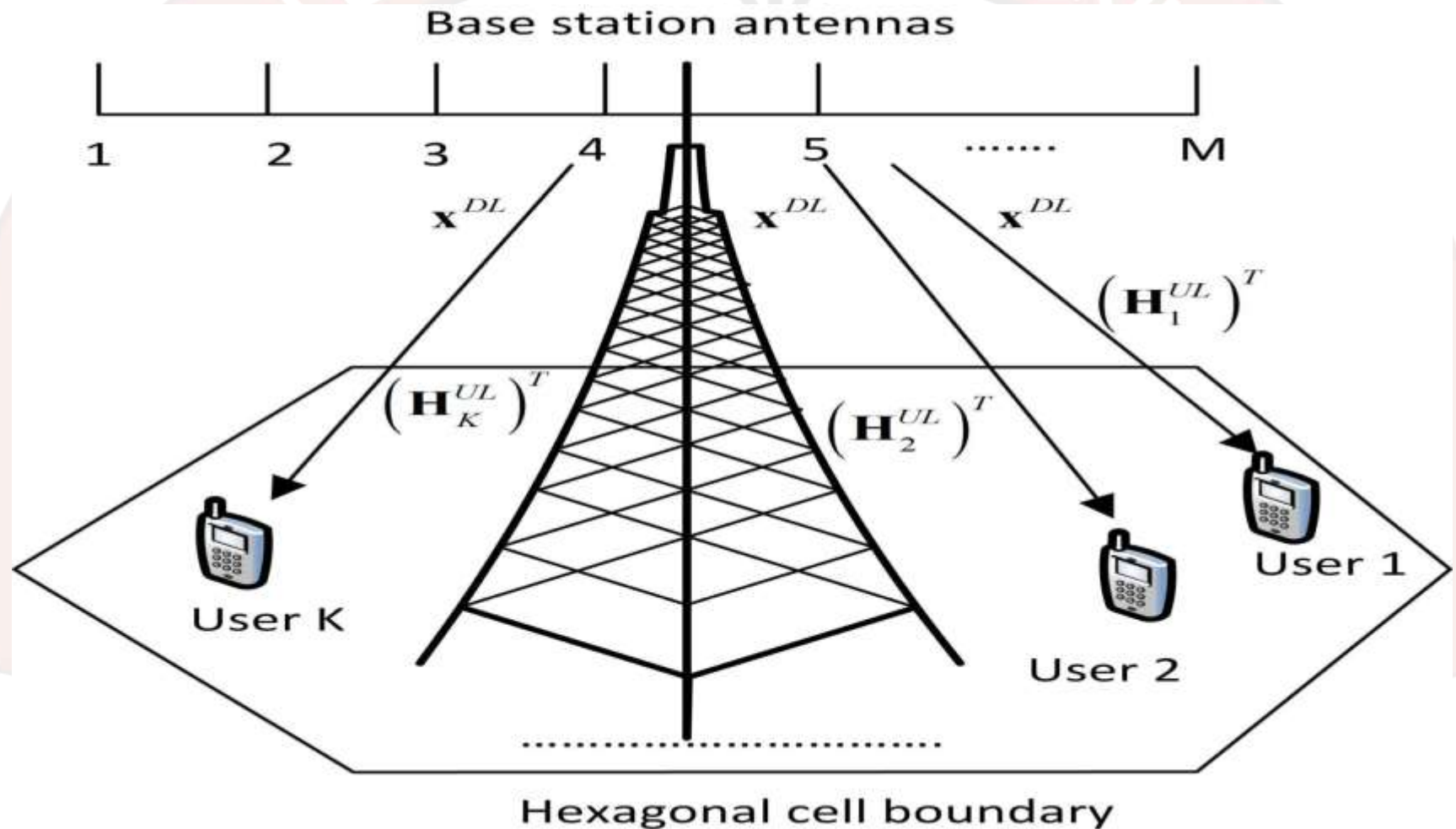
- Multiple access for K mobile users
- K users transmit signal to base station
- Each user may have single or multiple antenna
- Single antenna: one symbol per transmission
- Multiple antenna: Symbol vector per transmission

Multuser MIMO uplink



MU-MIMO system with transmitter employing M antennas serving K users with single antenna (uplink)

Multiuser MIMO downlink

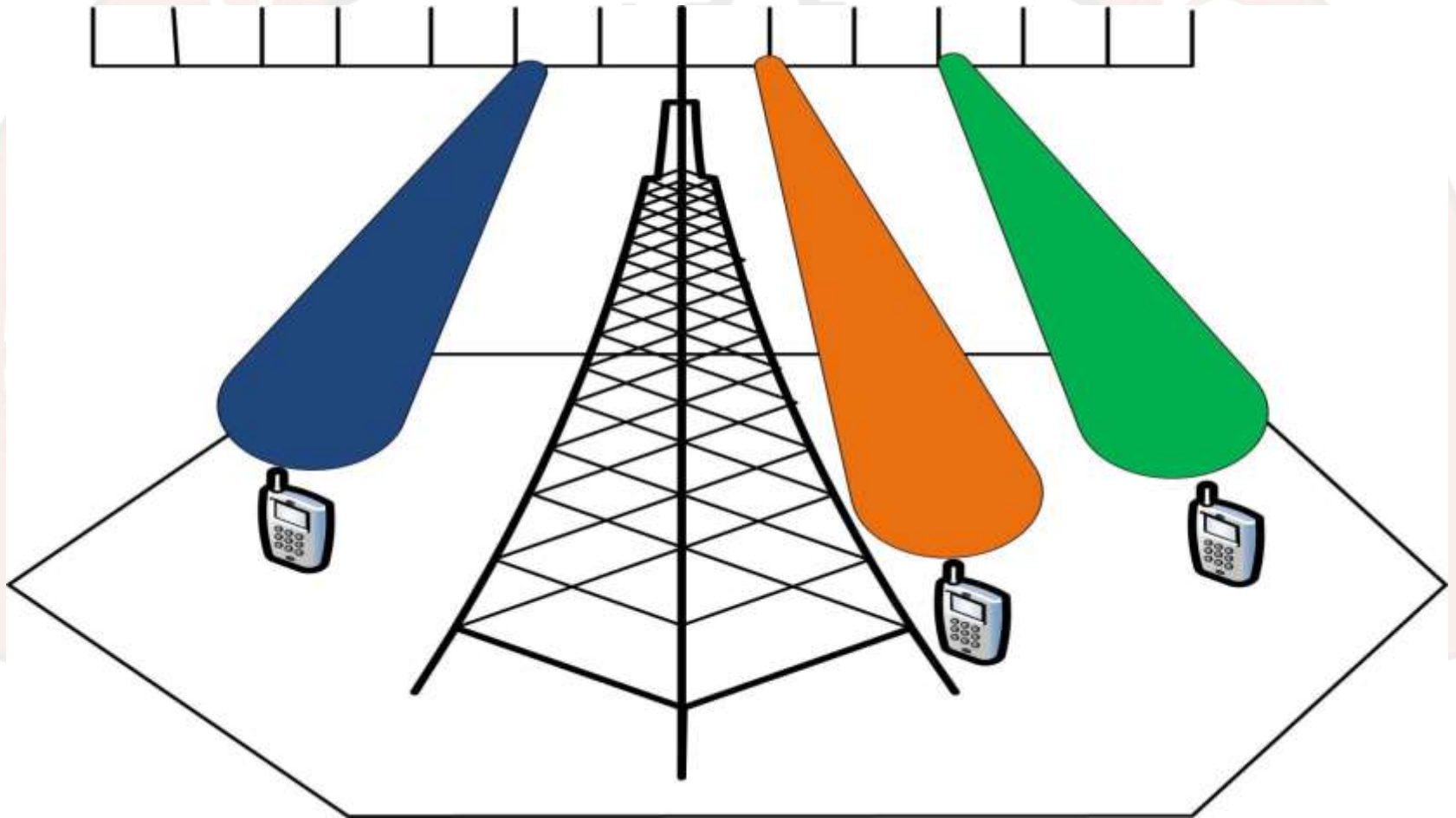


MU-MIMO system with transmitter employing M antennas serving K users with single antenna (downlink)

Massive MIMO, Large MIMO Systems

- **Massive MIMO or large MIMO systems technology is being developed for use in many wireless links where to provide additional data capacity or signal enhancement.**
- Large MIMO systems, often referred to as massive MIMO systems, can be defined as those that use tens or hundreds of antennas in the communication terminals.
- Traditional MIMO systems may have two or four, some may even have eight antennas, but this has been the limit on early systems that have adopted MIMO.
- The concept of massive MIMO or large MIMO systems is entering many areas of development as it is able to offer some distinct advantages.

Massive MIMO



MU-MIMO system for a single base station employing $M=14$ antennas serving $K=3$ users with single antenna (downlink transmission)

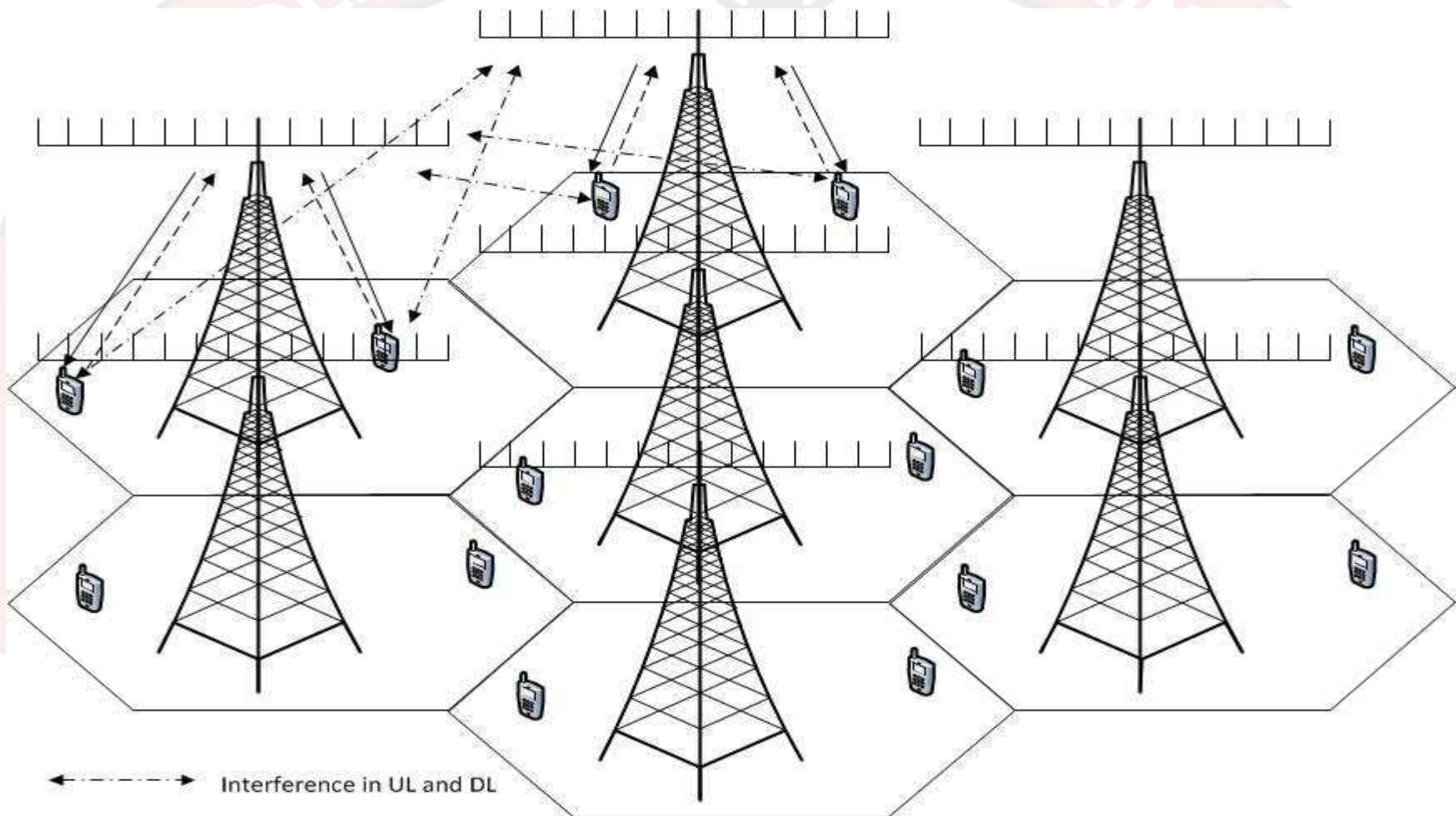
Massive MIMO Benefits

- There are many advantages to using large MIMO technology. Using more antennas in a MIMO system creates more degrees of freedom in the spatial domain and therefore this enables greater improvement in performance to be achieved:
- ***Increasing data rate:*** The increase in the number of antennas allows for a greater number of paths to be used and hence a much greater level of data to be transferred within a given time.
- ***Increasing basic link signal to noise ratio:*** One of the basic advantages of the use of MIMO systems is that it can be used to improve the signal to noise ratio of the overall system. The use of large MIMO or massive MIMO enables this to be taken to a greater level. There is also an increase in hardening against intentional jamming as a result of the large diversity.

- ***Use of spatial modulation:*** The number of RF chains needed for a massive MIMO system can be reduced without compromising the spectral efficiency by using spatial modulation. Spatial modulation is a form of modulation that only requires the use of one transit chain for multiple antennas. Effectively it uses one antenna from an array at a time for transmission.

Spatial modulation adopts a simple but effective coding mechanism which establishes a one to one mapping between blocks of transmitted information bits and the spatial positions of the transmitter antenna in the overall antenna array.

Massive MIMO: Multicell network



Multi-cell MIMO based cellular network (BS equipped with $M=14$ antennas and single antenna MS or user, each cell has $K=2$ users for illustration purpose)

mmWave Massive MIMO

- To be implemented at mmWave frequency region
- Shorter wavelength – smaller antenna size – allows large number of antennas at single terminal
- One of the technology candidate for 5G communication

mmWave Massive MIMO

5G mobile technology requirements and comparison with 4G

Parameter	Unit	5G	4G
Area traffic capacity	Mbps/m ²	10	0.1
Peak data rate	Gbps	20	1
User experienced data rate	Mbps	100	10
Spectrum efficiency		3X	1X
Energy efficiency		100X	1X
Connection density	devices/km ²	10 ⁶	10 ⁵
Latency	ms	1	10
Mobility	km/h	500	350

Conclusion

- Multiple-input multiple-output, or MIMO, is a radio communications technology or RF technology that is being mentioned and used in many new technologies these days.
- Wi-Fi, LTE; Long Term Evolution, and many other radio, wireless and RF technologies are using the new MIMO wireless technology to provide increased link capacity and spectral efficiency combined with improved link reliability using what were previously seen as interference paths.



**SHRI RAMDEOBABA COLLEGE OF ENGINEERING AND
MANAGEMENT,
KATOL ROAD, NAGPUR, MAHARASHTRA, INDIA**

Communication System Analysis(ECTH 41)

Dr. (Mrs.) Mridula Korde

Associate Professor, Department of Electronics and
Communication Engineering, RCOEM, Nagpur

Congestion Control in Asynchronous Transfer Mode(ATM)

What is Congestion?

- Congestion in a network is a state in which performance degrades due to the saturation of network resources such as communication links, processor cycles, and memory buffers.
- Network congestion has been well recognized as a *resource-sharing problem*.
- *In a packet switched* network, resources are shared among all the hosts attached to it, including switch processors, communication channels, and buffer spaces.
- These three driving forces of data transmission in network communication can also be potential bottlenecks that cause congestion in the network

What Is Congestion?

- Congestion occurs when the number of packets being transmitted through the network approaches the packet handling capacity of the network
- Congestion control aims to keep number of packets below level at which performance falls off dramatically
- Data network is a network of queues
- Generally 80% utilization is critical
- Finite queues mean data may be lost

- Networks need to serve all user requests for data transmission, which are often unpredictable and bursty with regard to transmission starting time, rate, and size.
- On the other hand, any physical resource in the network has a finite capacity, and must be managed for sharing among different transmissions.
- Consequently, network congestion will result if the resources in the network cannot meet all of the users' current demands.
- In simple terms, if, for any time interval, the total sum of demands on a resource is more than its available capacity, the resource is said to be congested for that interval.
- **Mathematically speaking: Demand > Available Resources**

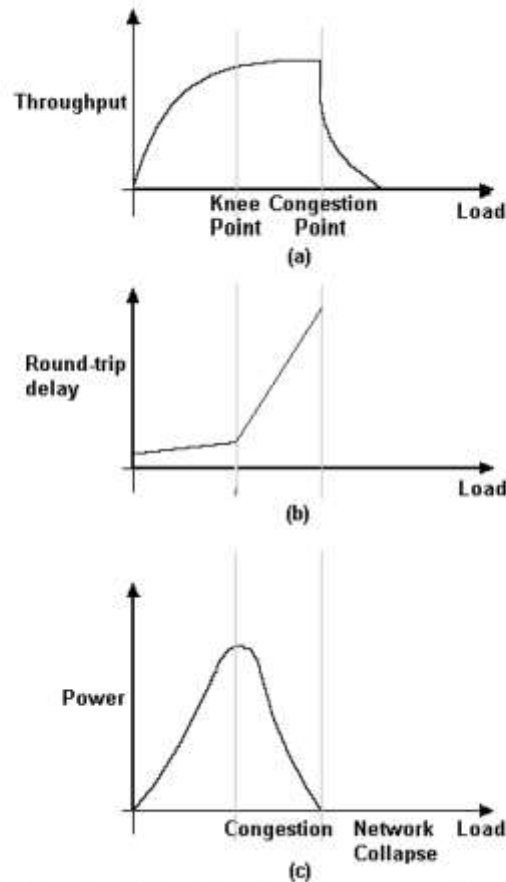


Figure 1.1: Network performance vs. offered traffic load.

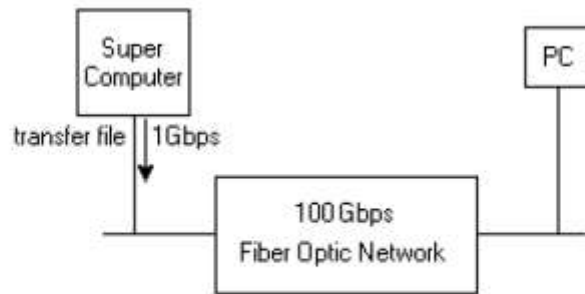
A more formal and quantitative definition for network congestion is based on the performance behaviour of a network. Figure 1.1(a) shows the throughput-load relationship in a packet-switching network

As the load is small and within the subnet carrying capacity, network throughput generally keeps up with the increase of the load until the offered load reaches to the knee point, where the increase of the throughput becomes much slower than the increase of the load.

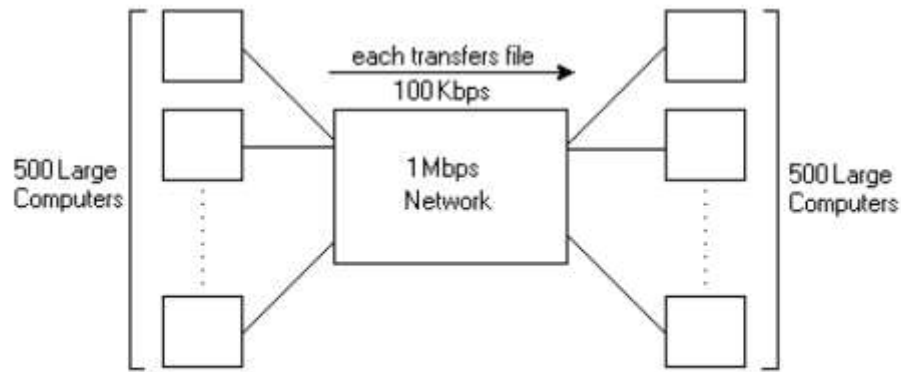
If the load keeps increasing up to the capacity of the network, the queues on switching nodes will build up, potentially resulting in packets being dropped, and throughput will eventually arrive at its maximum and then decrease sharply to a low value (possibly zero).

It is at this point that the network is said to be congested.

Figures 1.1(b) and 1.1(c) illustrate the relationships between the round-trip delay, and the resource power with respect to the offered load.



(a) Flow Control Problem



(b) Congestion Control Problem

Flow control, in contrast, relates to the point-to-point traffic between a given sender and a given receiver.

Its job is to make sure that a fast sender cannot continually transmit data faster than the receiver can absorb it. Flow control nearly always involves some direct feedback from the receiver to the sender to tell the sender how things are doing at the other end.

The reason congestion control and flow control are often confused is that some congestion control algorithms operate by sending messages back to the various sources telling them to slow down when the network gets into trouble.

Thus a host can get a “slow down” message either because the receiver cannot handle the load, or because the network cannot handle it.

Myths about Congestion Control

- Congestion occurs when the demand is greater than the available resources.
- Therefore, it is believed that as resources become less expensive, the problem of congestion will be solved automatically. This has led to the following myths:
 1. *Congestion is caused by a shortage of buffer space and will be solved when memory becomes cheap enough to allow infinitely large memories.*
 2. *Congestion is caused by slow links. The problem will be solved when high-speed links become available.*

- The congestion problem can not be solved with a large buffer space
- With infinite-memory switches, as shown in Figure 1.3, the queues and the delays can get so long that by the time the packets come out of the switch, most of them have already timed out, dropped, and have been retransmitted by higher layers.

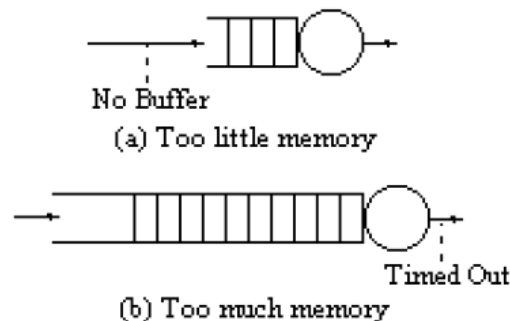


Figure 1.3: Too much memory in the intermediate nodes is as harmful as too little memory

- The congestion problem can not be solved with high-speed links
- In Figure 1.4, with the high-speed link, the arrival rate to the first router became much higher than the departure rate, leading to long queues, buffer overflows, and packet losses that cause the transfer time to increase.
- The point is that high-speed links cannot stay in isolation. The low-speed links do not go away as the high-speed links are added to a network. The protocols have to be designed specifically to ensure that this increasing range of link speeds does not degrade the performance.

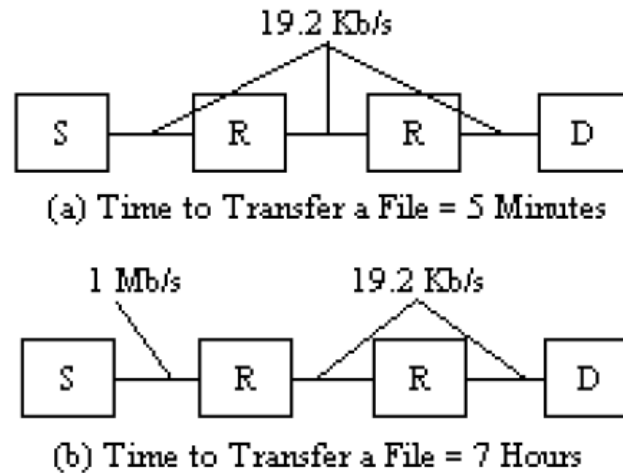


Figure 1.4: Introducing a high-speed link may reduce the performance.

- **Congestion occurs even if all links and processors are of the same speed** Our arguments above may lead some to believe that a balanced configuration with all processors and links at the same speed will probably not be susceptible to congestion.

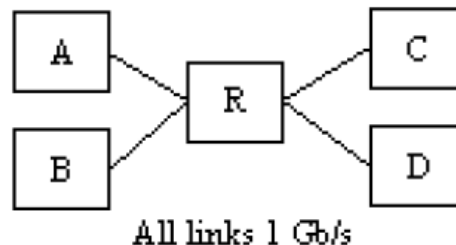
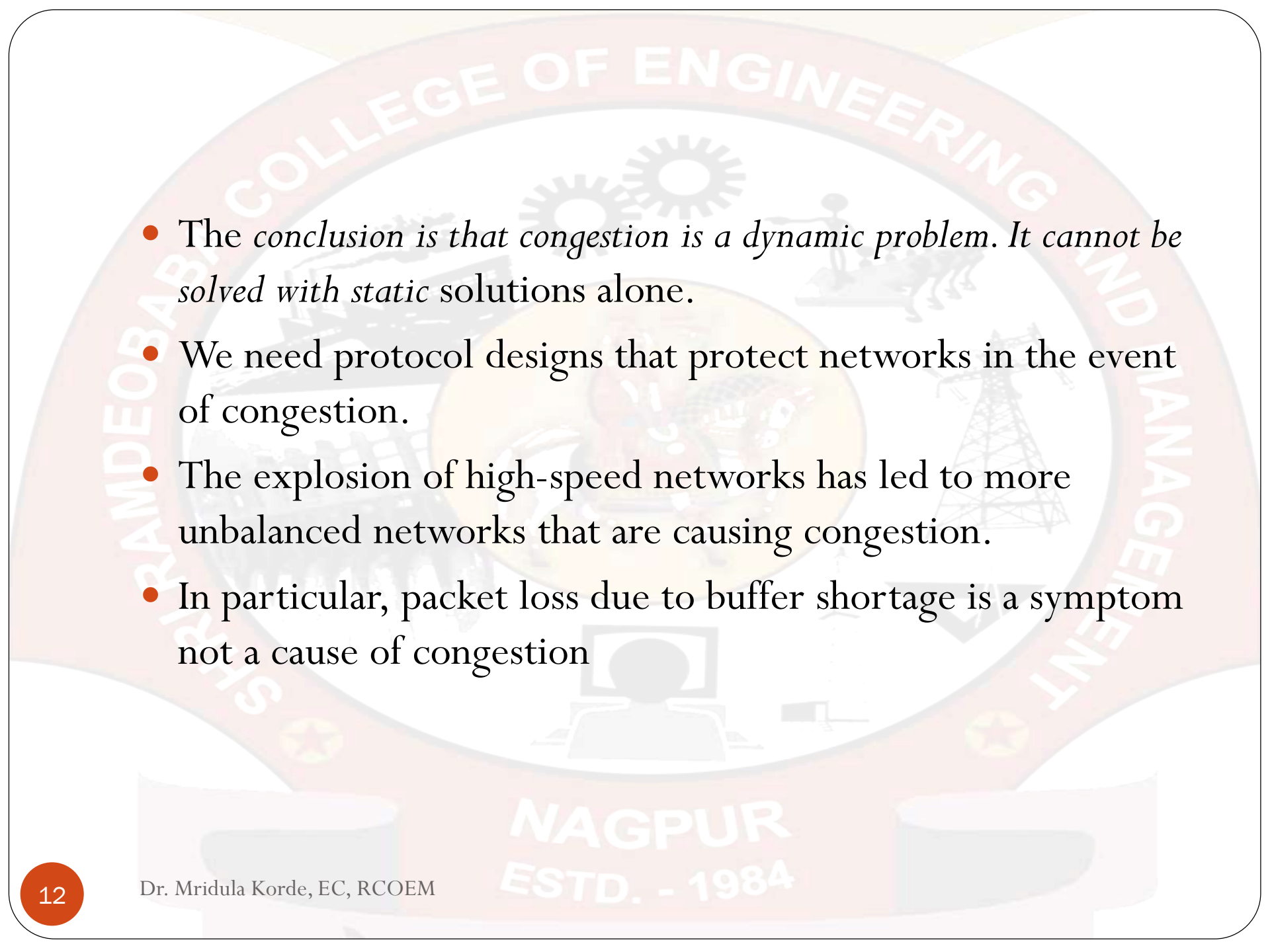


Figure 1.5: A balanced configuration with all processors and links at the same speed is also susceptible to congestion.

- 
- The *conclusion is that congestion is a dynamic problem. It cannot be solved with static solutions alone.*
 - We need protocol designs that protect networks in the event of congestion.
 - The explosion of high-speed networks has led to more unbalanced networks that are causing congestion.
 - In particular, packet loss due to buffer shortage is a symptom not a cause of congestion

TRAFFIC MANAGEMENT

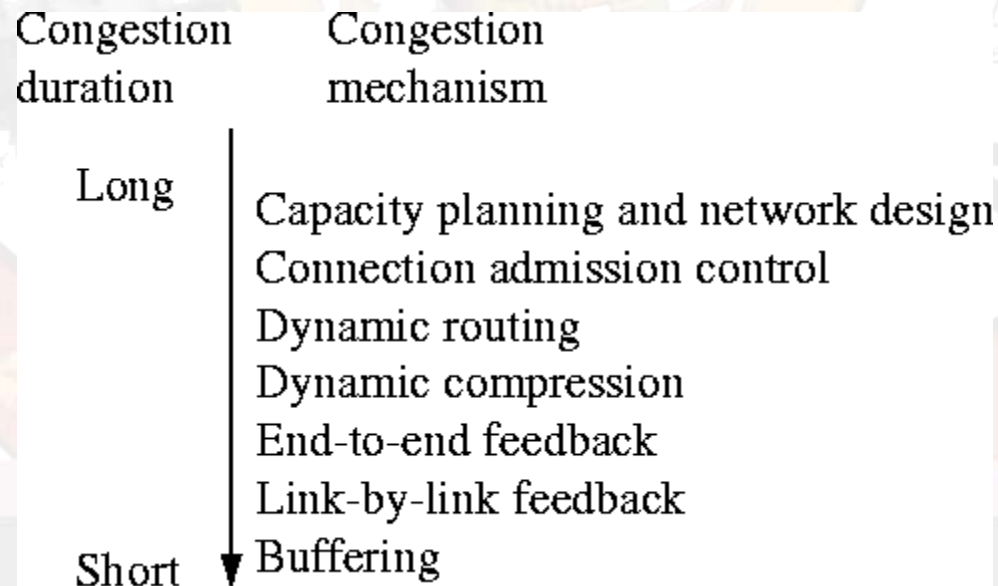
Role of Traffic Management

- ATM technology is intended to support a wide variety of services and applications.
- Proper traffic management helps ensure efficient and fair operation of networks in spite of constantly varying demand and ensure that users get their desired quality of service.
- One of the challenges in designing ATM traffic management was to maintain the QoS for various classes while attempting to make maximal use of network resources.
- Traffic management is required even if the network is underloaded. The problem is especially difficult during periods of heavy load particularly if the traffic demands cannot be predicted in advance.
- This is why congestion control, although only a part of the traffic management issues, is the most essential aspect of traffic management.

Congestion Control Methods

- Congestion happens whenever the input rate is more than the available link capacity:
- **$\text{Sum}(\text{Input Rate}) > \text{Available link capacity}$**
- Most congestion control schemes consist of adjusting the input rates to match the available link capacity (or rate).

- Figure 1 shows how the duration of congestion affects the choice of the method. The best method for networks that are almost always congested is to install higher speed links and redesign the topology to match the demand pattern.
- **Figure 1:** Congestion techniques for various congestion durations



Congestion Schemes

Fast Resource Management

- This requires sources to send a resource management (RM) cell requesting the desired bandwidth before actually sending the cells.
- If a switch cannot grant the request it simply drops the RM cell; the source times out and resends the request.
- If a switch can satisfy the request, it passes the RM cell on to the next switch.
- Finally, the destination returns the cell back to the source which can then transmit the burst.
- Here, the burst has to wait for at least one round trip delay at the source even if the network is idle (as is often the case).

Delay-Based Rate Control

- This requires that the sources monitor the round trip delay by periodically sending resource management (RM) cells that contain timestamp.
- The cells are returned by the destination.
- The source uses the timestamp to measure the roundtrip delay and to deduce the level of congestion.
- This approach, has the advantage that no explicit feedback is expected from the network and, therefore, it will work even if the path contained non-ATM networks.

Backward Explicit Congestion Notification (BECN)

- This method consists of switches monitoring their queue length and sending an RM cell back to source if congested.
- The sources reduce their rates by half on the receipt of the RM cell.
- If no cells are received within a recovery period, the rate for that VC is doubled once each period until it reaches the peak rate.
- This scheme was dropped because it was found to be unfair. The sources receiving cells were not always the ones causing the congestion.

Early Packet Discard

- This method is based on the observation that a packet consists of several cells.
- It is better to drop all cells of one packet than to randomly drop cells belonging to different packets.
- In AAL5, when the first bit of the payload type bit in the cell header is 0, the third bit indicates "end of message (EOM)." When a switch's queues start getting full, it looks for the EOM marker and it drops all future cells of the VC until the "end of message" marker is seen again.
- It was pointed out that the method may not be fair in the sense that the cell to arrive at a full buffer may not belong to the VC causing the congestion.

Credit-Based Approach

- The approach consists of per-link, per-VC, window flow control.
- Each link consists of a sender node (which can be a source end system or a switch) and a receiver node (which can be a switch or a destination end system).
- Each node maintains a separate queue for each VC.
- The receiver monitors queue lengths of each VC and determines the number of cells that the sender can transmit on that VC.
- This number is called ``credit.“
- The sender transmits only as many cells as allowed by the credit.
- If there is only one active VC, the credit must be large enough to allow the whole link to be full at all times.
- **$\text{Credit} \geq \text{Link Cell Rate} \times \text{Link Round Trip Propagation Delay}$**

Summary

- Congestion control is important in high speed networks.
- Due to larger bandwidth-distance product, the amount of data lost due to simultaneous arrivals of bursts from multiple sources can be larger.
- For the success of ATM, it is important that it provides a good traffic management for both bursty and non-bursty sources.



**SHRI RAMDEOBABA COLLEGE OF ENGINEERING AND
MANAGEMENT,
KATOL ROAD, NAGPUR, MAHARASHTRA, INDIA**

Communication System Analysis(ECTH 41)

Dr. (Mrs.) Mridula Korde

Associate Professor, Department of Electronics and
Communication Engineering, RCOEM, Nagpur

Asynchronous Transfer Mode(ATM)

ATM – definition

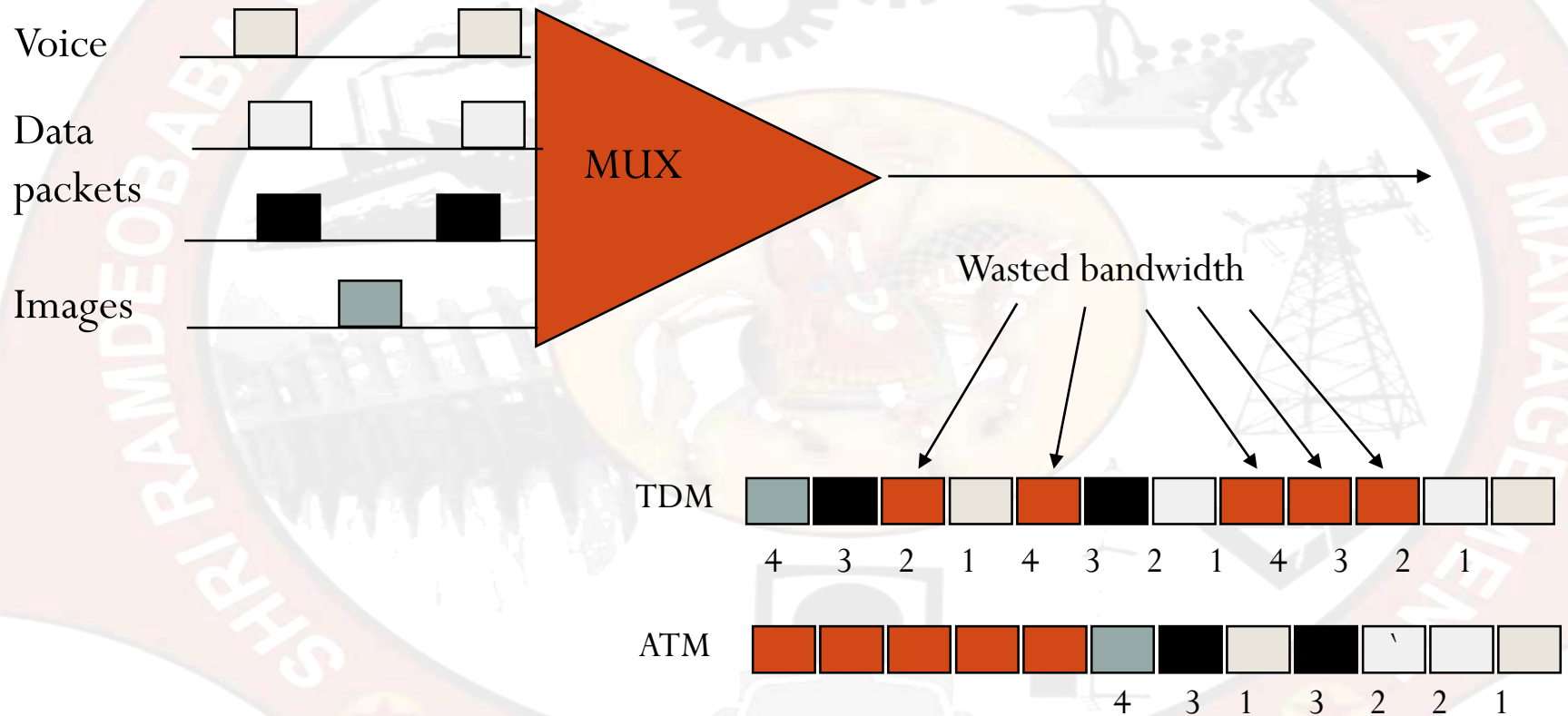
"A **TRANSFER MODE** in which information is organized into cells; it is **ASYNCHRONOUS** in the sense that the recurrence of cells containing information from an individual user is not necessarily periodic".

- Low-level network layer—above physical layer, below AAL (ATM adaptation layer)
- Single transport mechanism for different types of traffic (voice, data, video, etc.)

Synchronous Transfer Mode

- Pre-assigned “slots,” frame boundaries, global timing
- Slots identified by position from the start of the frame
- BW allocated in units of slots
- Idle slots wasted
- Efficient for Constant Bit Rate traffic

Asynchronous Transfer Mode (ATM)



Stallings “High-Speed Networks”

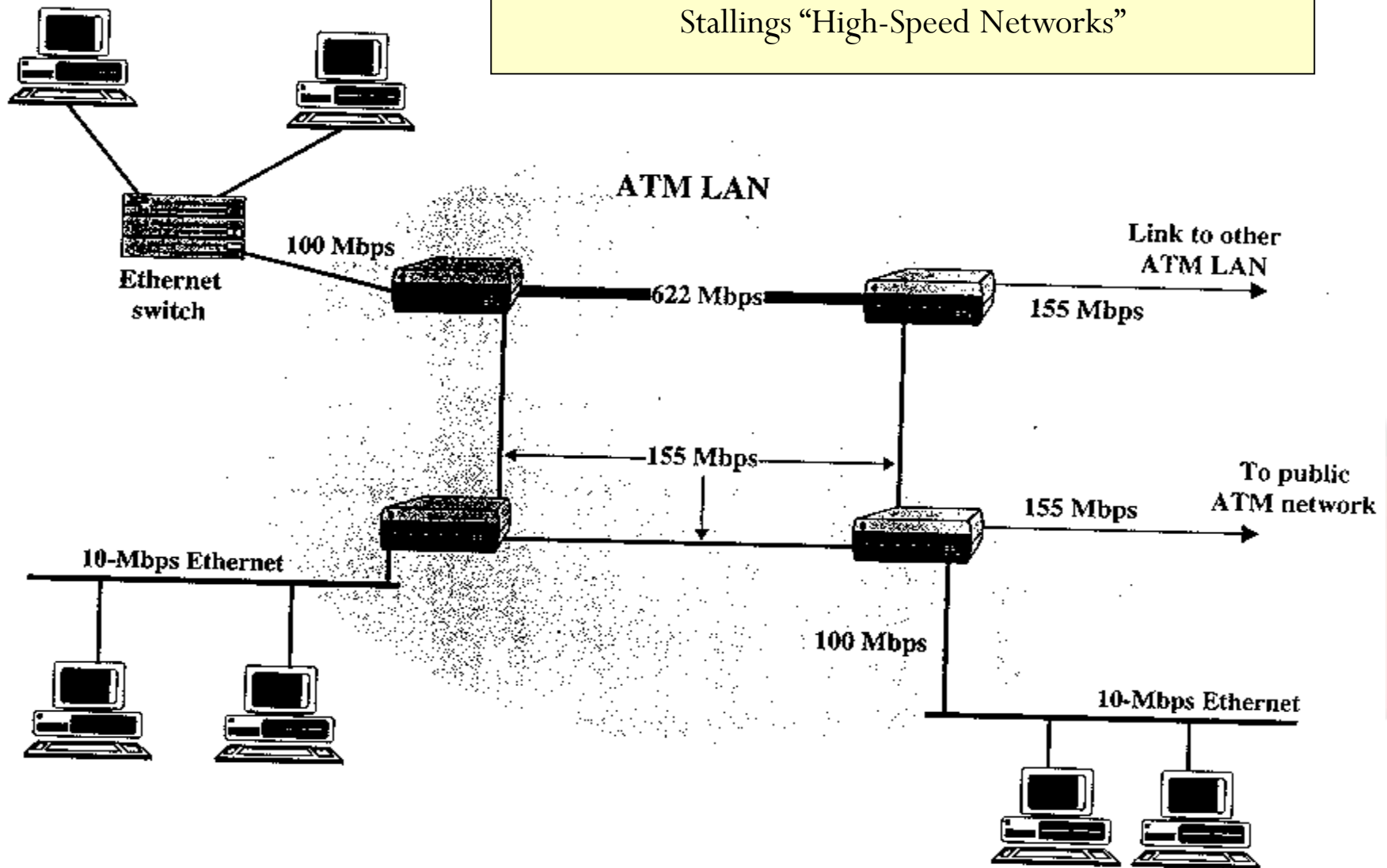


Figure 5.9 Example ATM LAN configuration.

Stallings “High-Speed Networks”

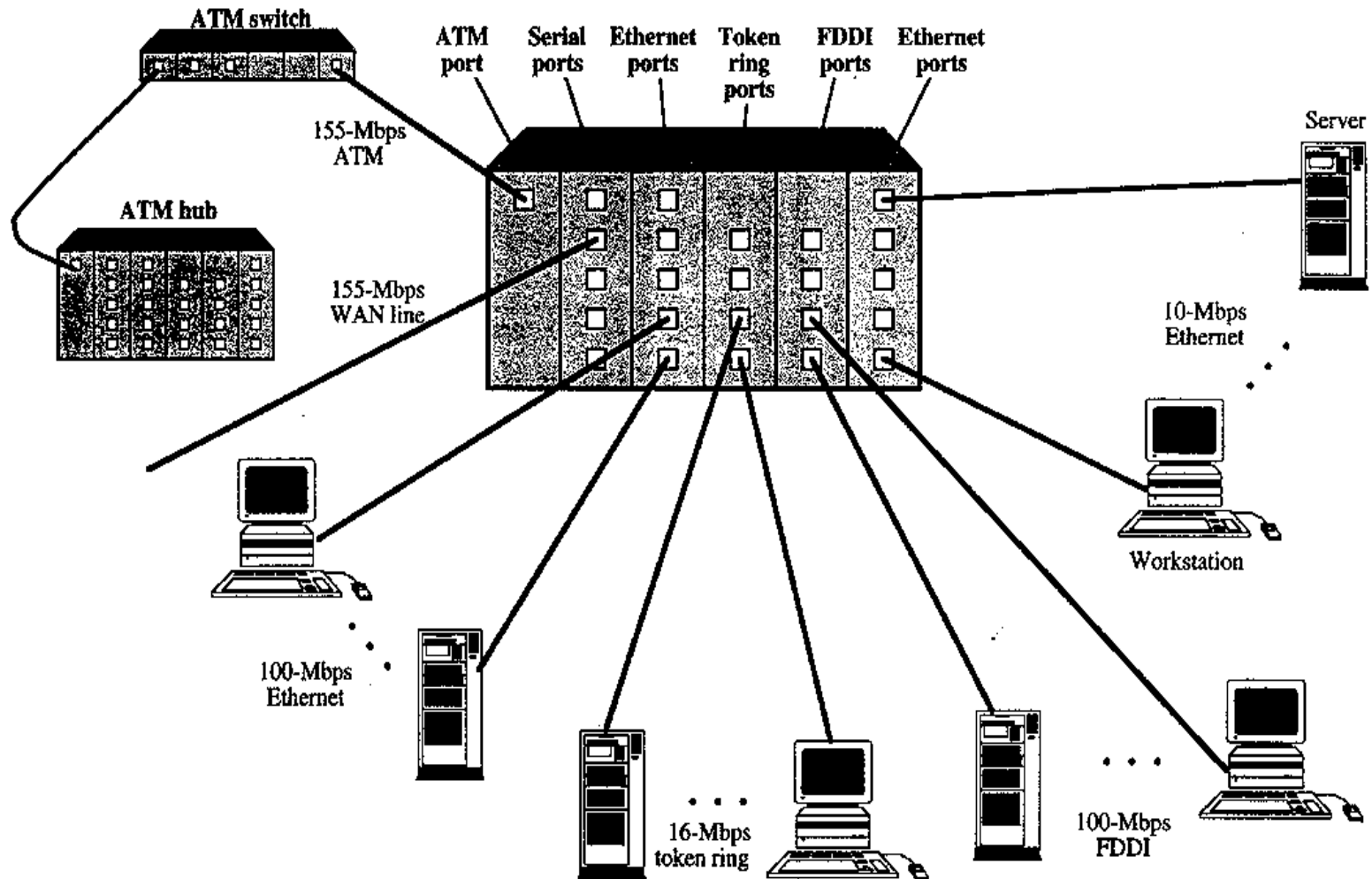
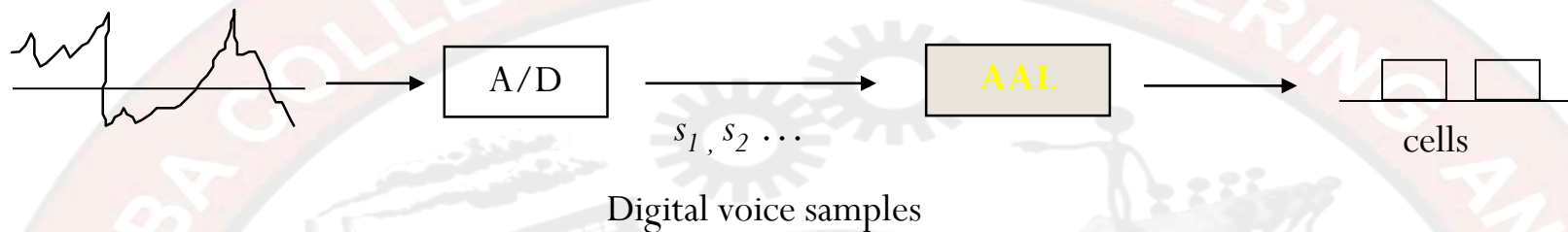


Figure 5.10 ATM LAN hub configuration.

Voice

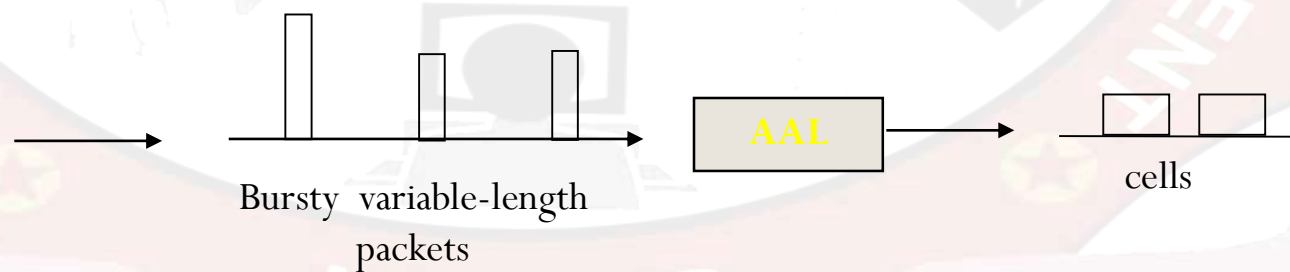
ATM Adaptation Layers



Video



Data



ATM

- ATM standard is widely accepted by common carriers as mode of operation for communication.
- ATM is a form of cell switching using small fixed-sized packets.

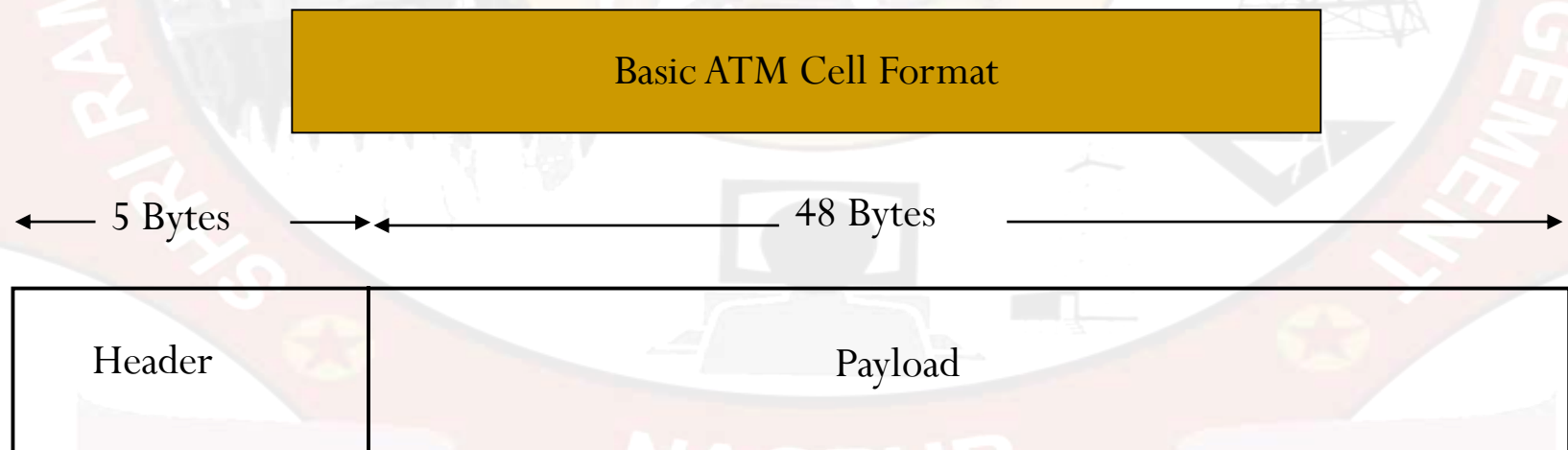
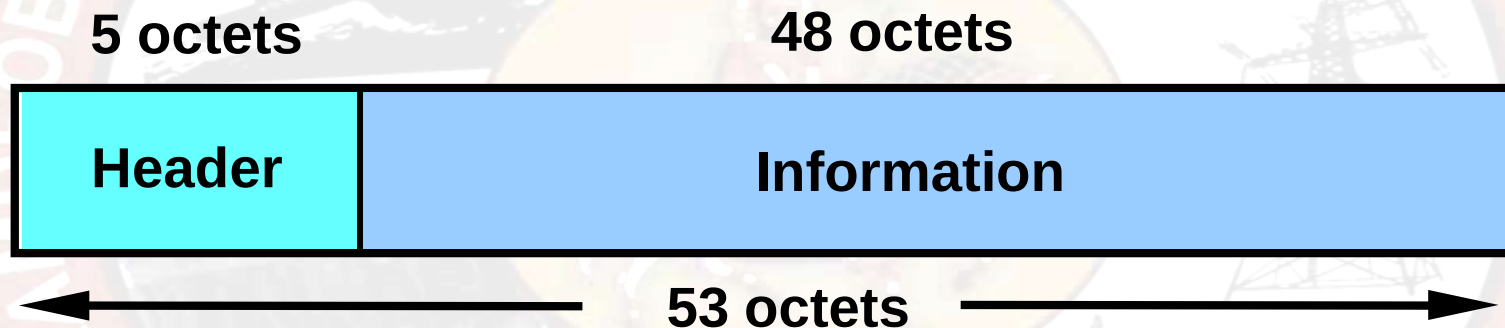


Figure 9.1

What is Asynchronous Transfer Mode (ATM)?

- Asynchronous Transfer Mode (ATM) is a connection-oriented, high-speed, low-delay switching and transmission technology that uses short and fixed-size packets, called cells, to transport information.



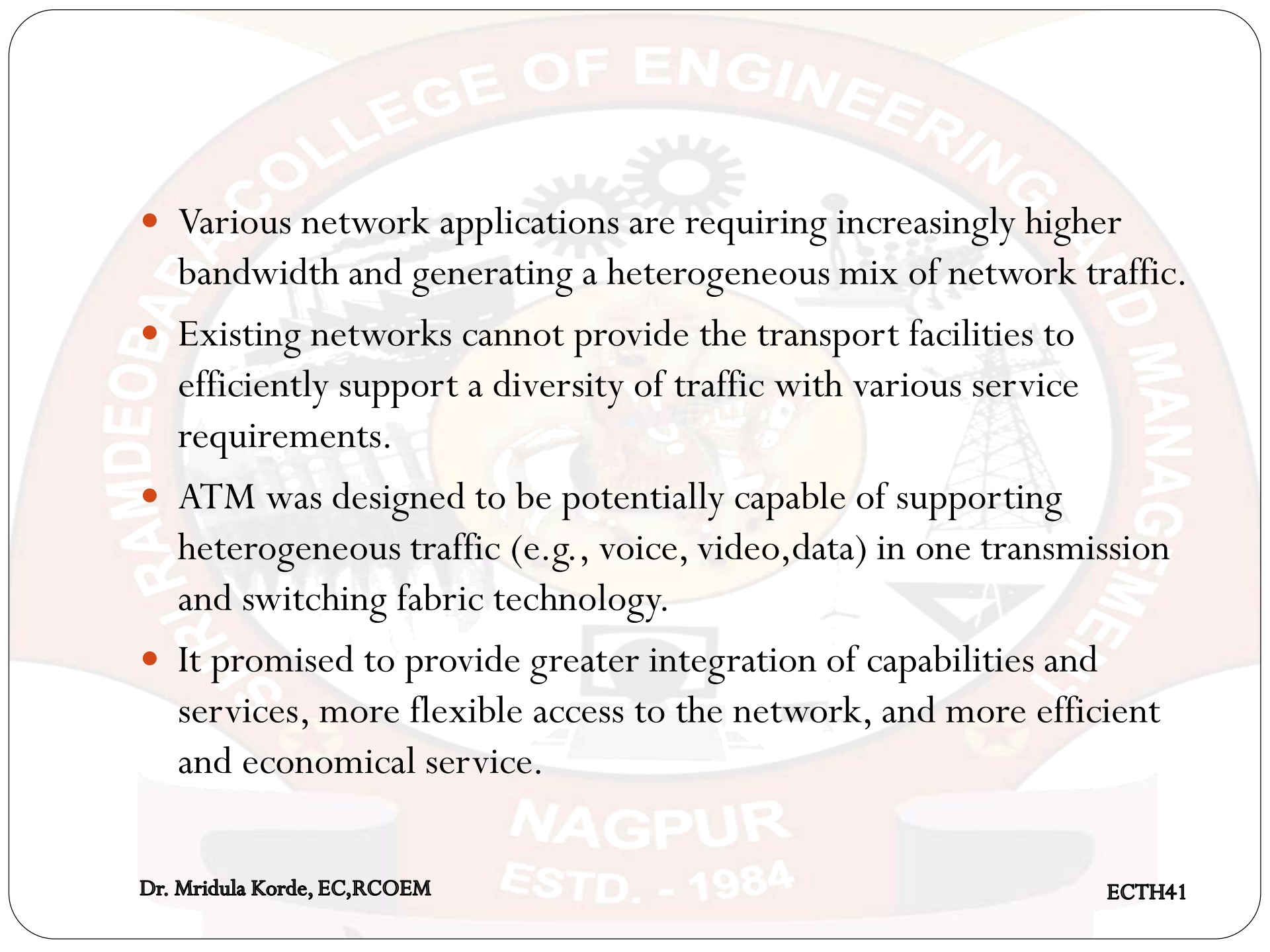
- Using the cell switching technique, ATM combines the benefits of both circuit switching (low and constant delay, guaranteed capacity) and packet switching (flexibility, efficiency for bursty traffic) to support the transmission of multimedia traffic such as voice, video, image, and data over the same network.

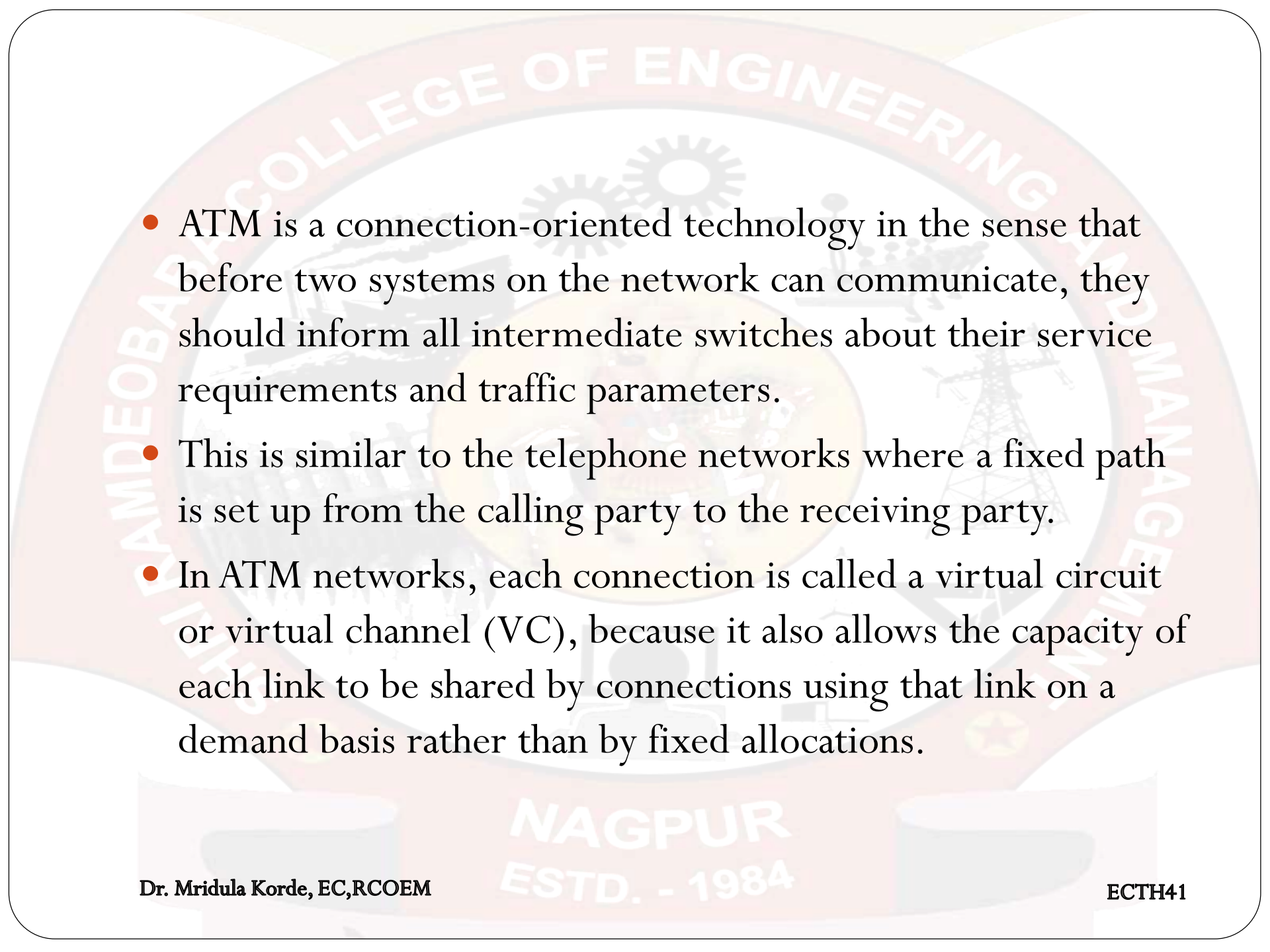
Why ATM?

- International standard-based technology (for interoperability)
- Low network latency (for voice, video, and real-time applications)
- Low variance of delay (for voice and video transmission)
- Guaranteed quality of service
- High capacity switching (multi-giga bits per second)
- Bandwidth flexibility (dynamically assigned to users)
- Scalability (capacity may be increased on demand)
- Medium not shared for ATM LAN (no degradation in performance as traffic load or number of users increases)
- Supports a wide range of user access speeds
- Appropriate (seamless integration) for LANs, MANs, and WANs
- Supports audio, video, imagery, and data traffic (for integrated services)

Introduction

- Asynchronous transfer mode (ATM) is a cell-oriented switching and multiplexing technology that uses fixed-length (53 byte; 48 bytes of data, and 5 bytes of header information) packets — called cells — to carry various types of traffic, such as data, voice, video, multimedia, and so on, through multiple classes of services.
- ATM is a connection-oriented technology, in which a connection is established between the two endpoints before the actual data exchange begins.

- 
- Various network applications are requiring increasingly higher bandwidth and generating a heterogeneous mix of network traffic.
 - Existing networks cannot provide the transport facilities to efficiently support a diversity of traffic with various service requirements.
 - ATM was designed to be potentially capable of supporting heterogeneous traffic (e.g., voice, video, data) in one transmission and switching fabric technology.
 - It promised to provide greater integration of capabilities and services, more flexible access to the network, and more efficient and economical service.

- 
- ATM is a connection-oriented technology in the sense that before two systems on the network can communicate, they should inform all intermediate switches about their service requirements and traffic parameters.
 - This is similar to the telephone networks where a fixed path is set up from the calling party to the receiving party.
 - In ATM networks, each connection is called a virtual circuit or virtual channel (VC), because it also allows the capacity of each link to be shared by connections using that link on a demand basis rather than by fixed allocations.

ATM Protocol Reference Model

- The ATM protocol reference model is based on standards developed by the ITU.
- Communication from higher layers is adapted to the lower ATM defined layers, which in turn pass the information onto the physical layer for transmission over a selected physical medium.
- The protocol reference model is divided into three layers: the ATM adaptation layer (AAL), the ATM layer, and the physical layer, as shown in Figure 1 [4].
- The three management planes user/control plane, layer management and plane management, are shown in Figure 2 [4].

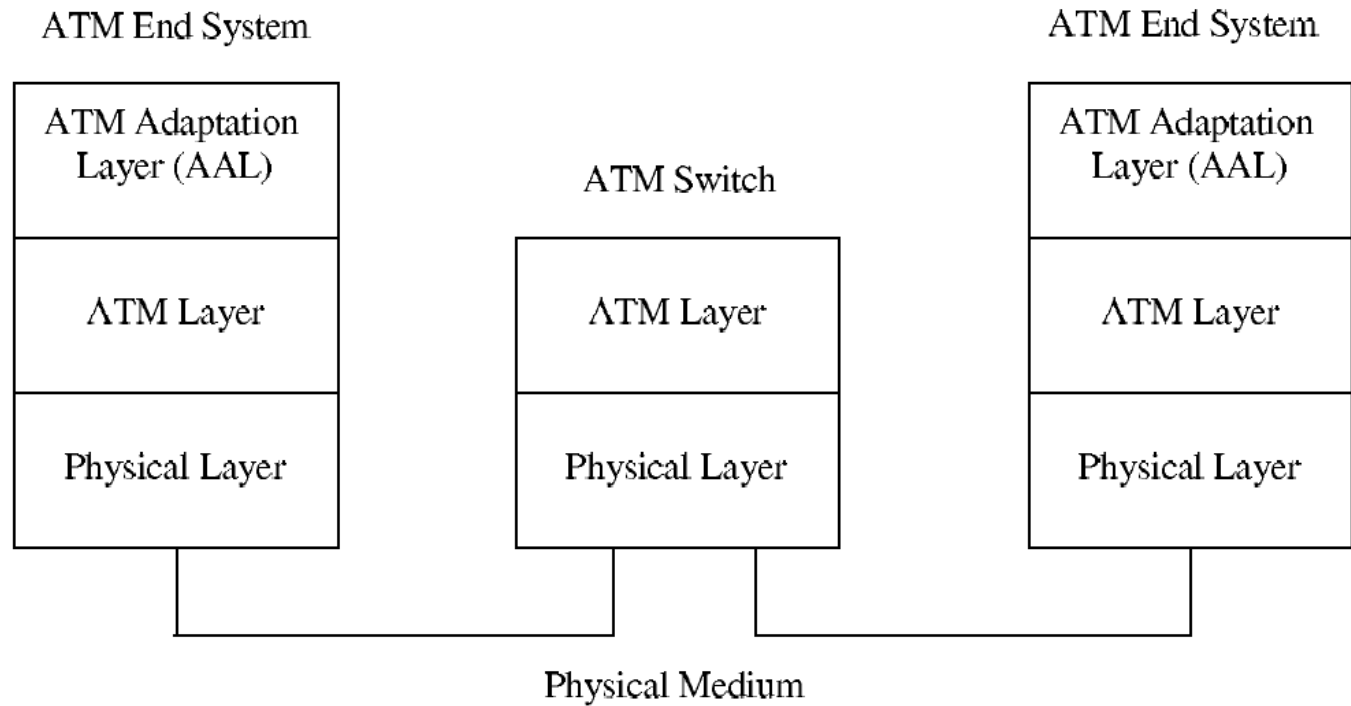


Figure 1: ATM protocol structure

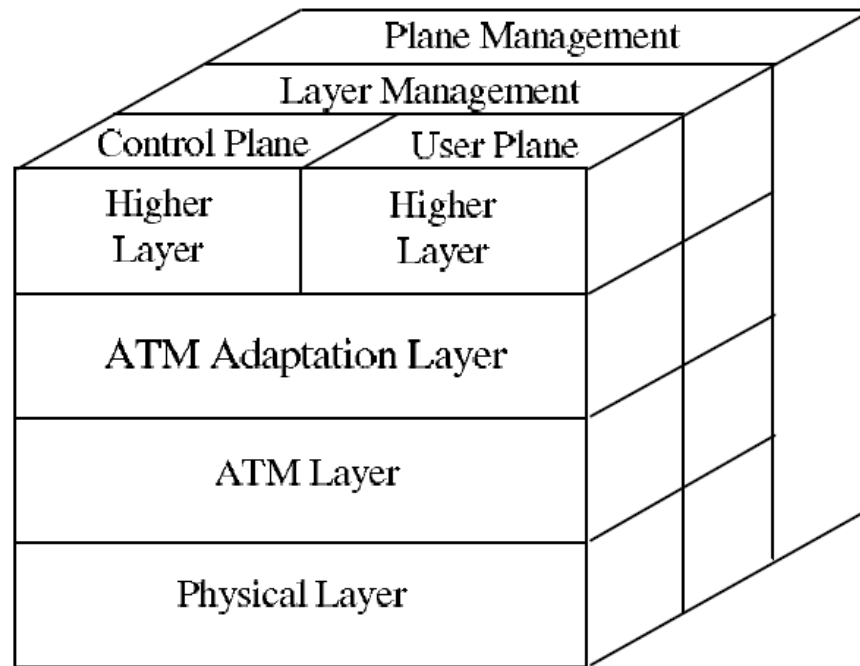


Figure 2: ATM model

- Physical Layer

Concerned with specifications of the transmission medium and signal encoding . Data rates specified include 155 and 622 Mbps with other data rates possible.

- ATM Layer

Defines transmission of data in fixed size cells and also defines the logical connections (Virtual circuits and virtual paths).

- ATM Adaptation Layer (AAL)

Supports transfer protocols not based on ATM. It maps higher layer information into ATM cells to be transported over an ATM network, then collects information from ATM cells for delivery to higher layers (e.g. a IP packet can be mapped to ATM cells).

(There are 3 planes in the protocol architecture: the *User* plane is for user traffic including flow and error control; the *Control* plane is for connection control; the *Management* plane manages the system as a whole and coordinates the planes and layers).

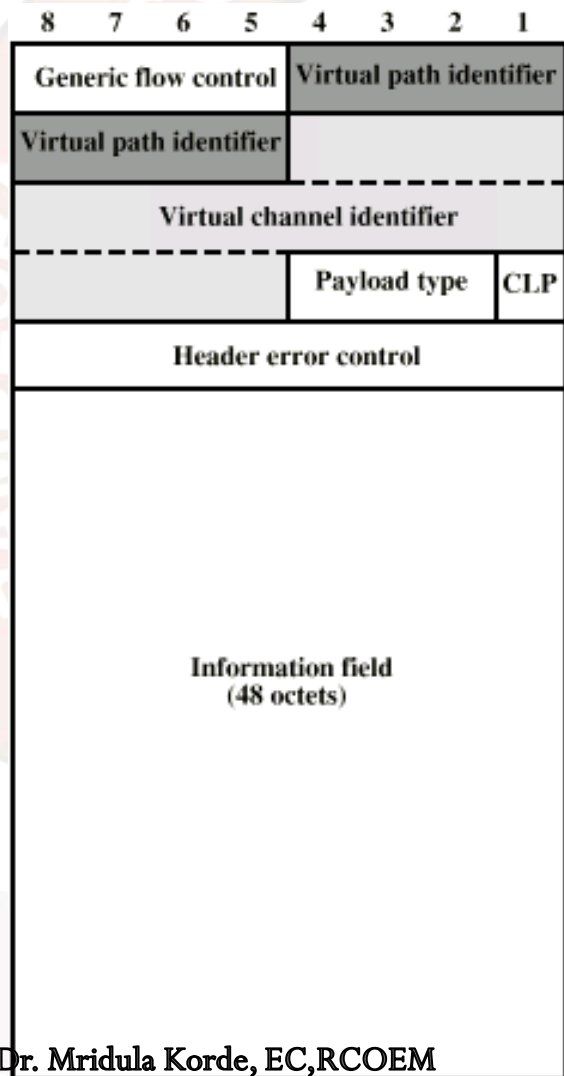
The AAL Layer

- The ATM adaptation layer (AAL) interfaces the higher layer protocols to the ATM Layer.
- It relays ATM cells both from the upper layers to the ATM layer and vice versa. When relaying information received from the higher layers to the ATM layer, the AAL segments the data into ATM cells.

The ATM Layer

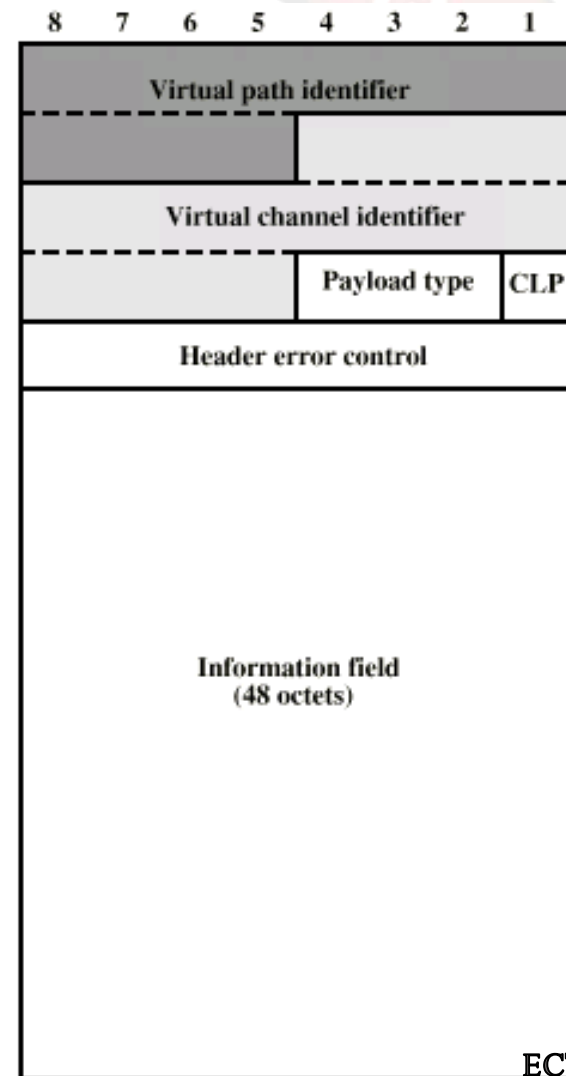
- The ATM layer provides an interface between the AAL and the physical layer.
- This layer is responsible for relaying cells from the AAL to the physical layer for transmission and from the physical layer to the AAL for use at the end systems.
- When it is inside an end system, the ATM layer receives a stream of cells from the physical layer and transmits cells with new data.
- When it is inside a switch, the ATM layer determines where the incoming cells should be forwarded to, modifies the corresponding connection identifiers, and forwards the cells to the next link.
- Moreover, it buffers incoming and outgoing cells, and handles various traffic management functions.

ATM Cell Format



(a) User-Network Interface

Dr. Mridula Korde, EC,RCOEM



(b) Network-Network Interface

ECTH41

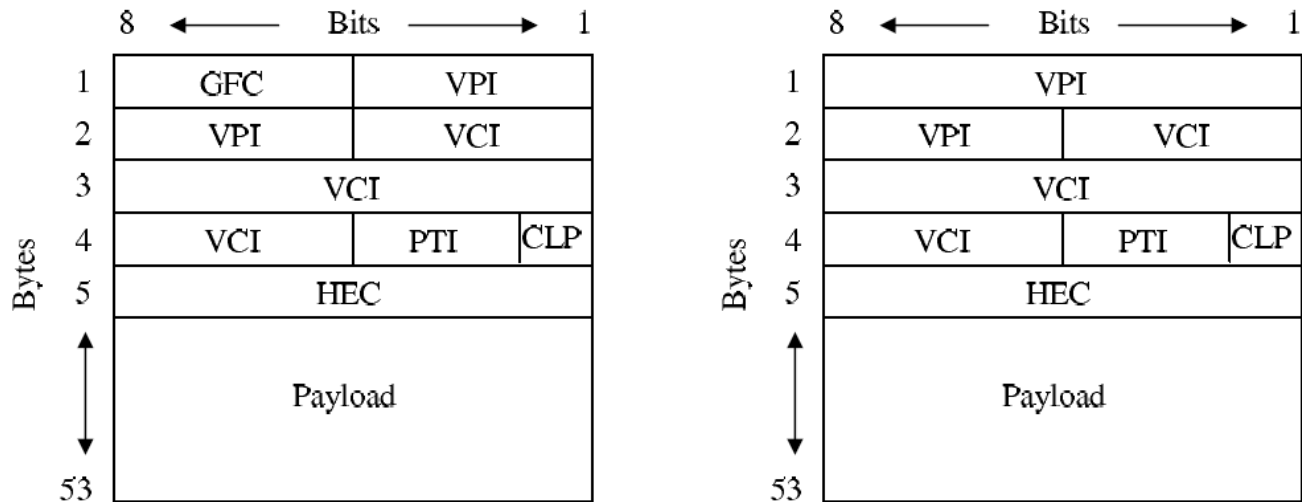
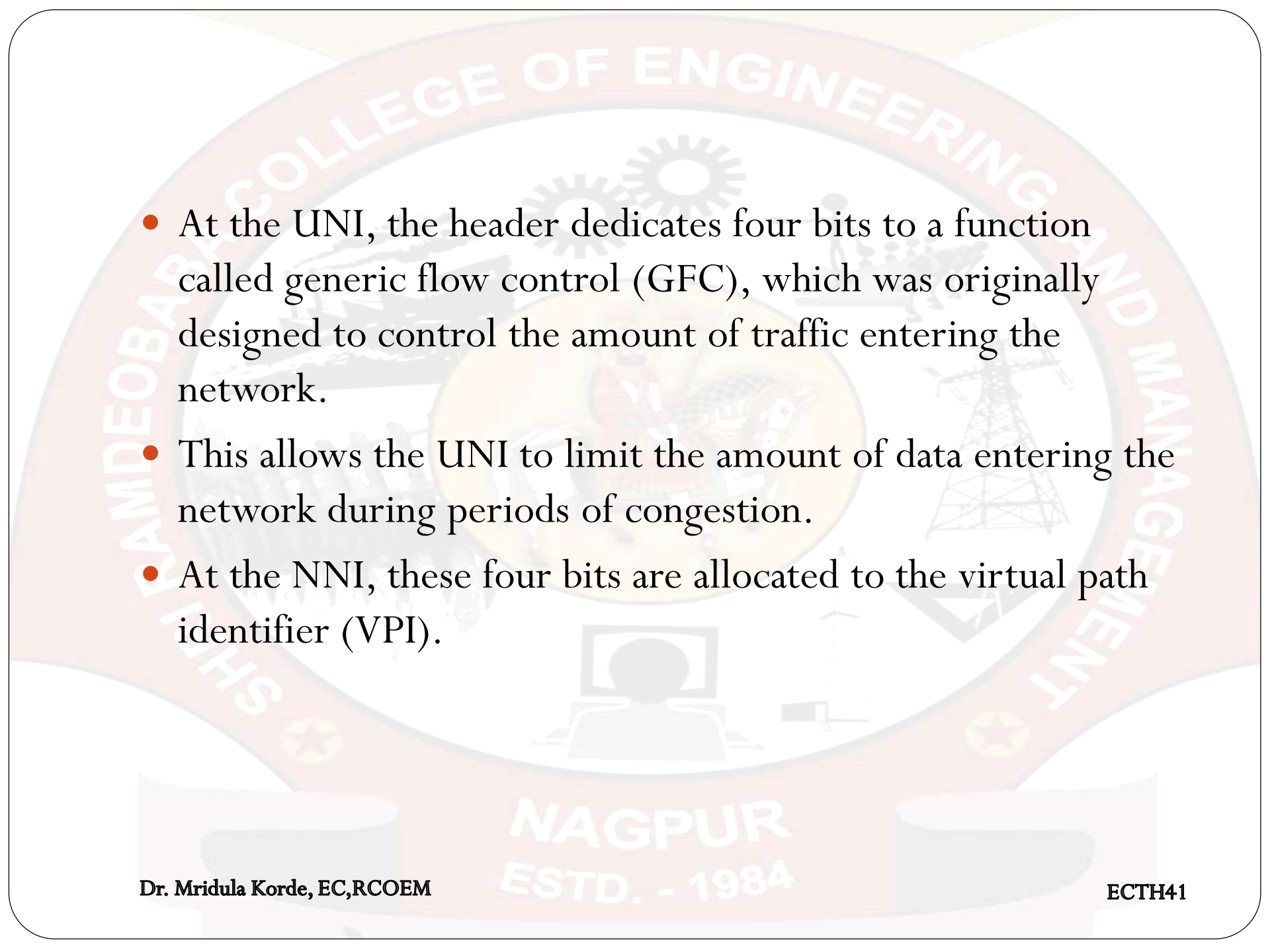


Figure 3: UNI (left) and NNI (right) ATM cell format

The fields in the ATM cell header define the functionality of the ATM layer. The format of the header for ATM cells has two different forms, one for use at the user-to-network interface (UNI) [10, 9] and the other for use internal to the network, the network-to-node interface (NNI), as shown in Figure 3.

- 
- At the UNI, the header dedicates four bits to a function called generic flow control (GFC), which was originally designed to control the amount of traffic entering the network.
 - This allows the UNI to limit the amount of data entering the network during periods of congestion.
 - At the NNI, these four bits are allocated to the virtual path identifier (VPI).

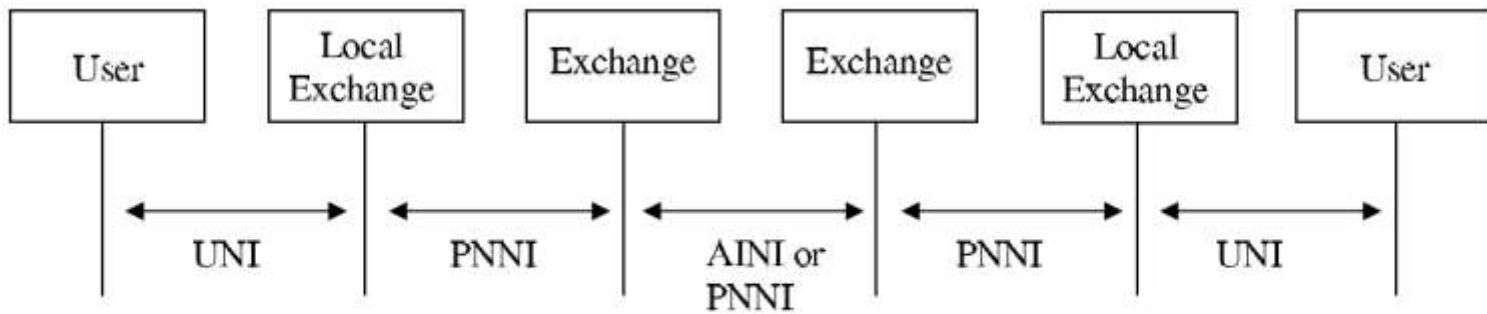


Figure 4: ATM network interfaces

Figure 4 gives an illustration of ATM Network Interfaces.

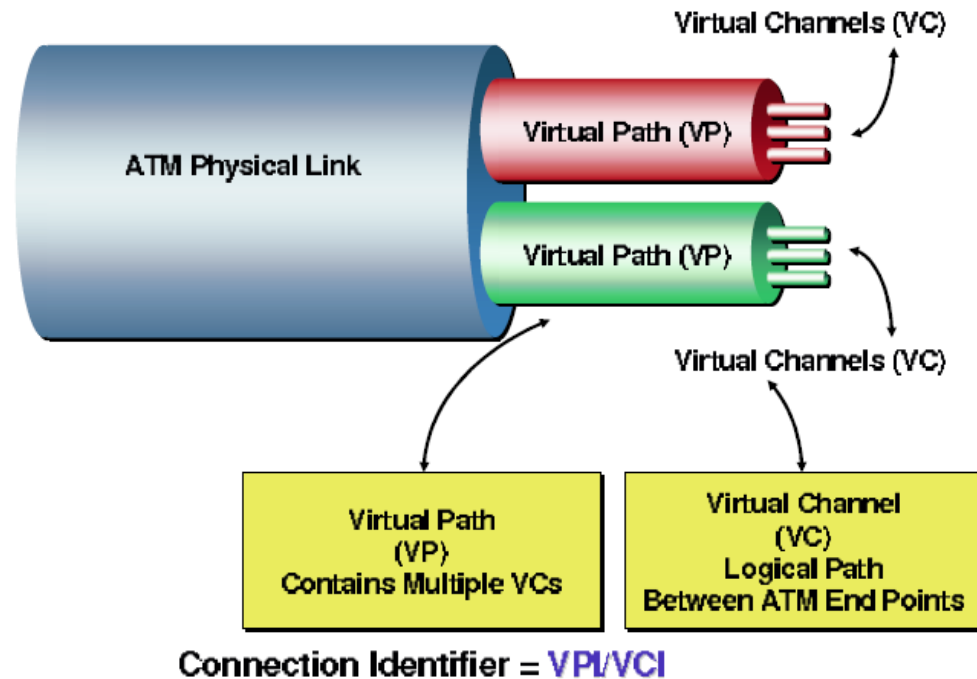
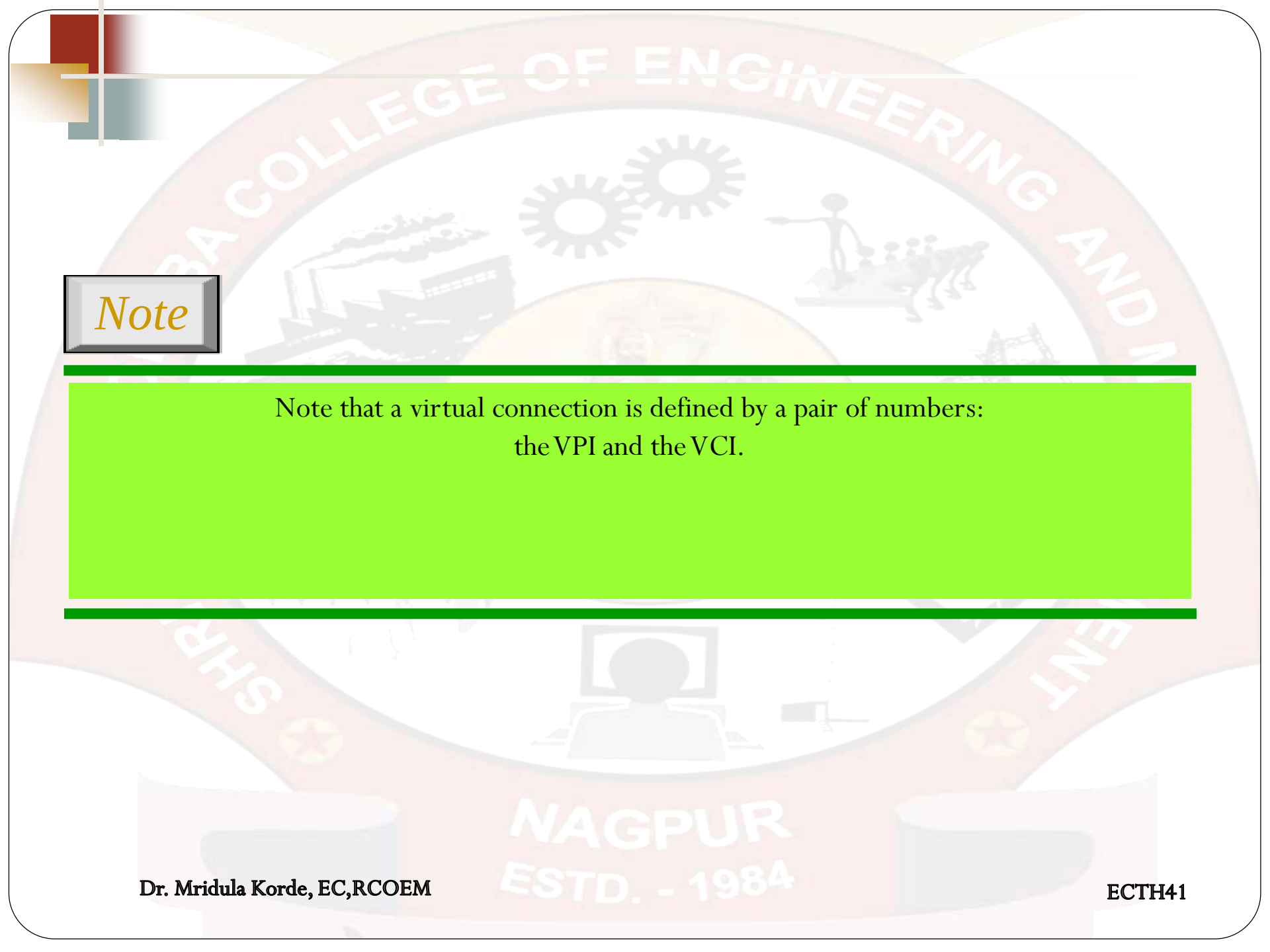
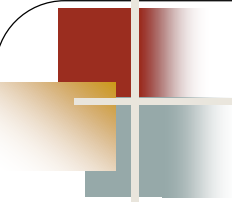


Figure 5: Virtual path and virtual channels

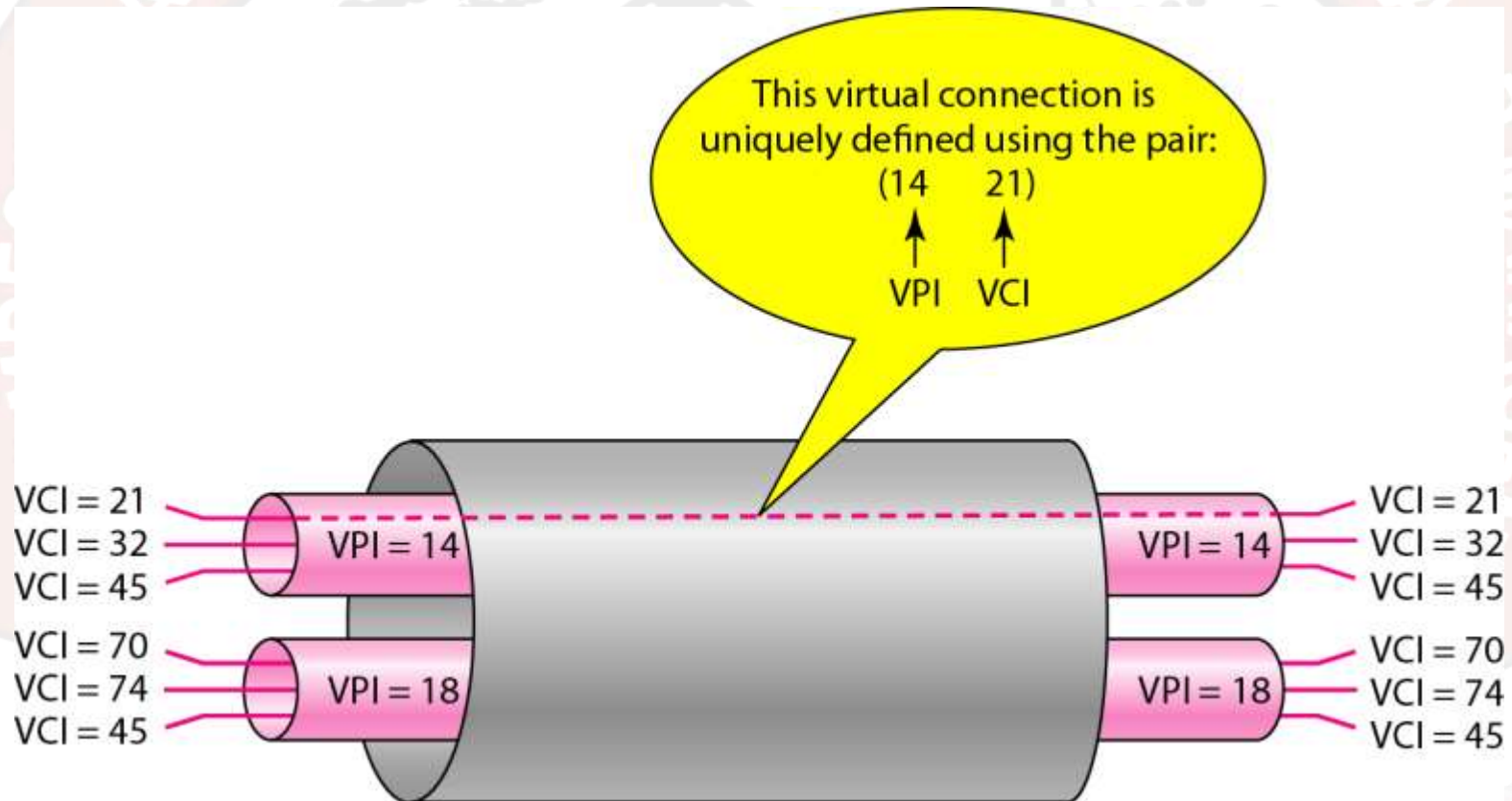
The VPI and the virtual channel identifier (VCI) together, as shown in Figure 5, form the routing field, which associates each cell with a particular channel or circuit, see Figure 6.



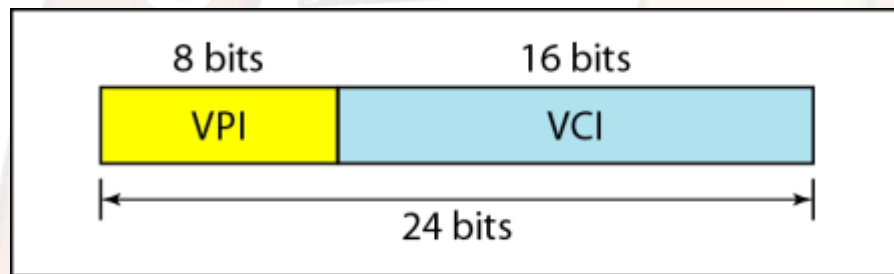
Note

Note that a virtual connection is defined by a pair of numbers:
the VPI and the VCI.

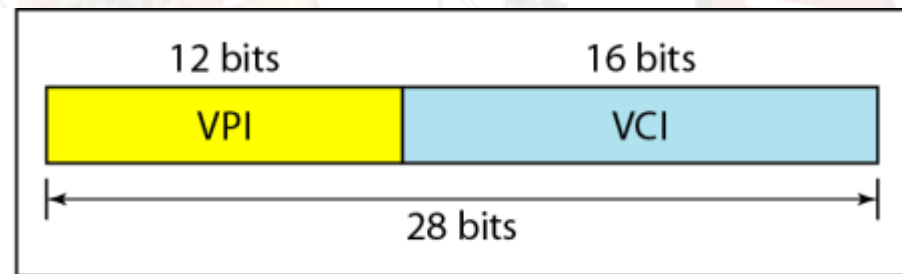
Connection identifiers



Virtual connection identifiers in UNIs and NNIs

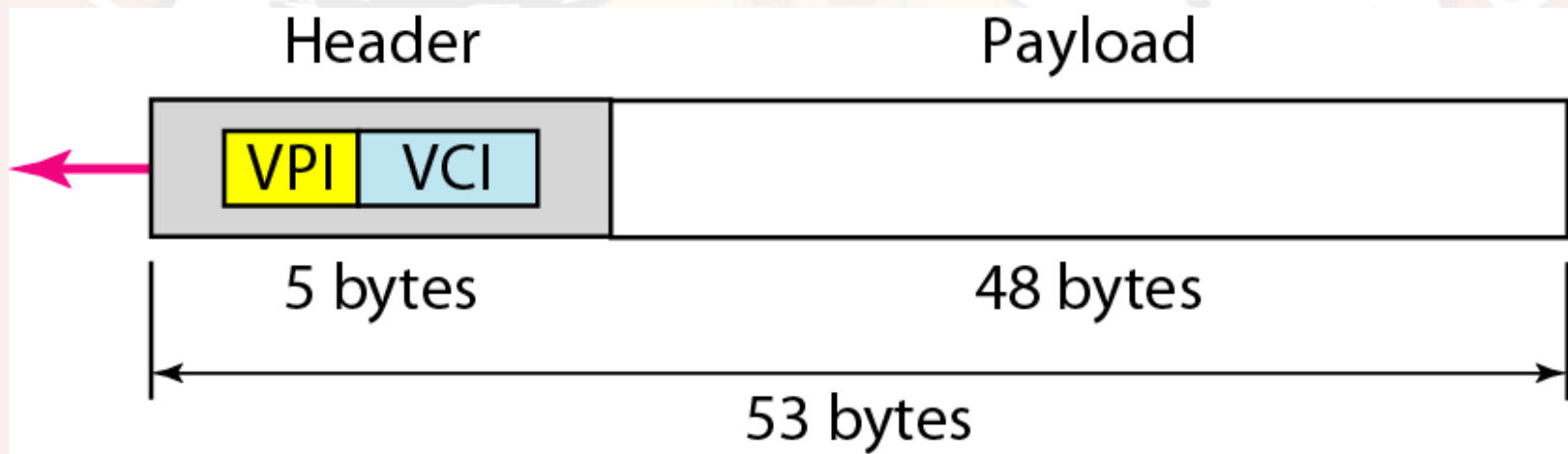


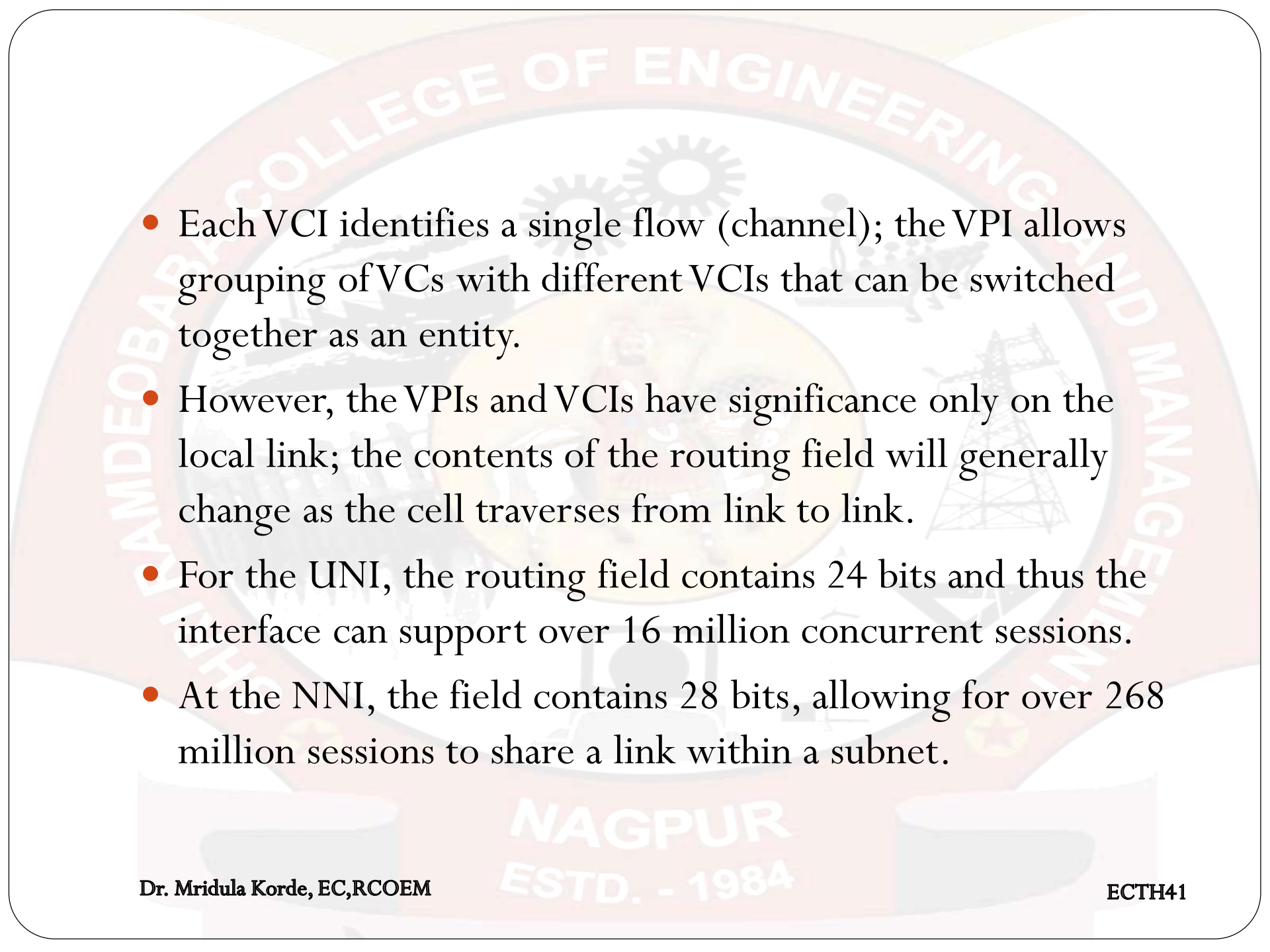
a. VPI and VCI in a UNI

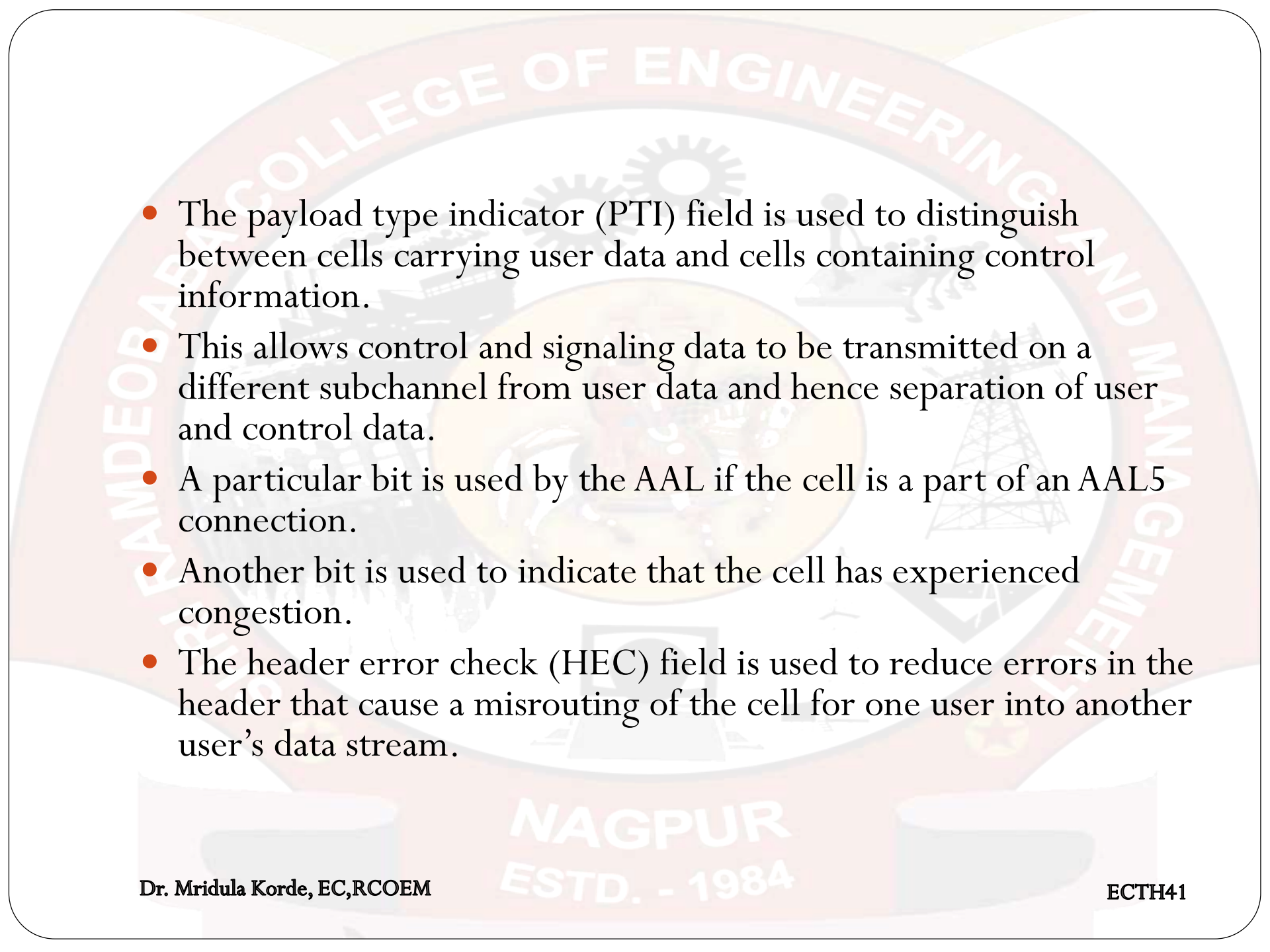


b. VPI and VCI in an NNI

An ATM cell



- 
- Each VCI identifies a single flow (channel); the VPI allows grouping of VCs with different VCIs that can be switched together as an entity.
 - However, the VPIs and VCIs have significance only on the local link; the contents of the routing field will generally change as the cell traverses from link to link.
 - For the UNI, the routing field contains 24 bits and thus the interface can support over 16 million concurrent sessions.
 - At the NNI, the field contains 28 bits, allowing for over 268 million sessions to share a link within a subnet.

- 
- The payload type indicator (PTI) field is used to distinguish between cells carrying user data and cells containing control information.
 - This allows control and signaling data to be transmitted on a different subchannel from user data and hence separation of user and control data.
 - A particular bit is used by the AAL if the cell is a part of an AAL5 connection.
 - Another bit is used to indicate that the cell has experienced congestion.
 - The header error check (HEC) field is used to reduce errors in the header that cause a misrouting of the cell for one user into another user's data stream.

The Physical Layer

- The physical layer defines the bit timing and other characteristics for encoding and decoding the data into suitable electrical/optical waveforms for transmission and reception on the specific physical media used.
- In addition, it also provides cell delineation function, header error check (HEC) generation and processing, performance monitoring, and payload rate matching of the different transport formats used at this layer.
- The Synchronous Optical Network (SONET), a synchronous transmission structure, is often used for framing and synchronization at the physical layer.

Traffic Management

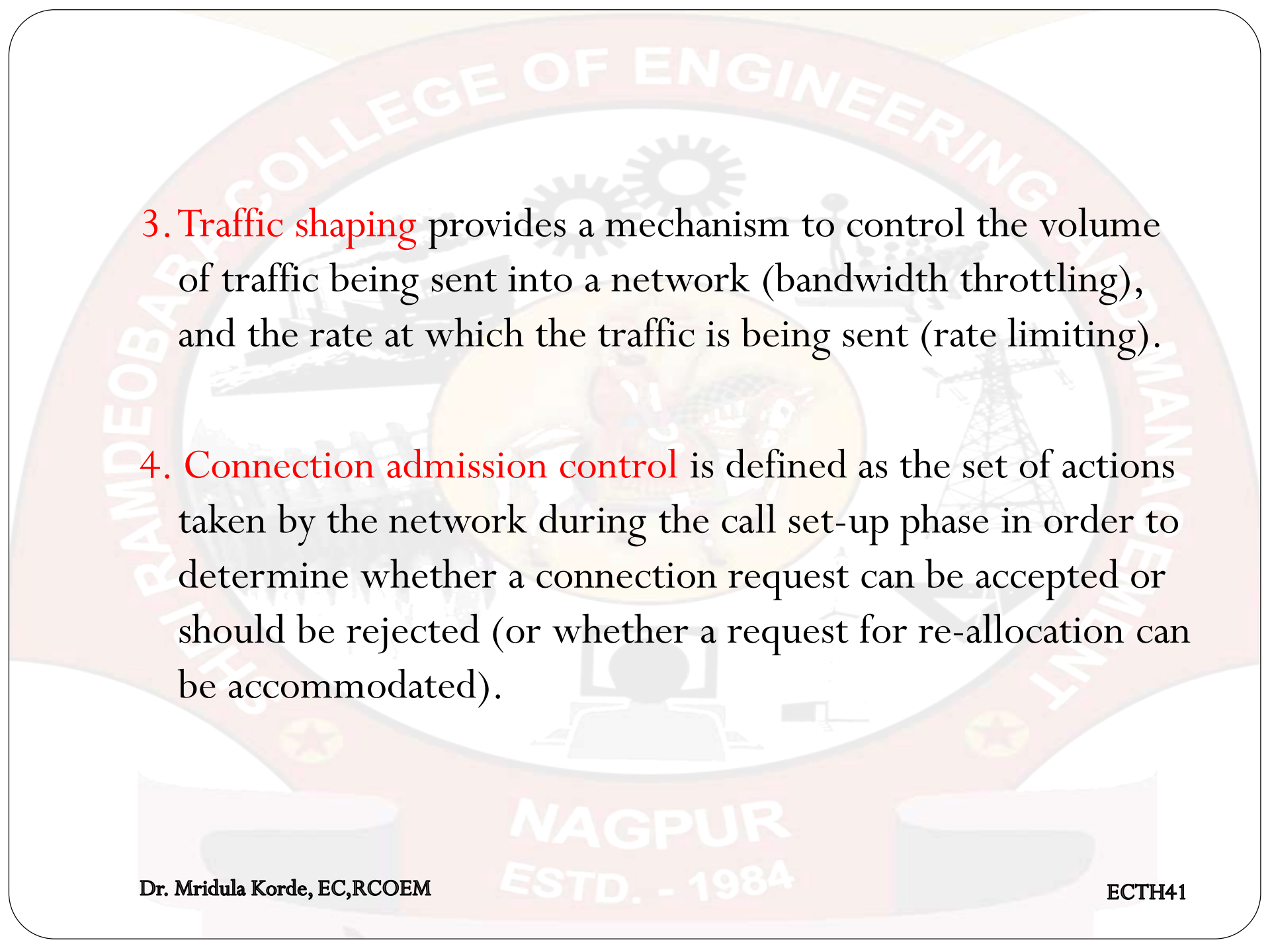
- In order for ATM networks to deliver guaranteed quality of service (QoS) on demand while maximizing the utilization of available network resources, effective traffic management mechanisms are needed.

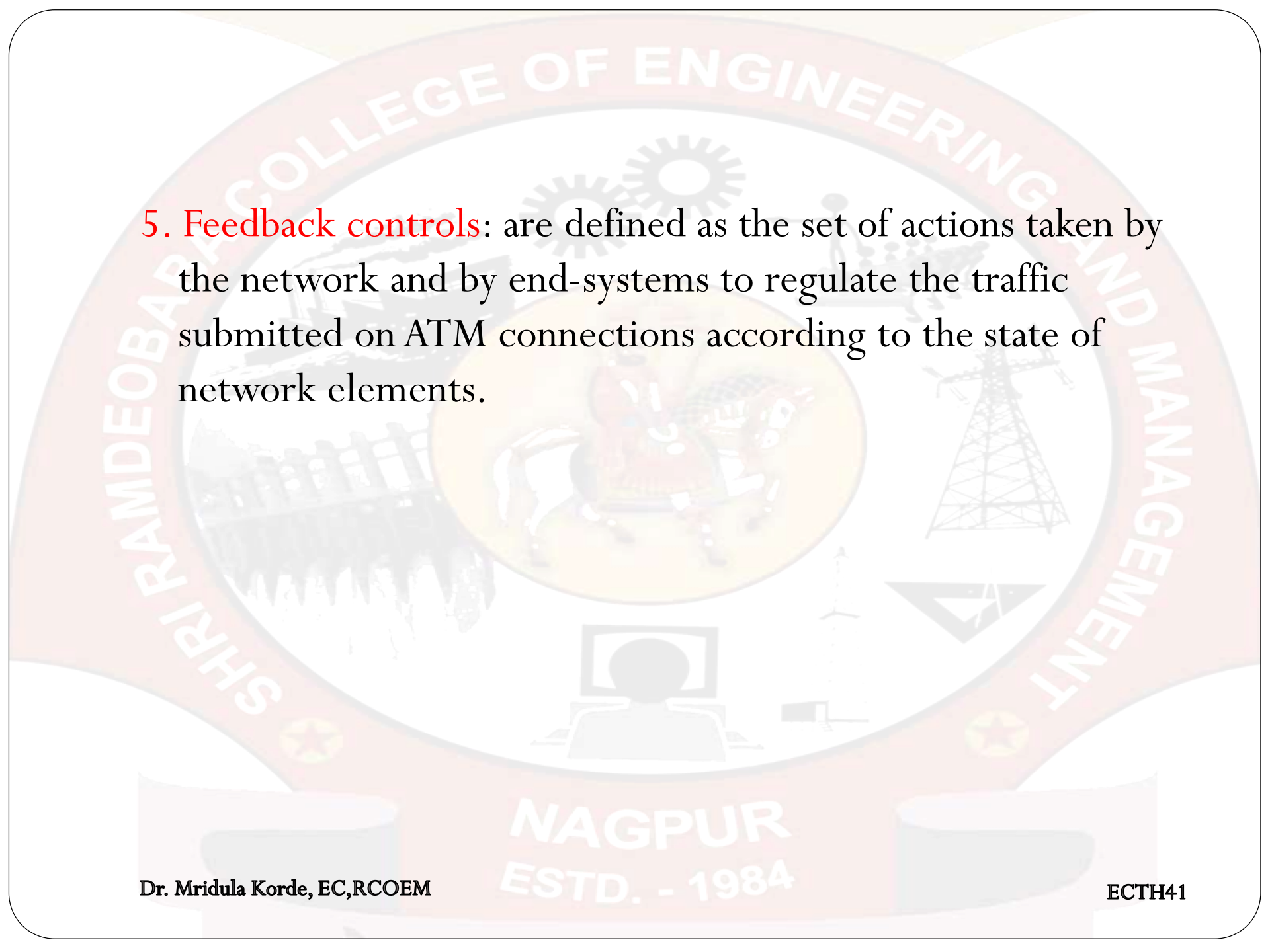
1. **Network Resource Management:** is used in broadband networks to keep track of the way link resources are allocated to connections.

The two primary resources that are tracked by network resource management are capacity (bandwidth) and connection identifiers.

2. Traffic policing: is monitoring network traffic for conformity with a traffic contract.

- An application that wishes to use the broadband network to transport traffic must first request a connection, which involves informing the network about the characteristics of the traffic and the quality of service (QOS) required by the application.
- This information is stored in a traffic contract.
- If the connection request is accepted, the application is permitted to use the network to transport traffic.
- The main purpose of this function is to protect the network resources from malicious connections and to enforce the compliance of every connection to its negotiated traffic contract.

- 
3. **Traffic shaping** provides a mechanism to control the volume of traffic being sent into a network (bandwidth throttling), and the rate at which the traffic is being sent (rate limiting).
4. **Connection admission control** is defined as the set of actions taken by the network during the call set-up phase in order to determine whether a connection request can be accepted or should be rejected (or whether a request for re-allocation can be accommodated).



5. **Feedback controls**: are defined as the set of actions taken by the network and by end-systems to regulate the traffic submitted on ATM connections according to the state of network elements.

Quality of Service Attributes

While setting up a connection on ATM networks, users can negotiate with the network the following parameters related to the desired quality of service:

Peak-to-peak cell delay variation (peak-to-peak CDV): is the delay experienced by a cell between network entry and exit points is called the cell transfer delay.

It includes propagation delays, queueing delays at various intermediate switches, and service times at queueing points.

Maximum cell transfer delay (maxCTD): is a measure of variance of CTD. High variation implies larger buffering for delay sensitive traffic such as voice and video

Cell loss ratio (CLR): The percentage of cells that are lost in the network because of error or congestion and are not delivered to the destination, i.e.,

- $CLR = \# \text{ Lost Cells} / \# \text{ Transmitted Cells}$

ATM Cell Format

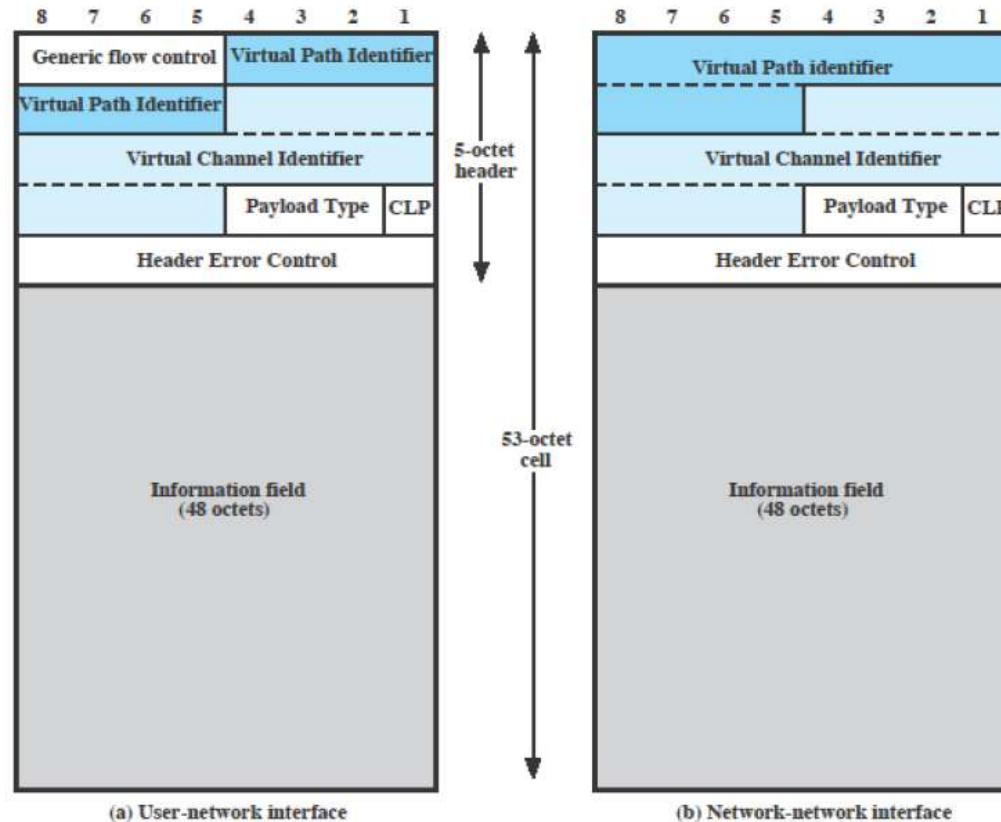


Figure 11.6 ATM Cell Format

ATM Cell Format ...

Fixed size cell has a 5-byte header and 48-byte data field.

- Use of small cells may reduce queuing delay for a high-priority cell (since it waits less if it arrives slightly behind a lower priority cell that has gained access to the transmitter).
- Fixed size cells can be switched more efficiently (i.e. easier to implement switching mechanism in hardware and build more scalable switches). This is important for the high data rates of ATM.

Header format

GFC (4-bits): *Generic Flow Control* to assist the customer in controlling traffic flow based on QoS. This field is only retained at the user-network interface (UNI) and not retained at the network-network interface (NNI) between ATM switches.

VPI (8-bits at the UNI and 12-bits at the NNI allowing more VPCs to be represented within the network): *Virtual Path Identifier* is a routing field for the network.

VCI (16-bits): *Virtual Channel Identifier* is used for routing to and from an end user.

ATM Cell Format ...

- Header format (continued)

PTI (3-bits): *Payload Type* indicates type of information in the information field.

The MSB indicates the type of ATM cell that follows. This bit set to 0 indicates user data; a bit set to 1 indicates network management data.

The middle bit indicates whether the cell experienced congestion in its journey from source to destination. This bit is also called the Explicit Forward Congestion Indication (EFCI) bit. This bit is set to 0 by the source; if an interim switch experiences congestion while routing the cell, it sets the bit to 1. After it is set to 1, all other switches in the path leave this bit value at 1.

000	User data cell, congestion not experienced, AAU* =0
001	User data cell, congestion not experienced, AAU=1
010	User data cell, congestion experienced, AAU=0
011	User data cell, congestion experienced, AAU=1
100	Network management segment cell
101	Network management end-to-end cell
110	Resource management cell
111	Reserved for future function

(*AAU represents ATM-user to ATM-user and is represented by the LSB of the *Payload Type* field. When AAU = 1, it indicates that the cell carries network management information. This allows for the insertion of network management cells into a user's VCC without affecting the user's data. So it provides in-band control information. This bit indicates the last cell in a block for AAL5 in user ATM cells.)

ATM Cell Format ...

Header format (continued)

CLP (1-bit): *Cell Loss Priority* provides guidance to the network in the event of congestion. CLP = 0 indicates a cell of high priority which should not be discarded. CLP = 1 indicates a cell that is subject to discard.

- When the CLP bit is set to 1, the interim switches sometimes discard the cell in congestion situations.
- An ATM user sets the CLP bit to 1 when a cell is created to indicate a lower priority cell. The ATM switch can set the CLP to 1 if the cell exceeds the negotiated traffic parameters of a VCC. Later if congestion is experienced, the cell that has been marked with CLP = 1 is subject to discard in preference to cells that fall within agreed traffic limits.

HEC (8-bits): *Header Error Code* is calculated based on the remaining 32-bits of the header. Polynomial used for the checksum is $x^8 + x^2 + x + 1$. In the case of ATM, the input to calculate the checksum is only 32-bits compared to the 8-bits for the code.

- The HEC only checks the ATM header and not the ATM payload. Checking the payload for errors is the responsibility of upper layer protocols.

Examples of ATM Workgroup Switches



**Marconi LE-155
ATM Workgroup Switch**



**Cisco Systems LS-1010
ATM Workgroup Switch**

Examples of ATM Backbone Switches



**Marconi ASX-4000
ATM Backbone Switch**



**Marconi ASX-1000
ATM Backbone Switch**



**Marconi ASX-200
ATM Backbone Switch**

Examples of ATM WAN Switches



**Cisco Systems
Catalyst 8500 Multiservice
Switch Router**



**Lucent GX 550
ATM Switch**

Examples of ATM Edge Devices



**Cisco Systems
Catalyst 5500 Switch**



**Marconi ES-3810
10/100 Ethernet Switch**



**Marconi ESX-3000
Campus Switch**