

Unit 1

WAVE PROPAGATION AND ANTENNA

Effects of Environment

Reflection: the result of a wave striking an object and bouncing off of it is called as reflection

Refraction: Refraction is the bending of a wave when it passes from one medium to another. The bending is caused due to the differences in density between the two substances.

Interference: It is a phenomenon where two or more waves overlap to form a resultant wave of either the same, greater or lesser amplitude.

Diffraction: it is the bending of waves around a barrier.

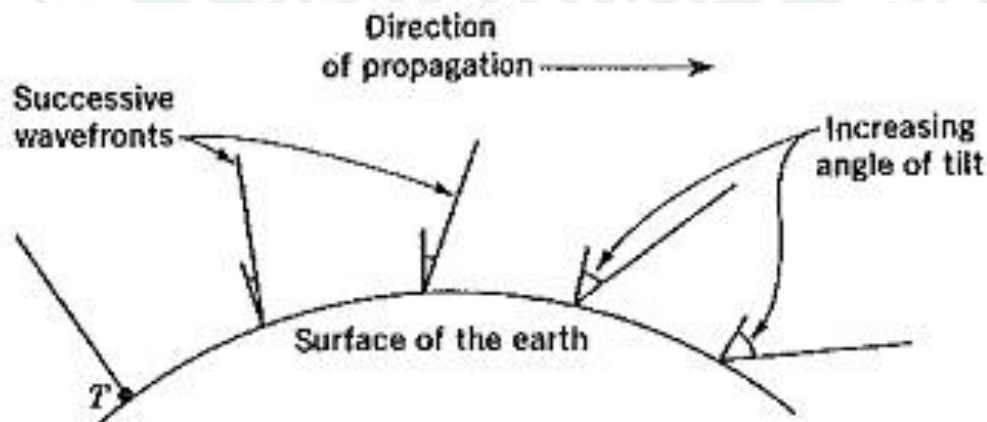
Absorption: the transfer of the energy of a wave to matter as the wave passes through it. Wave is absorbed by a material or may disappear.

Attenuation: it is the loss of signal strength

Classification based on modes of propagation

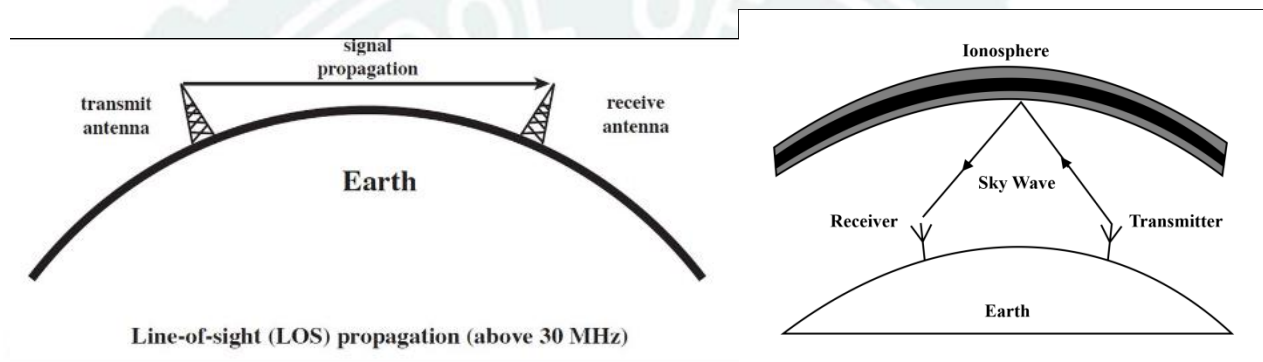
Ground wave propagation

Ground wave propagation of the wave follows the contour of earth. Such a wave is called as **direct wave**. The wave sometimes bends due to the Earth's magnetic field and gets reflected to the receiver. Such a wave can be termed as **reflected wave**.



Skywave Propagation

The waves, which are transmitted from the transmitter antenna, are reflected from the ionosphere. It consists of several layers of charged particles ranging in altitude from 30- 250 miles above the surface of the earth. Such a travel of the wave from transmitter to the ionosphere and from there to the receiver on Earth is known as **Sky Wave Propagation**. Ionosphere is the ionized layer around the Earth's atmosphere, which is suitable for sky wave propagation.



Line of Sight Propagation

It is a characteristic of wave propagation which means waves travel in a direct path from the source to the receiver.

Critical Frequency: it is the highest frequency radio wave, which sent normally towards layer of ionosphere gets reflected back and returns to the earth.

Maximum useable frequency: It is the highest radio frequency that can be used for transmission between two points via reflection from the ionosphere at a specified time.

Skip Distance: it can be defined as the minimum distance between the point of transmission and the point of reception of a skywave.

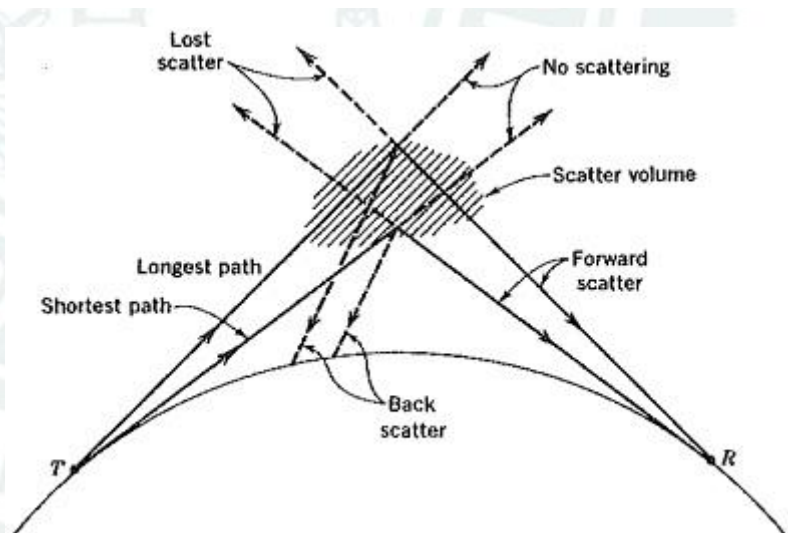
Fading: the decrease in the quality of the signal can be termed as fading. Fading occurs when there are significant variations and phase over time.

Duct Propagation: A radio wave propagation technique that allows the transmission of UHF and VHF electromagnetic waves through the region near the tropospheric layer of the atmosphere is known as **duct propagation**. Basically in duct propagation, despite being reflected from the ionosphere or gliding over the surface of the earth, the waves propagate from an end to another by undergoing successive refraction from the troposphere.

Troposphere scatter propagation

Tropospheric scatter, also known as troposcatter, is a method of communicating with microwave radio signals over considerable distances – often up to 500 kilometres (310 mi) and further depending on frequency of operation, equipment type, terrain, and climate factors. This method of propagation uses the tropospheric scatter phenomenon, where radio waves at UHF and SHF frequencies are randomly scattered as they pass through the upper layers of the troposphere. Radio signals are transmitted in a narrow beam aimed just above the horizon in the direction of the receiver station. As the signals pass through the troposphere, some of the energy is scattered back toward the Earth, allowing the receiver station to pick up the signal.

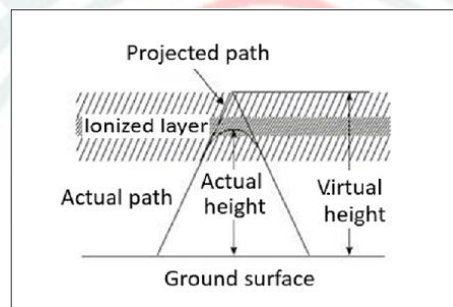
Normally, signals in the microwave frequency range travel in straight lines, and so are limited to line-of-sight applications, in which the receiver can be 'seen' by the transmitter. Communication distances are limited by the visual horizon to around 48–64 kilometres (30–40 mi). Troposcatter allows microwave communication beyond the horizon. It was developed in the 1950s and used for military communications until communications satellites largely replaced it in the 1970s.



Because the troposphere is turbulent and has a high proportion of moisture, the tropospheric scatter radio signals are refracted and consequently only a tiny proportion of the transmitted radio energy is collected by the receiving antennas. Frequencies of transmission around 2 GHz are best suited for tropospheric scatter systems as at this frequency the wavelength of the signal interacts well with the moist, turbulent areas of the troposphere, improving signal-to-noise ratios.

Actual Height: the actual path of the wave in the ionized layer is a curve and is due to refraction of wave. The height from this curve to earth surface is called actual height.

Virtual Height: It is the height from which the radio wave appears to be reflecting



Radiation Mechanism of an antenna

The sole functionality of an antenna is power radiation or reception. Antenna (whether it transmits or receives or does both) can be connected to the circuitry at the station through a transmission line. The functioning of an antenna depends upon the radiation mechanism of a transmission line.

A conductor, which is designed to carry current over large distances with minimum losses, is termed as a transmission line. For example, a wire, which is connected to an antenna. A transmission line conducting current with uniform velocity, and the line being a straight line with infinite extent, radiates no power.

For a transmission line, to become a waveguide or to radiate power, has to be processed as such

- If the power has to be radiated, though the current conduction is with uniform velocity, the wire or transmission line should be bent, truncated or terminated.
- If this transmission line has current, which accelerates or decelerates with a time-

varying constant, then it radiates the power even though the wire is straight.

- The device or tube, if bent or terminated to radiate energy, then it is called as waveguide. These are especially used for the microwave transmission or reception.

Two wire Mechanism

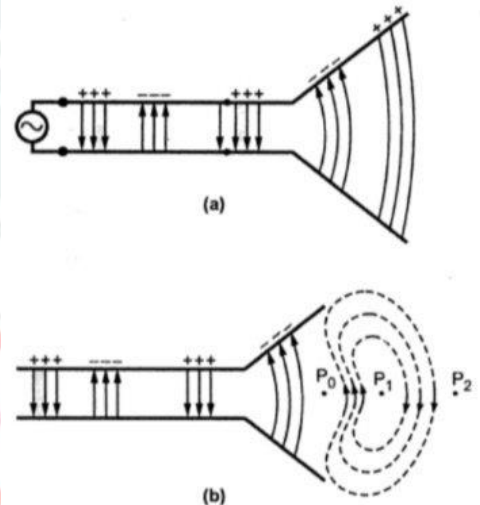
Consider an antenna driven by a transmission line with a voltage source conductor.

The electric lines of force are tangent to the electric field. The strength of electric lines of force is proportional to electric field intensity.

As the free electrons have a tendency to get separated from the atom, the electric lines of force operate on free electrons of each conductor and force to displace. Due to the charge movement, the current is produced and it produces magnetic field intensity.

The electric field lines travel from positive charges to negative charges while the magnetic field lines form closed loops encircling current-carrying conductors.

The charge distribution is due to electric field lines. Assuming a sinusoidal source between the two conductors, the electric field between the conductor is also sinusoidal.



Antenna Gain: Antenna gain is the ability of the antenna to radiate more or less in any direction compared to a theoretical antenna. If an antenna could be made as a perfect sphere, it would radiate equally in all directions. Such an antenna is theoretically called an isotropic antenna and does not in fact exist.

Directive Gain: It is defined as the ratio of radiation intensity due to the test antenna to isotropic antenna.

Directivity: The ratio of the radiation intensity in a given direction from an antenna to the radiation intensity averaged over all direction is termed as directivity.

Antenna Apertures: It is defined as the measure of the ability of the antenna to effectively receive the power radiated towards it.

Input impedance of antenna: Input impedance is defined as the ratio of the voltage and current at the pair of the input antenna terminals.

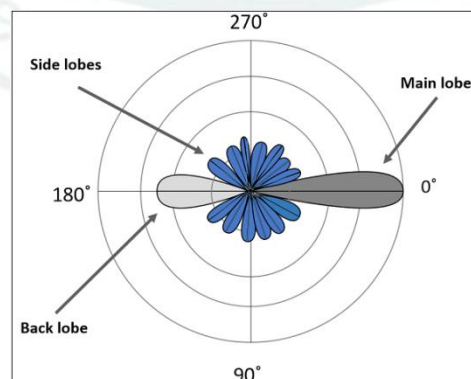
Polarisation: The polarization of an antenna is loosely defined as the direction of the electromagnetic fields produced by the antenna as energy radiates away from it. These directional fields determine the direction in which the energy moves away from or is received by an antenna.

Bandwidth: The bandwidth of an antenna refers to the range of frequencies over which the antenna can operate correctly.

Beamwidth: The beamwidth is the angular separation at which the magnitude of the directivity pattern decreases by a certain value from the peak of the main beam.

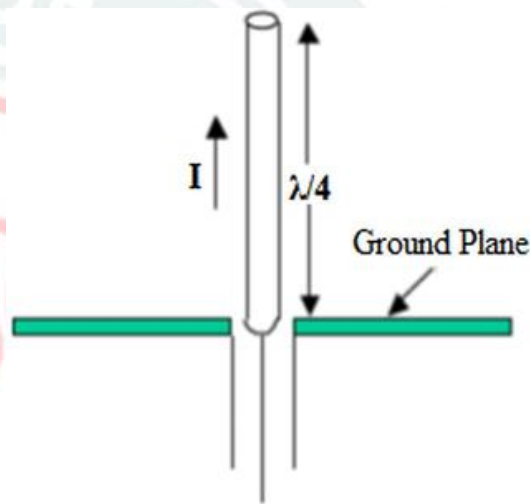
Radiation Pattern: The energy radiated by an antenna is represented by the Radiation pattern of the antenna. Radiation Patterns are diagrammatical representations of the distribution of radiated energy into space, as a function of direction.

- The major part of the radiated field, which covers a larger area, is the main lobe or major lobe. This is the portion where maximum radiated energy exists. The direction of this lobe indicates the directivity of the antenna.
- The other parts of the pattern where the radiation is distributed side wards are known as side lobes or minor lobes. These are the areas where the power is wasted.
- There is other lobe, which is exactly opposite to the direction of main lobe. It is known as back lobe, which is also a minor lobe. A considerable amount of energy is wasted even here.



Monopole Antenna

Monopole antennas, constitute a group of derivatives of dipole antennas. Here, only half of the dipole antenna is needed for operation. A metal ground plane (ideally of infinite size) is used, with respect to which the excitation voltage is applied to the half structure. The half structure for a regular dipole antenna is called a monopole antenna, in reference to the presence of only one physical side. A similar half structure for a folded dipole antenna is called a folded monopole antenna. The presence of the ground plane allows the monopole antenna to operate as electrically equivalent to a dipole antenna. The ground plane equivalently replaces the lower half by an imaging principle, similar to creating an optical image through a mirror. For the currents in the monopole and dipole structures to be the same, one needs the source voltage of the equivalent dipole antenna to be twice that of the monopole antenna. As a result, the input impedance of the monopole structure is half that of the corresponding dipole structure.



Due to the imaging principle, the polarization of radiation and radiation patterns of a monopole antenna is the same as that of its equivalent dipole antenna. However: the monopole antenna has a field only in the top half of the space, having zero radiation below the ground plane.

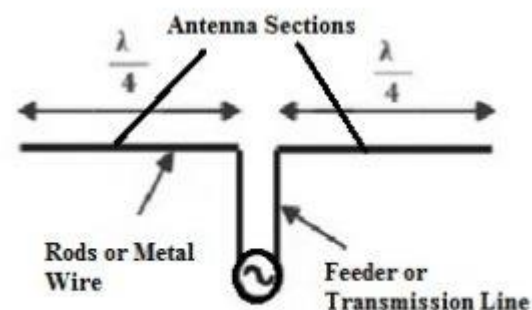
Dipole Antenna

A dipole antenna is a type of RF antenna that includes two conductive elements like wires or rods where the metal wire length is half of the highest wavelength approximately in free space at the operation of frequency. At the center of the antenna, the conductive materials are separated through an insulator which is called an antenna section. The RF voltage source is given to the middle of the antenna then the voltage & current supplying throughout the two conductive elements generate an electromagnetic or radio signal and this signal is radiated

outside of the antenna. At the center of this antenna, the voltage is minimum and current is maximum whereas the voltage is maximum & current is minimum at the dipole antenna two ends.

A dipole antenna includes two conductive elements like wires and rods or wires where the feeder at the center & radiating sections of the antenna are on either side. The metal wires' length is half of the highest wavelength that is $\lambda/2$ within free space at the frequency of operation. The conductive element in the antenna is split in the middle into two sections through an insulator which is called an antenna section. These sections are simply connected to a coaxial cable or feeder at the middle of the antenna. We know that wavelength is the distance among two consecutive highest or lowest points.

Once the RF voltage source is applied to the center of the two sections in the antenna then the flow of voltage & current throughout the two conductive elements can generate an electromagnetic or radio wave signal to be radiated outside of the antenna. At the center of this antenna, the voltage is minimum and the current is maximum. In opposition, the current is minimum & the voltage is maximum at the antenna's ends. This is the current distribution of dipole antenna. So, this antenna converts the signals from electrical to RF electromagnetic & emits them at the transmitting end & changes RF electromagnetic signals into electrical at the receiving side.



Omnidirectional Antenna

An omnidirectional antenna is a wireless transmitting or receiving antenna that radiates or intercepts radio-frequency (RF) electromagnetic fields equally well in all horizontal directions in a flat, two-dimensional (2D) geometric plane. Omnidirectional antennas are used in most consumer RF wireless devices, including cellular telephone sets and wireless routers.

Yagi antenna

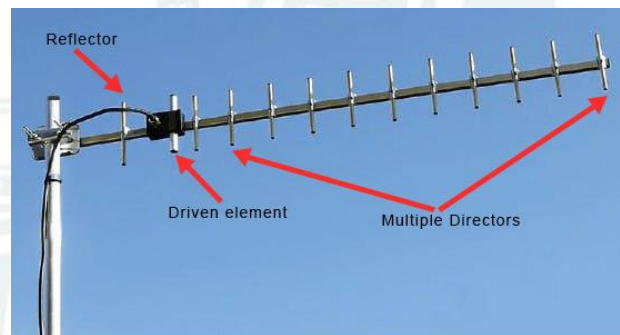
Yagi-Uda antenna is the most commonly used type of antenna for TV reception over the last few decades. It is the most popular and easy-to-use type of antenna with better performance, which is famous for its high gain and directivity. The frequency range in which the Yagi-Uda antennas operate is around 30 MHz to 3GHz which belong to the VHF and UHF bands.

It is seen that there are many directors placed to increase the directivity of the antenna. The feeder is the folded dipole. The reflector is the lengthy element, which is at the end of the structure. The center rod like structure on which the elements are mounted is called as boom. The element to which a thick black head is connected is the driven element to which the transmission line is connected internally, through that black stud. The single element present at the back of the driven element is the reflector, which reflects all the energy towards the direction of the radiation pattern. The other elements, before the driven element, are the directors, which direct the beam towards the desired angle.

Advantages

The following are the advantages of Yagi-Uda antennas –

- High gain is achieved.
- High directivity is achieved.
- Ease of handling and maintenance.
- Less amount of power is wasted.
- Broader coverage of frequencies.



Disadvantages

The following are the disadvantages of Yagi-Uda antennas –

- Prone to noise.
- Prone to atmospheric effects.

Applications

The following are the applications of Yagi-Uda antennas –

- Mostly used for TV reception.
- Used where a single-frequency application is needed

Rhombus antenna

The type of antenna which is formed by connecting long wires in the orientation of a rhombus is known as a rhombic antenna. Its basis of operation is similar to travelling wave radiator. It is also known as a diamond antenna. These are highly directional broadband antenna that radiates maximal power along the axis.

Rhombic antenna generally operates on high frequency and very high-frequency ranges. The operating range exists between 30 MHz to 300 MHz.

Construction

A rhombic antenna is a type of travelling wave antenna which utilizes the principle of travelling wave radiator. It consists of 4 conducting long wires connected in the form of rhombus or diamond. The complete antenna structure is placed horizontally above the surface of the ground. A rhombic antenna may also be considered as a combination of 2 V antennas or inverted V antennas forming an obtuse angle. The excitation to the antenna is provided through feed lines. This feed line can be two-wire transmission lines or coaxial wire that excites the whole structure. Also, to avoid reflections of the travelling wave the opposite end of the antenna wrt feed line is terminated with a properly adjusted resistor. This leads to cause the absence of standing waves in any of the legs of the antenna. The value of load resistance is generally around 600 to 800 Ω .

While designing the rhombic antenna, necessarily it has to be kept in mind that length of all the four conductive wires must be equal, ranging between quarter to one wavelength or more. However, the opposite acute angles of the rhombus must be equal. The physical form and the size of the antenna are decided by the length of the side and the angle between the vertex. The feed point and the terminating point of the antenna must be properly aligned as the line joining the two forms the axis of the antenna. Generally, the antenna is terminated with a value equivalent to characteristic impedance thereby causing the non-resonant condition to establish. Thus radiation characteristics of the antenna are unidirectional.

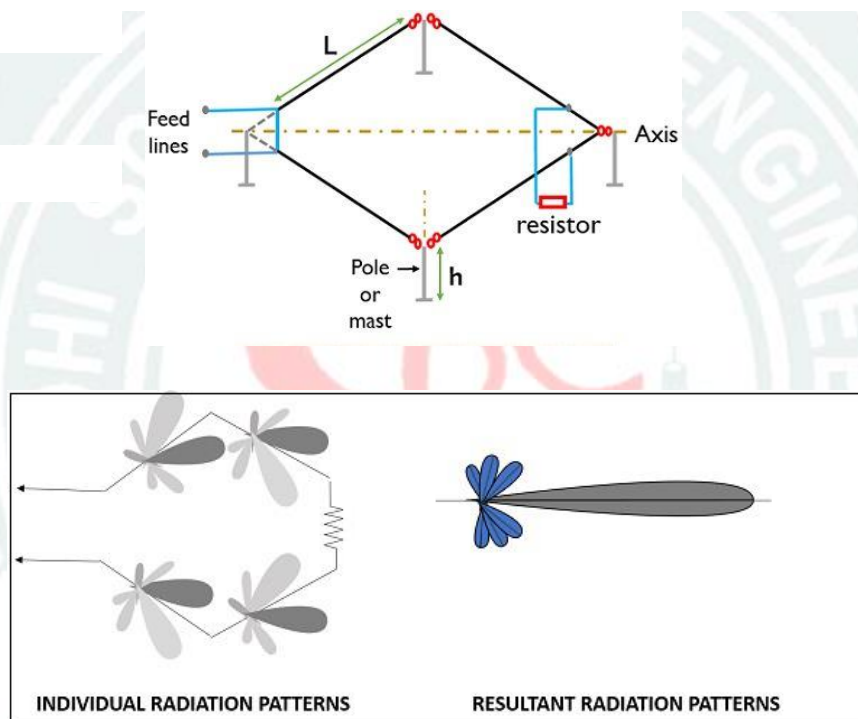
Working

When the power through the feed lines either two-wire transmission line or coaxial line is provided to the antenna. Then the generated current travels through the legs of the rhombic antenna. These current through the antenna generate radio waves that progress in one direction through the legs of the antenna.

We have already discussed that resistive termination is done at the end of the antenna. However, sometimes the terminating end is kept open. This is so because terminating the end reduces reflections of the transmitted wave. But on resistive

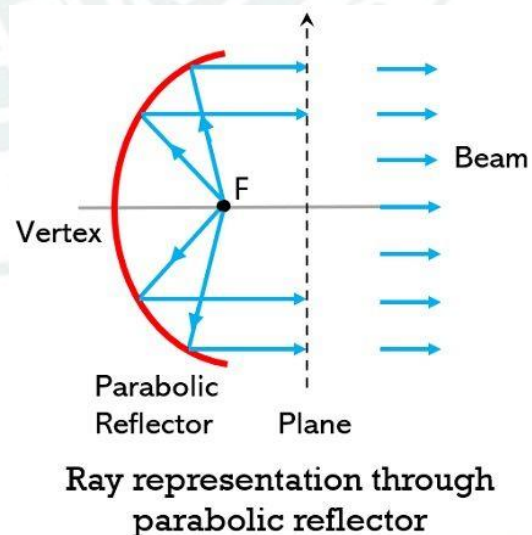
terminating, a certain amount of power is wasted. This power loss is around 35 to 50%. In the same way if despite terminating we kept the ends open then this causes the reflection of the waves back towards the feed point.

Thus in case of the rhombic antenna, the losses are almost equivalent in both the cases. However, in case of open ends, the antenna exhibits bidirectional behaviour. As the radiation is emitted not only in the front direction but also in the back direction. While if the rhombic antenna is resistor terminated then the antenna possesses unidirectional nature and the radiation is perfectly radiated in the forward direction only. Thereby supporting point to point communication.



Dish Antenna

The crucial function of the parabolic reflector is to change the spherical wave into a plane wave. Feed element is present at the focus. At the focus when a feed antenna is placed which is nothing but an isotropic source then the waves are emitted from the source.



The radiating element used at the focus is generally dipole or horn antenna which is used to illuminate the reflecting surface. The waves emitted from the source, incident on the surface of the reflector and are further reflected back as a plane wave. The plane waves travel in the direction parallel to the axis.

Advantage

It reduces minor lobe.

It offers high gain and directivity.

The amount of power wastage is comparatively less than other antennas.

Disadvantage

The size of the structure is quite large.

The overall cost of the system is high.

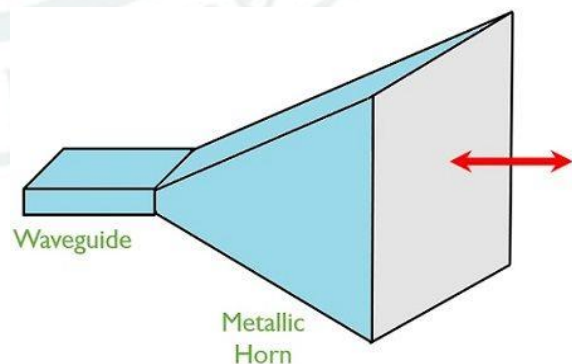
Horn Antenna

It is an antenna that consists of a flaring metal waveguide shaped like a horn to direct radio waves in a beam. These operate in ultra high and super high frequencies ranging between 300MHz to 30 GHz.

It consists of a short length of rectangular or cylindrical metal tube closed at one end, flaring into an open ended conical or pyramidal shaped horn on the other end. The radio waves are usually introduced into the waveguide by a coaxial cable attached to the side. In case of propagation through the waveguide, the propagating field is constrained by the walls of the waveguide. Thereby the field fails to spread spherically while this is not the case with free space propagation.

When the traversing field reaches the end of the waveguide then also it propagates in the same manner. At the end of the structure, spherical waveforms are achieved.

The region is said to be the transition region where the guided propagation changes to free space propagation. The free space impedance is different from the impedance of the waveguide thus flaring of the



waveguide is one to reduce reflections.

Advantage

They can operate over a wide range of frequencies

Good directivity

High gain

Disadvantage

These antennas will radiate energy in spherical wavefront shape, thus this antenna does not have a directive or sharp beam.

Unit 2

TRANSMISSION LINE

A **transmission line** is a connector which transmits energy from one point to another. The study of transmission line theory is helpful in the effective usage of power and equipment.

There are basically four types of transmission lines:

- Two-wire parallel transmission lines
- Coaxial lines
- Strip type substrate transmission lines
- Waveguides

While transmitting or while receiving, the energy transfer has to be done effectively, without the wastage of power. To achieve this, there are certain important parameters which has to be considered.

CHARACTERISTIC IMPEDANCE

If a uniform lossless transmission line is considered, for a wave travelling in one direction, the ratio of the amplitudes of voltage and current along that line, which has no reflections, is called as **Characteristic impedance**.

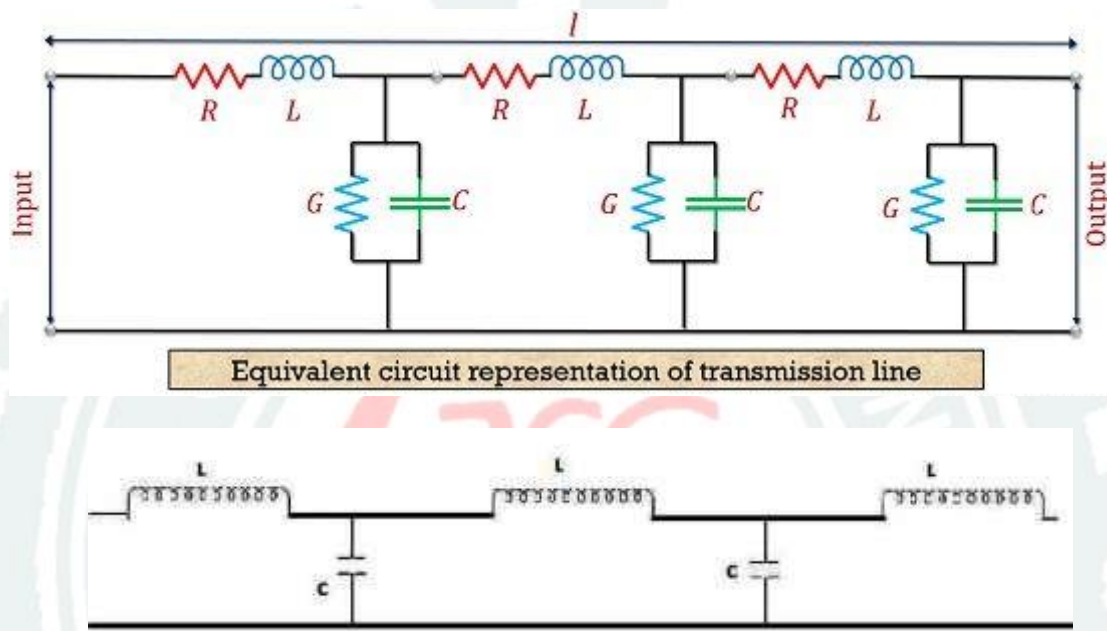
It is denoted by **Z_0**

Equivalent circuit of Transmission Line

We know that the conductors are present in a two wire line. Dielectrics are also present between them. It is also clear that conductors can be of any length. Conductors also have some diameter. If both the length and diameter are associated with the conductor then resistance and inductance must be present there. If wires are separated from each other by placing the dielectric between them then the leakage of charge will take place, because the dielectric that we are using is an insulating material which can't be a perfect insulator. This can be explained well by introducing the concept of shunt conductance. It is denoted by G .

All the quantities i.e. Resistance, Capacitance and Shunt conductance are calculated with respect to the length of the conductor.

The inductive resistance has a greater value than that of resistance with respect to the radio frequencies. On the other hand the value of the capacitive susceptance is also much more than that of the shunt capacitance. These all the quantities are working along the length of line. So, if we ignore both R and G then we can consider the circuit as lossless.



Characteristic impedance (Z_0) is the most important parameter for any transmission line. It is a function of geometry as well as materials and it is a dynamic value independent of line length; you can't measure it with a multimeter. It is related to the conventional distributed circuit parameters of the cable or conductors by: :

$$Z_0 = \sqrt{(R + j\omega L) / (G + j\omega C)}$$

where: R is the series resistance per unit length (Ω/m)

L is the series inductance (H/m)

G is the shunt conductance (mho/m)

C is the shunt capacitance (F/m).

L and C are related to the velocity factor by:

velocity of propagation = $1/\sqrt{LC}$ =

For an ideal, lossless line $R = G = 0$ and Z_0 reduces to $\sqrt{L/C}$. Practical lines have some losses which attenuate the signal, and these are quantified as an attenuation factor for a specified length and frequency

Losses in Transmission line

In a transmission line we face the following types of losses:

- (i) **Radiation loss:** It happens when the distance between the conductors in the transmission line is comparable to the wavelength. In such cases the electromagnetic and electrostatic field of the conductors acts as small antennas which conducts out energy to the nearby conducting materials.
- (ii) **Conductor loss:** Conductor losses are often called power loss. It is mainly due to the resistance of the conductor. Conductor also is also due to frequency which is known as skin effect. Skin effect is the tendency of the alternating current by which they tend to increase the current density near the surface more than that at the core i.e. the current appears to flow on the skin of the conductor.
- (iii) **Dielectric heating loss:** Dielectric heating loss is due to the potential difference between the two conductors of a transmission line. When air is the dielectric, loss is negligible. However, in case of solid conductors it increases with the frequency.

Standing Wave Ratio

A standing wave is the combination of these waves moving in opposite directions and with the same amplitude and frequency. Since the waves are "superimposed" on each other, interference is created and their energies (voltage or current) are either added or canceled out by each other. Standing waves are created when a transmission line does not terminate correctly. As a result, the traveling wave (also known as the incident wave) gets reflected -- completely or partially -- at the receiving end. Together, the incident and reflected waves give rise to standing waves along the line.

At some points along the transmission line, the two waves are in phase, so they add together, leading to maximum voltage and current. These points are known as the voltage or current maxima. At other points, the two waves are out of phase, so the resultant voltage and current will fall to the minimum (voltage or current minima). Ultimately, the amplitude of standing waves indicates the amount of reflection along the transmission line.

Standing wave ratio (SWR) is the ratio of the maximum magnitude or amplitude of a standing wave to its minimum magnitude. It indicates whether there is an impedance mismatch between the load and the internal impedance on a radio frequency (RF) transmission line, or waveguide. Such mismatches indicate that

there are standing waves along the line that can reduce its power transmission efficiency.

VSWR

Voltage standing wave ratio (VSWR) is defined as the ratio between transmitted and reflected voltage standing waves in a radio frequency (RF) electrical transmission system. It is a measure of how efficiently RF power is transmitted from the power source, through a transmission line, and into the load.

$$\text{VSWR} = |V^{\text{MAX}}|/|V^{\text{MIN}}|$$

Reflection Coefficient

Reflection coefficient is a parameter that describes how much of a wave is reflected by an impedance discontinuity in the transmission medium. It is equal to the ratio of the amplitude of the reflected wave to the incident wave, with each expressed as phasors.

Stub Matching

A stub is a piece of transmission line.

It is possible to connect sections of open or short circuited line called stub in shunt with the main line at some point to effect impedance matching. This is called stub matching.

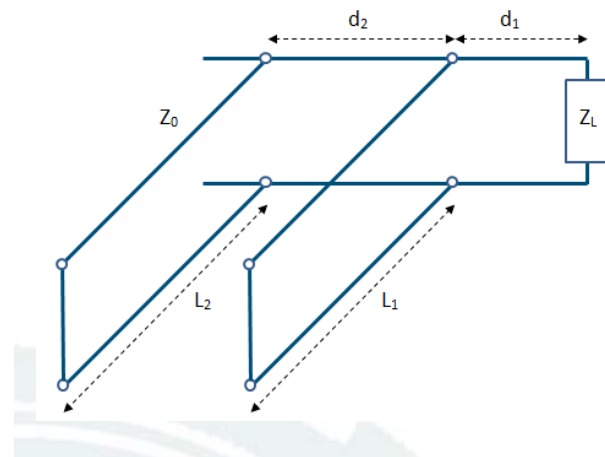
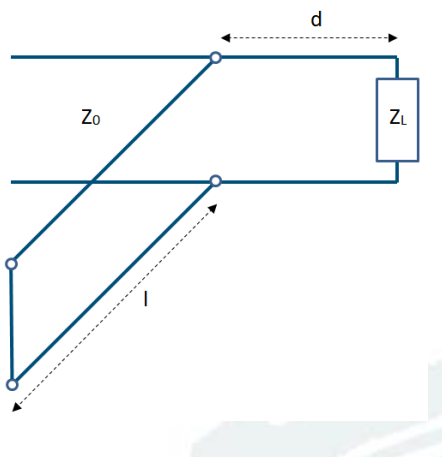
It has two advantages:

- (a) The length and characteristic impedance of the line remains unchanged.
- (b) Adjustable susceptance can be added in shunt with the transmission line.

Stub matching is of two types:

- (i) Single stub matching
- (ii) Double stub matching

Single stub matching:- in this method, a stub of certain fixed length is placed at some distance from the load. It is used only for a fixed frequency because for any change in frequency, the location of the stub has to be changed, which is not done.



Double stub matching: in this method, two stubs of variable length are fixed at certain positions. As the load changes, only the lengths of the stubs are adjusted to achieve matching. This is widely used in laboratory practice as a single frequency matching.

Unit 3

TELEVISION ENGINEERING

Aspect Ratio

The ratio between width to height of rectangle picture frame adopted in TV system is known as aspect ratio.

Aspect ratio = Width / height = 4/3 or 4: 3

Reasons for having this ratio is

1. Most of the objects are moving only in horizontal plane.
2. Our eye can see the movement of object comfortably only in horizontal plane than in vertical plane.
3. The frame size of motion picture already existing is having the aspect ratio of 4 : 3

Flicker

The sensation produced by incident light on the nerves of the eyes retina does not cease immediately. It persists for about 1/25th of a second (.062 Sec.) This storage characteristic is called as persistence of vision of eye. Flicker means if the scanning rate of picture is low, the time taken to move one frame to another frame will be high. This results in alternate bright and dark picture in the screen. This is called “Flicker”. To avoid flicker, the scanning rate of the picture should be increased i.e. 50 frames/Sec.

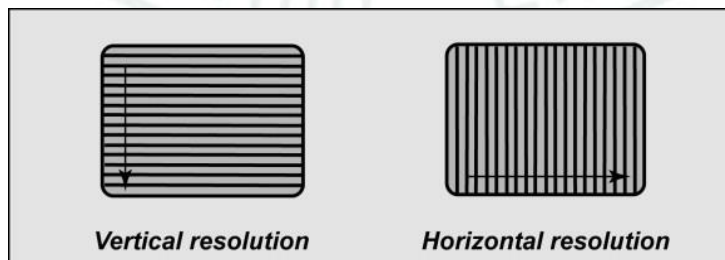
Resolution

The ability of the image reproducing system to resolve the fine details of the picture distinctly in both horizontal and vertical direction is called as “resolution”.

• **Vertical Resolution:** The ability to resolve and reproduce fine details of picture in vertical direction is called as Vertical resolution.

• **Horizontal Resolution :**

The ability of the system to resolve maximum number of picture elements along the scanning determines the horizontal resolution



Video Bandwidth

It is the range of frequencies between 0 and the highest frequency used to transmit a live television image.

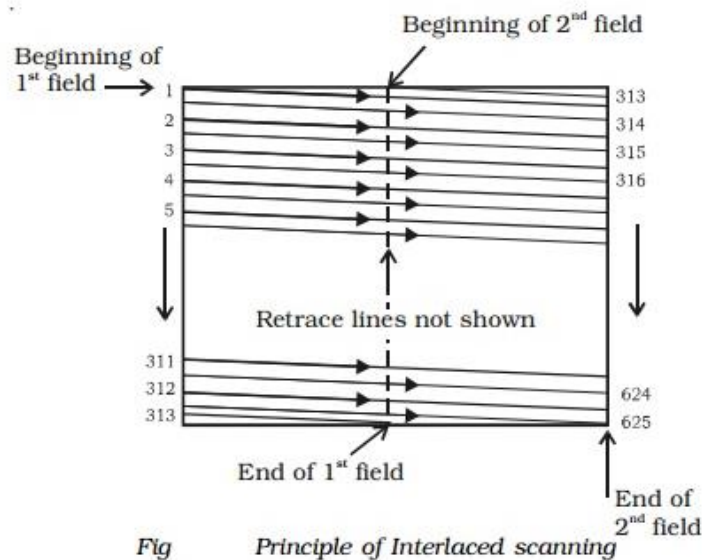
Scanning

Scanning is the process used to convert the optical into electrical signal. Fastest movement of electron beam on the image is called scanning.

Interlaced Scanning:

To reduce flicker, the vertical scanning is done 50 times per second in TV system. However only 25 frames are scanned per sec.

In interlaced scanning the 625 lines are grouped into two fields. They are called as even field and odd field. Each field contains 312.5 lines. Even field contains even numbered lines and odd field contains odd numbered lines. During first scanning line numbers 1, 3, 5 are scanned. During next scan, line numbers 2, 4, 6.....are scanned. That is alternate lines are scanned every time. So to cover each frame, scanning is done two times. Here the vertical rate of scanning is increased twice. So it will reduce flicker

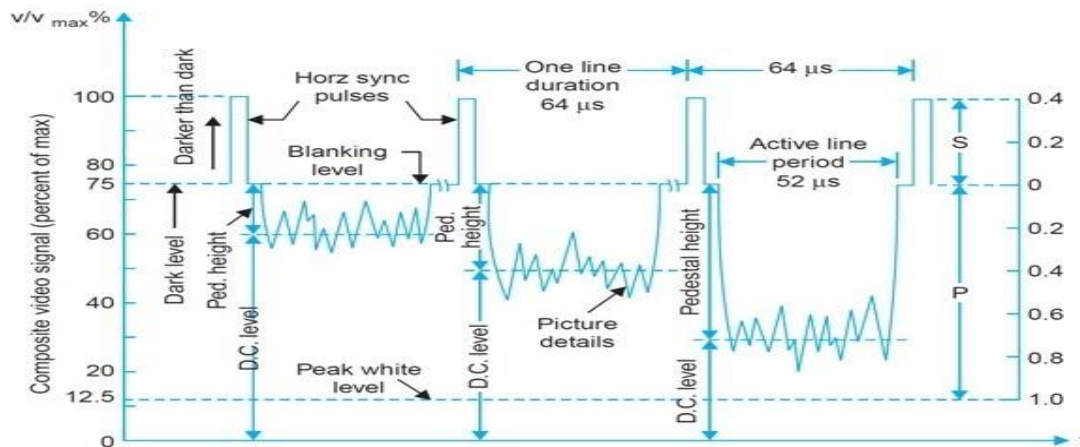


Composite Video Signal

CVS consists of,

- Camera signal corresponding to the picture to be transmitted.
- Blanking pulses to made the retrace invisible.

- Sync pulse to synchronize the transmitter and receiver.



A horizontal synchronizing pulse is needed at the end of each active line period whereas a vertical sync pulse is required after each field is scanned. The amplitude of both horizontal and vertical sync pulses is kept the same to obtain higher efficiency of picture signal transmission but their duration is chosen to be different for separating them at the receiver.

TV Transmitter

Television Camera:

Its function is to convert optical image of television scene into electrical signal by the scanning process.

Video Amplifier:

Video amplifier amplifies the video signal.

AM Modulating Amplifier

The video signals are amplified by the modulating amplifier to get the modulated signal.

Audio Amplifier

Audio amplifier amplifies the electrical form of audio signal from the microphone.

FM Modulating Amplifier:

Sound signal from audio amplifier is frequency modulated by FM Modulating amplifier.

FM Sound Transmitter:

FM modulated amplified signal is transmitted through this FM sound transmitter to transmitting antenna through the combining network.

Crystal Oscillator:

Crystal Oscillator generates the allotted picture carrier frequency.

RF Amplifier:

RF amplifier amplifies the picture carrier frequency generated by crystal oscillator to required level.

Power Amplifier:

Power amplifier varies according to the modulating signal from AM modulating amplifier.

Scanning and Synchronizing Circuit

Scanning is the process where picture elements are converted into corresponding varying electrical signal

Combining Network

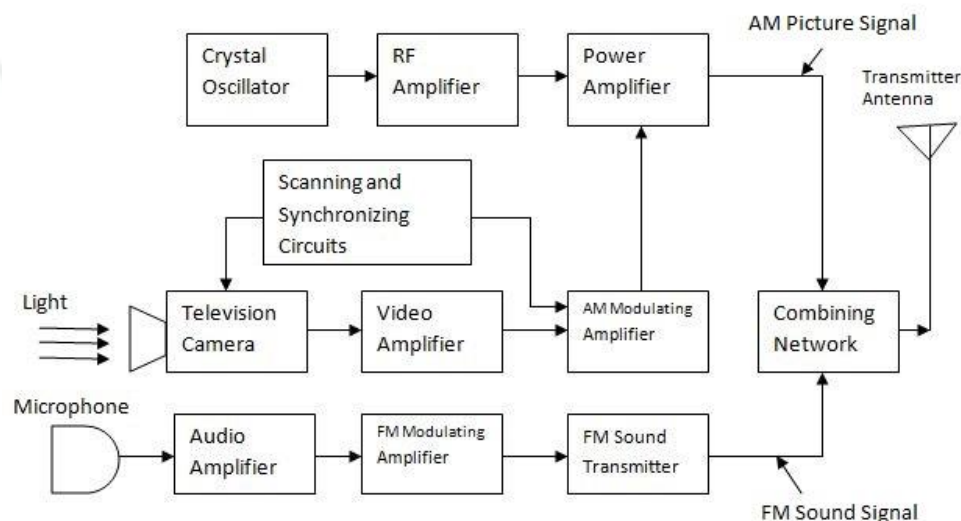
Combining network is used to isolate the AM picture and FM sound signal during transmission.

Transmitting Antenna:

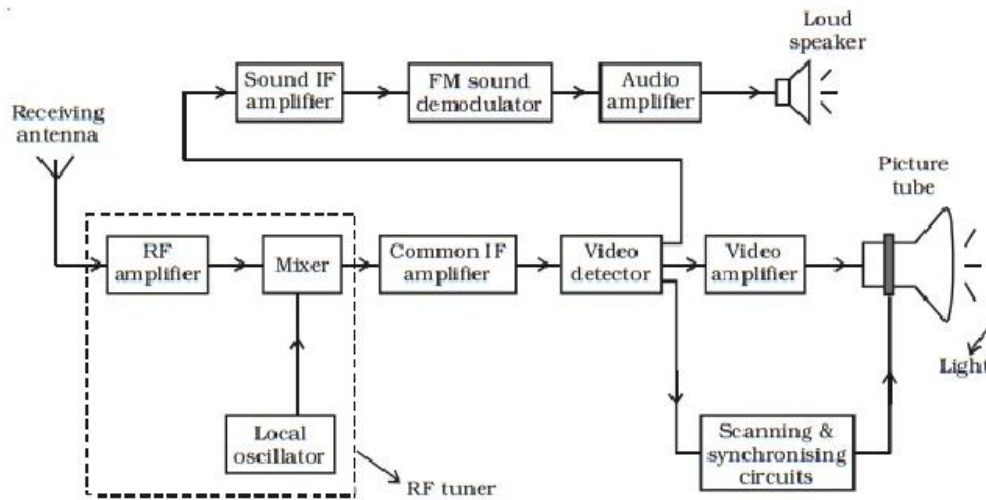
Transmitting antenna receives the AM picture signal and FM sound signal from combining network for radiation as electromagnetic waves.

Microphone:

Converts sound associated with picture being televised into proportionate electrical signal



Monochrome TV Receiver



RF Tuner:

RF Tuner selects the desired channel frequency band from the receiving antenna.

Receiver Antenna:

Receiver antenna intercepts the radiated RF signals and sends it to RF Tuner.

Common Amplifier:

There are 2 or 3 stages of IF amplifiers.

Video Detector:

Used to detect video signals coming from last stage of IF amplifiers.

Video Amplifier:

It amplifies the detected video signal to the level required.

Scanning and Synchronizing Circuit:

Scanning is the process where picture elements are converted into corresponding varying electrical signals.

Sound IF Amplifier:

Detected audio signal is separated and selected for its IF range and amplified.

FM Sound Demodulator:

FM Sound signal is demodulated in this stage.

Audio Amplifier:

FM demodulated audio signal is amplified to the required level to feed into the loud speaker.

Loud Speaker:

Loud Speaker converts FM demodulated amplifier signal associated with picture being televised into proportionate sound signal.

Picture Tube:

In picture tube the amplified video signal is converted back into picture elements.

Scanning:

Scanning is the process used to convert the optical into electrical signal. Fastest movement of electron beam on the image is called scanning

Need for Synchronization:

At any time the same co-ordinate will be scanned by the electron beam in both the camera tube and picture tube. Otherwise distorted picture will be seen on the screen. So synchronization between the transmitter and receiver is needed. For that we are using Sync pulses.

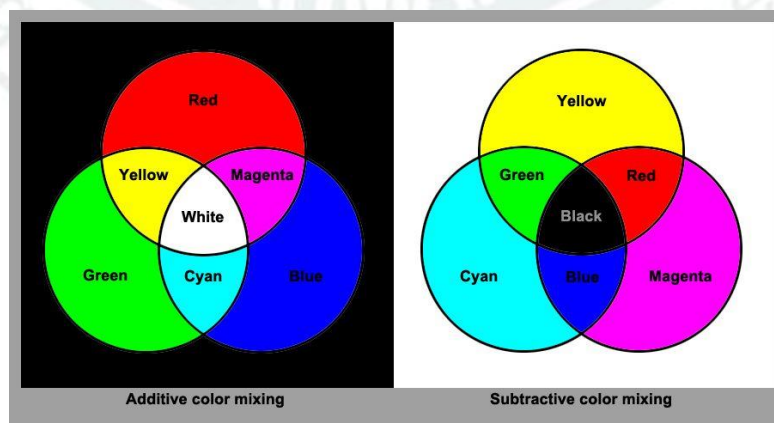
At the receiver side these pulses are identified, separated and used for triggering the oscillator circuit.

Color TV Signal

The three basic colors are called as primary colors. They are Red, Green and Blue. To get different color shading we have to mix primary colors. We have two types of mixing they are Additive Mixing and subtractive mixing.

Additive Mixing:

In this method two or three primary colors are mixed together to form a new color. By mixing primary colors with different intensities we can obtain all types of colors. By mixing 30% Red, 59% Green and 11% blue we can get white color.



Luminance

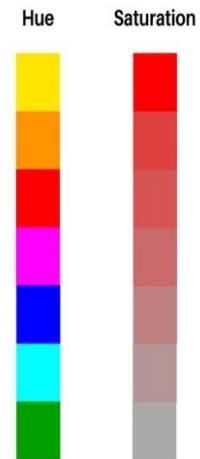
It is the amount of light intensity as perceived by the eye regardless of the color. It is also called as brightness signal, Y signal and white signal.

Hue

It is the prominent spectral color. E.g green leaf has a green hue

Saturation

It will indicate the spectral purify of color i.e. it will indicate how much white mixed with a particular color.



Chrominance

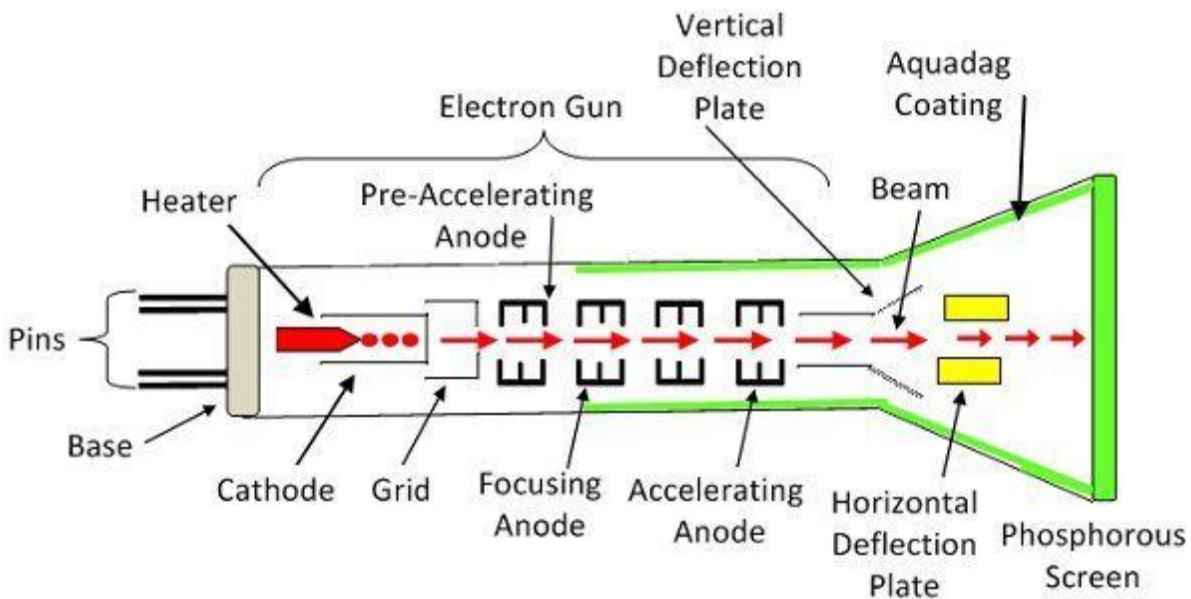
Hue and saturation together are called as chrominance or chroma signal.

Cathode Ray Tube TVs

The CRT is a display screen which produces images in the form of the video signal. It is a type of vacuum tube which displays images when the electron beam through electron guns are strikes on the phosphorescent surface. In other Words, the CRT generates the beams, accelerates it at high velocity and deflect it for creating the images on the phosphorous screen so that the beam becomes visible.

The working of CRT depends on the movement of electrons beams. The electron guns generate sharply focused electrons which are accelerated at high voltage. This high-velocity electron beam when strikes on the fluorescent screen creates luminous spot

After exiting from the electron gun, the beam passes through the pairs of electrostatic deflection plate. These plates deflected the beams when the voltage applied across it. The one pair of plate moves the beam upward and the second pair of plate moves the beam from one side to another. The horizontal and vertical movement of the electron are independent of each other, and hence the electron beam positioned anywhere on the screen. The working parts of a CRT are enclosed in a vacuum glass envelope so that the emitted electron can easily move freely from one end of the tube to the other.



Plasma Display Panel

Plasma Display Panel is composed of two parallel sheets of glass that enclose a mixture of discharge gases composed of helium, neon, and xenon. On the inner side of the glass, plates are ribs, which help keep the glass plates parallel.

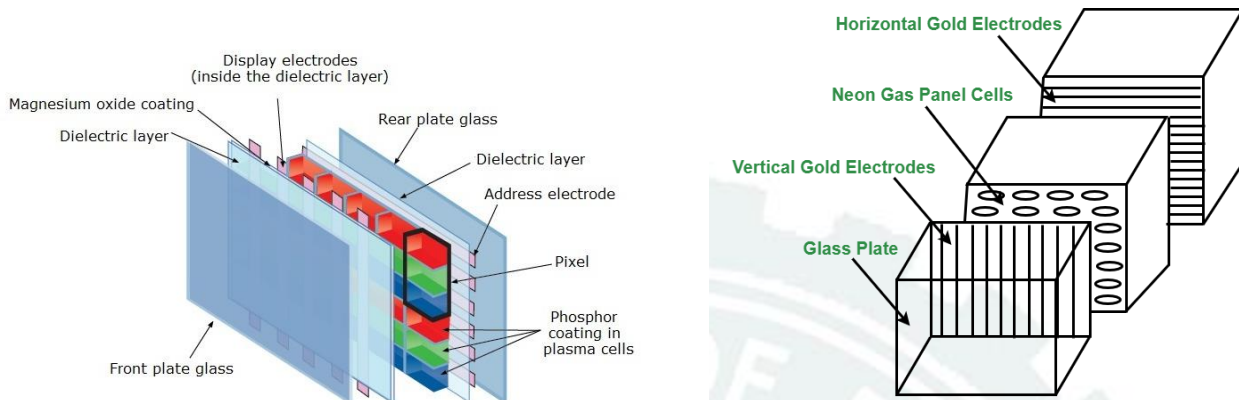
Groups of electrodes sit at right angles between the panes forming rectangular compartments, or cells, between the glass sheets. Phosphorus is embedded within each cell that individually emits red, green, or blue light and collectively creates a single color pixel.

Selectively applying voltages to the electrodes causes them to generate a discharge in the panel's dielectric layer and on its protective surface. This generates ultraviolet light that excites the phosphors, stimulating them to emit light.

Advantages of Plasma Display Panel :

- Plasma Display Panel are thin lightweight and take up less space than other displays which makes them easy to install anywhere.
- Plasma Display Panel offers uniform brightness.
- They show images without distortion and avoid problems such as misregistered colors and lack of focus.

- They also resist interference from magnetic fields, are free from distortion and are viewable from a wide angle.

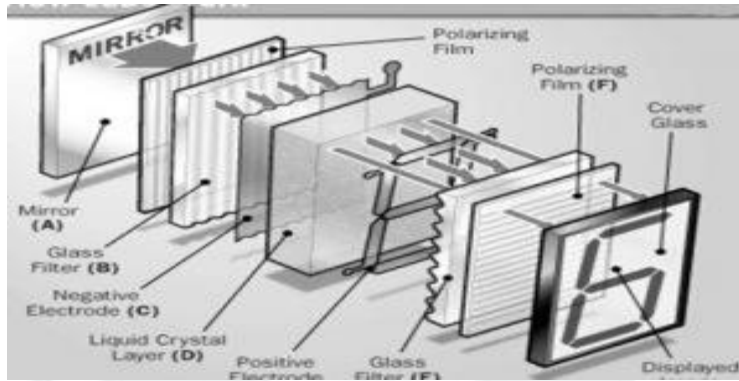


Liquid Crystal Display

It is a combination of two states of matter, the solid and the liquid. LCD uses a liquid crystal to produce a visible image. Liquid crystal displays are super-thin technology display screens that are generally used in laptop computer screens, TVs, cell phones, and portable video games. LCD's technologies allow displays to be much thinner when compared to a cathode ray tube (CRT) technology.

As mentioned above that we need to take two polarized glass pieces filter in the making of the liquid crystal. The glass which does not have a polarized film on the surface of it must be rubbed with a special polymer that will create microscopic grooves on the surface of the polarized glass filter. The grooves must be in the same direction as the polarized film.

Now we have to add a coating of pneumatic liquid phase crystal on one of the polarizing filters of the polarized glass. The microscopic channel causes the first layer molecule to align with filter orientation. When the right angle appears at the first layer piece, we should add a second piece of glass with the polarized film. The first filter will be naturally polarized as the light strikes it at the starting stage. Thus the light travels through each layer and guided to the next with the help of a molecule. The molecule tends to change its plane of vibration of the light to match its angle. When the light reaches the far end of the liquid crystal substance, it vibrates at the same angle as that of the final layer of the molecule vibrates. The light is allowed to enter into the device only if the second layer of the polarized glass matches with the final layer of the molecule.



How does LCDs work

The principle behind the LCDs is that when an electrical current is applied to the liquid crystal molecule, the molecule tends to untwist. This causes the angle of light which is passing through the molecule of the polarized glass and also causes a change in the angle of the top polarizing filter. As a result, a little light is allowed to pass the polarized glass through a particular area of the LCD.

Thus that particular area will become dark compared to others. The LCD works on the principle of blocking light. While constructing the LCDs, a reflected mirror is arranged at the back. An electrode plane is made of indium-tin-oxide which is kept on top and a polarized glass with a polarizing film is also added on the bottom of the device. The complete region of the LCD has to be enclosed by a common electrode and above it should be the liquid crystal matter.

Next comes the second piece of glass with an electrode in the form of the rectangle on the bottom and, on top, another polarizing film. It must be considered that both the pieces are kept at the right angles. When there is no current, the light passes through the front of the LCD it will be reflected by the mirror and bounced back. As the electrode is connected to a battery the current from it will cause the liquid crystals between the common-plane electrode and the electrode shaped like a rectangle to untwist. Thus the light is blocked from passing through. That particular rectangular area appears blank.

OLED(Organic LED)

OLED's are simple solid-state devices (more of an LED) comprised of very thin films of organic compounds in the electro-luminescent layer. These organic compounds have a special property of creating light when electricity is applied to it. The organic compounds are designed to be in between two electrodes. Out of

these one of the electrodes should be transparent. The result is a very bright and crispy display with power consumption lesser than the usual LCD and LED.

The components in an OLED differ according to the number of layers of the organic material. There is a basic single layer OLED, two layer and also three layer OLED's. As the number of layers increase the efficiency of the device also increases. The increase in layers also helps in injecting charges at the electrodes and thus helps in blocking a charge from being dumped after reaching the opposite electrode. Any type of OLED consists of the following components.

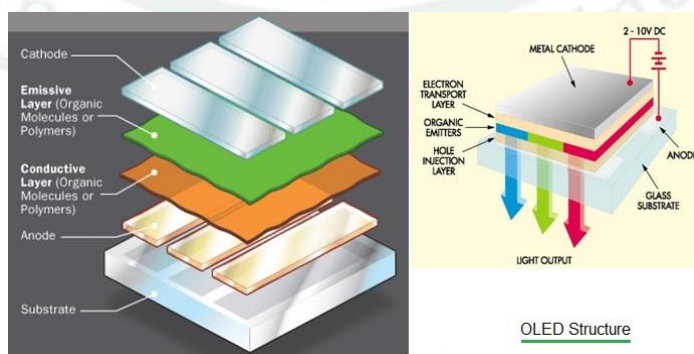
1. An emissive layer
2. A conducting layer
3. A substrate
4. Anode and cathode terminals.

As the emissive layer and the conducting layer is made up of organic molecules (both being different), OLED is considered to be an organic semi-conductor, and hence its name. The organic molecules have the property of conducting electricity and their conducting levels can be varied from that of an insulator to a conductor. The conducting layer is also an organic molecule, and the commonly used component is polyaniline.

The substrate most commonly used may be a plastic, foil or even glass.

The anode component should be transparent. Usually indium tin oxide is used. This material is transparent to visible light. It also has a great work function which helps in injecting holes into the different layers.

The cathode component depends on the type of OLED required. Even a transparent cathode can be used. Usually metals like calcium and aluminium used because they have lesser work functions than anodes which helps in injecting electrons into the different layers.



Working of an OLED

After the organic material has been applied to the substrate the real working of the OLED begins. The substrate is used to support the OLED. The anode is used to inject more holes when there is a path of current. The conducting layer is used to carry the holes from the anode. The cathode is used to produce electrons when current flows through its path. The emissive layer is the section where the light is produced. This layer is used to carry the electrons from the cathode.

First, the anode is kept positive w.r.t the cathode. Thus there occurs an electron flow from the cathode to the anode. This electron flow is captured by the emissive layer causing the anode to withdraw electrons from the conductive layer. Thus, there occurs a flow of holes in the conductive layer. As the process continues, the conductive layer becomes positively charged and the emissive layer becomes negatively charged.

A combination of the holes and electrons occur due to electrostatic forces. As the electrons are less mobile than the holes, the combination normally occurs very close to the emissive layer. This process produces light in the emissive region after there has been a drop in the energy levels of the electrons. The emissive layer got its name as the light produced in the emissive region has a frequency in the visible region. The colour of the light produced can be varied according to the type of organic molecule used for its process. To obtain colour displays, a number of organic layers are used. Another factor of the light produced is its intensity. If more current is applied to the OLED, the brighter the light appears. Take a look at the diagram given below. Now consider the process when the anode is negative w.r.t the cathode. This will not make the device work as there will not be any combination of the holes and electrons. The holes will move towards the anode and the electrons to the cathode.

QLED(Quantum Dot LED)

QLED TVs are televisions where quantum dot light-emitting diodes, or QLEDs, produce their display. Quantum dot light-emitting diodes are able to provide a clearer, brighter, and more efficient light compared to other forms of LEDs, such as traditional LEDs and organic LEDs.

QLED work using the principle of quantum dots. Quantum dots are incredibly small particles that are nanometres in size – that's 1 billion times smaller than 1 metre! In semiconductors, quantum dots are incredibly important,

because they can be used to produce electromagnetic radiation or light. In semiconductor physics, there are two important particles; electrons and holes ('positively-charged' electrons). When these particles combine, the combined energy is normally emitted as light, which is known as fluorescence. In different types of semiconductors, the energy can be small or large, which can result in light waves of varying wavelength, which can provide coherent, intense light. The first application of this technology was in laser physics, where scientists developed quantum dot lasers. However, engineers in the technology sector have developed this to use quantum dots to produce light that can act as a display itself, as opposed to using a standard LED as a backlight.

CATV systems and types and network

Cable television is a technical system for distribution of television by cable (coaxial, twisted pair or fiber optic) with potential for largest bandwidth and integrated return channel for interactive services. With the introduction of new technology, the CATV networks will have more active components than at present. Access of an end-user to services in the CATV network is realised by the help of access network. Access network must be enhanced to carry various multimedia services. There are several options to introduce fiber in the access network, to the cabinet/curb FTTCa/C with a last copper drop based on very high rate DSL on one hand and fibre to the building/fibre to the home FTTB/H on the other.

FTTCa/FFTC avoids the installation cost of fibre to the customer premises, but introduces a high exploitation cost, since network operator personnel will have to travel to the cabinet or curb unit for every alteration or maintenance action. Moreover, the process by which the operator becomes entitled to site a cabinet or curb unit in suitable position is complex, and powering will require a large investment. Because of these reasons it is assumed here that network operators will make the strategic choice of introducing fibre directly to the customer premises with FTTB/H. Residential broadband access network technology based on Asynchronous Transfer Mode (ATM) will soon reach commercial availability. The capabilities provided by ATM access network promise integrated services bandwidth available in excess of those provided by traditional networks. Other services such as desktop video teleconferencing and enhanced server-based application support can be added as part of future evolution of the network.

The Structure Of the Interactive Network :

At present, the great importance of CATV is in the best transferring of information, mainly in association with satellite transmission of TV and R signals. It indicates that possibilities of CATV are much bigger and they reach to other spheres. Therefore, it is necessary to come nearer CATV from another point of view. This point of view is the information approach, and new conception of CATV doesn't deal with system, which transfers the TV and R signals, but with the system transferring whatever information to arbitrary direction. By this new approach, primary sense of CATV fades and the network becomes the universal data network, where it is possible to create and realise almost arbitrary services. In interactive CATV, there are three levels, similarly as in distribution network.

In a city or another larger region, there is a set of mutually connected central nodes, which create the primary network. Each of them serves respective part of the city e.g. and habitation, a street or likewise. Interactive CATVs can be independent networks, which exist simultaneously with other ones. But, it is economically better and efficiently when the primary network of CATV is an access network of certain great national network (WAN, MAN, B-ISDN etc...). Communication in the primary network must be enough wideband and of larger distance than in primary network of distribution CATV. Therefore, there is directly offered using of optical fibers with standard transfer rates here, e.g. 155Mb/s, 622Mb/s or much.

The central node is connected to set of distribution nodes and so it is created the secondary network in star form. Each of distribution nodes covers set of flats or offices in small area, e.g. in multi-storey building. Wideband connection through optical fibres uses standard transfer rates, e.g. 155Mb/s, 622Mb/s or much. Physical layer is the same as in primary network.

Distribution node creates its own tertiary network in star form, which connects user's terminal devices. The transfer medium is coax cable, therefore for this subnetwork is possible use the existing lines which were created for distribution CATVs. Tertiary network is possible alternative to present LAN with transfer rate 10Mb/s. One connection in tertiary network contains several Ethernet channels. So, the connection of terminal user is indeed wideband.

Digital TV Signal

Digital TV is the transmission of television signals using digital rather than analog methods. A digital TV signal works like a computer transmission. A burst of binary data representing the broadcast is sent out and decoders on the other end read the data and create a picture based on the information. Interference in the analog transmission results in ghosting or static, but the picture is visible even when almost totally degraded.

Transmission of Digital TV signal

To transmit the digital signal it uses QAM or VSB modulation techniques. The digital TV signal is produced by a digital processor which converts the electrical analog signal from the television camera by sampling the electrical voltage many times per second to numbers.

These numbers in turn are encoded as sets of 0's and 1's . A digital camera performs these operations automatically. Digital TV signals are transmitted via copper cables, fiber optic cables or for over digital microwave radio for satellites relay links. Fiber optic cables can give very high capacities over larger distances a span reach of tens of km at 2 to 100 Mbit/s depending on the wavelength and the type of transmitter and detector.

Digital TV Receiver

It is a set top box that permits the reception of digital television. All the main hardware components of the receiver are connected to the system board. The tuner in the box receives a digital signal from a cable, a satellite, or terrestrial network and isolates a particular band.

The signal is then forwarded to a demodulator and converted to binary format. Once in binary format, the demodulator will check for error and forward the binary signal to a demultiplexer that will extract audio, video and data from the binary stream. Once the demultiplexer has finished with the signal the decoders will transform the digital bits into a format suitable for viewing on the television set.

Unit 4

MICROWAVE ENGINEERING

Microwaves propagate through microwave circuits, components and devices, which act as a part of Microwave transmission lines, broadly called as Waveguides. A hollow metallic tube of uniform cross-section for transmitting electromagnetic waves by successive reflections from the inner walls of the tube is called as a **Waveguide**.

A waveguide is generally preferred in microwave communications. Waveguide is a special form of transmission line, which is a hollow metal tube. Unlike a transmission line, a waveguide has no center conductor.

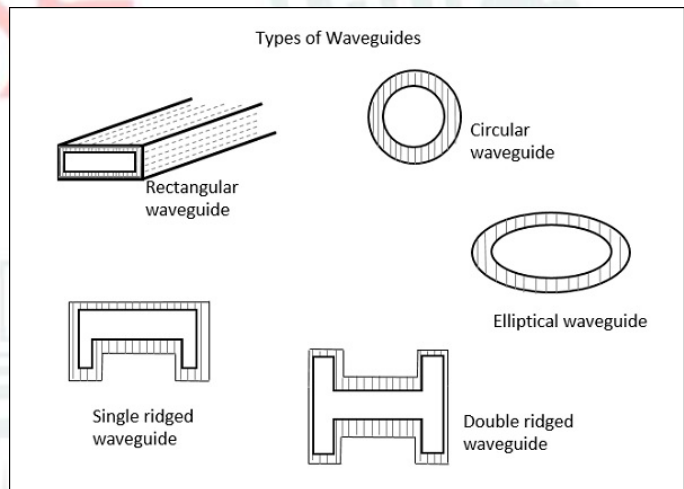
The main characteristics of a Waveguide are –

- The tube wall provides distributed inductance.
- The empty space between the tube walls provide distributed capacitance.
- These are bulky and expensive.

Types of Waveguides

There are five types of waveguides.

- Rectangular waveguide
- Circular waveguide
- Elliptical waveguide
- Single-ridged waveguide
- Double-ridged waveguide



Rectangular Waveguide

In electromagnetics, a waveguide confines electromagnetic signals within the structure, preventing spreading, losses, and signal transmission from one point to another. Usually, a basic waveguide can be constructed from a hollow conducting tube. If the conducting tube has a rectangular cross-section, then it forms the rectangular waveguide. They consist of a hollow metallic structure with a rectangular cross-section. A rectangular waveguide is usually constructed with a

length of $a > b$, where b is the breadth of the rectangle. A common trend for the dimension of a rectangular waveguide is $a=2b$.

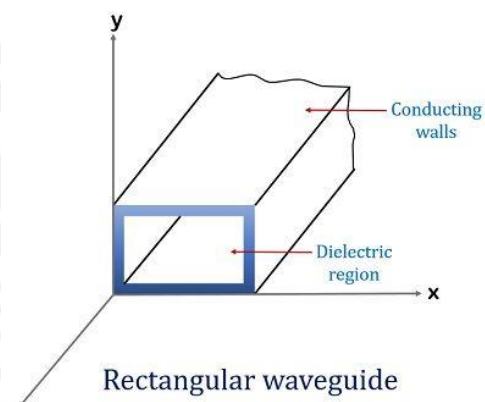
In a typical system, there may be an antenna at one end of a waveguide and a receiver or transmitter at the other end. The antenna generates electromagnetic waves, which travel down the waveguide to be eventually received by the load. The walls of the guide are conductors, and therefore reflections from them take place. It is of the utmost importance to realize that conduction of energy takes place not through the walls, whose function is only to confine this energy, but through the dielectric filling the waveguide, which is usually air. In discussing the behavior and properties of waveguides, it is necessary to speak of electric and magnetic fields, as in wave propagation, instead of voltages and currents, as in transmission lines. This is the only possible approach, but it does make the behavior of waveguides more complex to grasp.

Modes of Propagation

Rectangular waveguides are extensively used in radars, couplers, isolators, and attenuators for signal transmission. When electromagnetic waves are transmitted longitudinally through a rectangular waveguide, they are reflected from the conducting walls. The total reflection inside the rectangular waveguide results in either an electric field or magnetic field component in the direction of the propagation. There is no TEM mode in rectangular waveguides. The modes of propagation in a hollow rectangular waveguide with only one conductor are either TE or TM modes.

Advantages of Rectangular Waveguides

- Wide frequency bandwidth for single-mode propagation
- Low attenuation
- Excellent mode stability for fundamental propagation modes



Modes of propagation

A wave has both electric and magnetic fields. All transverse components of electric and magnetic fields are determined from the axial components of electric

and magnetic field, in the z direction. This allows mode formations, such as TE, TM, TEM and Hybrid in microwaves. Let us have a look at the types of modes.

TEM (Transverse Electromagnetic Mode)

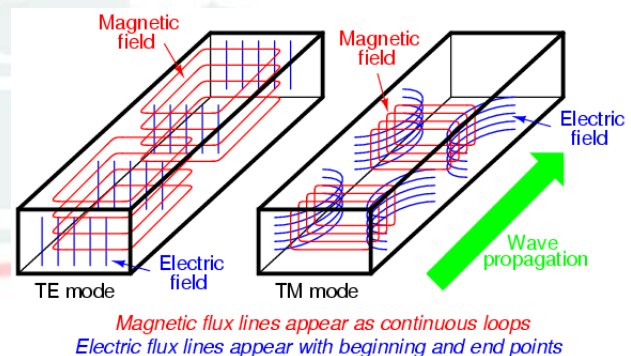
In this mode, both the electric and magnetic fields are purely transverse to the direction of propagation. There are no components in 'Z' direction.

$$E_z = 0 \text{ and } H_z = 0$$

TE (Transverse Electric Mode)

In this mode, the electric field is purely transverse to the direction of propagation, whereas the magnetic field is not.

$$E_z = 0 \text{ and } H_z \neq 0$$



TM (Transverse Magnetic Mode)

In this mode, the magnetic field is purely transverse to the direction of propagation, whereas the electric field is not.

$$E_z \neq 0 \text{ and } H_z = 0$$

HE (Hybrid Mode)

In this mode, neither the electric nor the magnetic field is purely transverse to the direction of propagation.

$$E_z \neq 0 \text{ and } H_z \neq 0$$

Multi conductor lines normally support TEM mode of propagation, as the theory of transmission lines is applicable to only those system of conductors that have a go and return path, i.e., those which can support a TEM wave. Waveguides are single conductor lines that allow TE and TM modes but not TEM mode. Open conductor guides support Hybrid waves. The types of transmission lines are discussed in the next chapter.

Circular Waveguide

A waveguide is a hollow metal tube (rectangular or circular in cross section) that transmits electromagnetic energy from one place to another. A waveguide with a circular cross-section is called as **Circular Waveguide**. It supports both

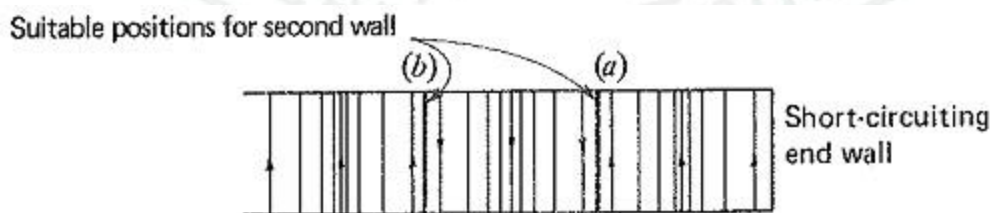
transverse electric (TE) and transverse magnetic (TM) modes. TE₁₁ is the dominant mode in a circular waveguide i.e., a signal in this mode propagates with the minimum degradation. The circular waveguide is easier to manufacture than rectangular waveguides and is relatively easy to install. It is usually used to connect a horn antenna with a reflector in tracking radars and for long-distance waveguide transmission above 10 GHz.

Operational Cavity Resonator

A cavity resonator is a piece of waveguide closed off at both ends with metallic planes. Where propagation in the longitudinal direction took place in the waveguide, standing waves exist in the resonator, and oscillations can take place if the resonator is suitably excited. Various aspects of cavity resonators will now be considered. Waveguides are used at the highest frequencies to transmit power and signals. Similarly, cavity resonators are employed as tuned circuits at such frequencies. Their operation follows directly from that of waveguides.

Operation:

Until now, waveguides have been considered from the point of view of standing waves between the side walls and traveling waves in the longitudinal direction. If conducting end walls are placed in the waveguide, then standing waves, or oscillations, will take place if a source is located between the walls. This assumes that the distance between the end walls is $n\lambda_p/2$, where n is any integer. The first wall ensures standing waves, and placement of the second wall permits oscillations, provided that the second wall is placed so that the pattern due to the first wall is left undisturbed. Thus, if the second wall is $\lambda_p/2$ away from the first, as in Figure 10-36, oscillations between the two walls will take place. They will then continue until all the applied energy is dissipated, or indefinitely if energy is constantly supplied. This is identical to the behavior of an LC tuned circuit.



It is thus seen that any space enclosed by conducting walls must have one (or more) frequency at which the conditions just described are fulfilled. In other words, any such enclosed space must have at least one resonant frequency. Indeed, the completely enclosed waveguide has become a cavity resonator with its own system of modes, and therefore resonant frequencies. The TE and TM mode-numbering system breaks down unless the cavity has a very simple shape, and it is preferable to speak of the resonant frequency rather than mode.

Each cavity resonator has an infinite number of resonant frequencies. This can be appreciated if we consider that with the resonator of Figure 10-35 oscillations would have been obtained at twice the frequency, because every distance would now be λ_p , instead of $\lambda_p/2$. Several other resonant frequency series will also be present, based on other modes of propagation, all permitting oscillations to take place within the cavity. Naturally such behavior is not really desired in a resonator, but it need not be especially harmful. The fact that the cavity can oscillate at several frequencies does not mean that it will. Such frequencies are not generated spontaneously; they must be fed in.

Applications:

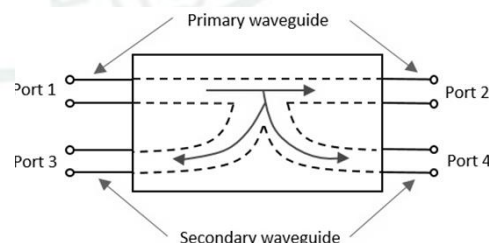
Cavity resonators are employed for much the same purposes as tuned LC circuits or resonant transmission lines, but naturally at much higher frequencies since they have the same overall frequency coverage as waveguides. They may be input or output tuned circuits of amplifiers, tuned circuits of oscillators, or resonant circuits used for filtering or in conjunction with mixers. In addition, they can be given shapes that make them integral parts of microwave amplifying and oscillating devices, so that almost all such devices use them.

Directional Coupler

A **Directional coupler** is a device that samples a small amount of Microwave power for measurement purposes. The power measurements include incident power, reflected power, VSWR values, etc.

Directional Coupler is a 4-port waveguide junction consisting of a primary main waveguide and a secondary auxiliary waveguide. The following figure shows the image of a directional coupler.

Directional coupler is used to couple the

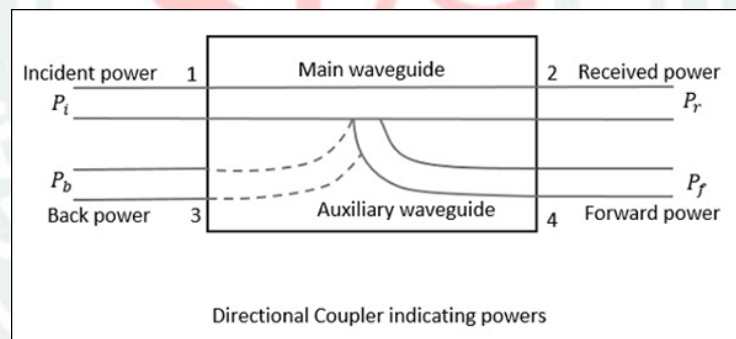


Microwave power which may be unidirectional or bi-directional.

Properties of Directional Couplers

- All the terminations are matched to the ports.
- When the power travels from Port 1 to Port 2, some portion of it gets coupled to Port 4 but not to Port 3.
- As it is also a bi-directional coupler, when the power travels from Port 2 to Port 1, some portion of it gets coupled to Port 3 but not to Port 4.
- If the power is incident through Port 3, a portion of it is coupled to Port 2, but not to Port 1.
- If the power is incident through Port 4, a portion of it is coupled to Port 1, but not to Port 2.
- Port 1 and 3 are decoupled as are Port 2 and Port 4.

Ideally, the output of Port 3 should be zero. However, practically, a small amount of power called **back power** is observed at Port 3. The following figure indicates the power flow in a directional coupler.



P_i = Incident power at Port 1
 P_r = Received power at Port 2
 P_f = Forward coupled power at Port 4
 P_b = Back power at Port 3

Following are the parameters used to define the performance of a directional coupler.

Coupling Factor (C)

The Coupling factor of a directional coupler is the ratio of incident power to the forward power, measured in dB.

Directivity (D)

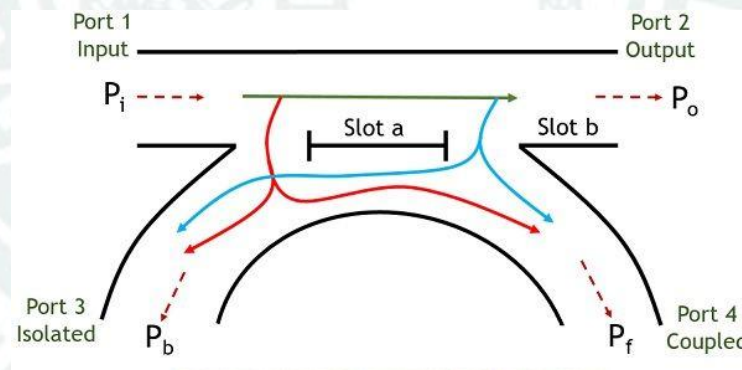
The Directivity of a directional coupler is the ratio of forward power to the back power, measured in dB.

Isolation

It defines the directive properties of a directional coupler. It is the ratio of incident power to the back power, measured in dB.

Two Hole Directional Coupler

This is a directional coupler with same main and auxiliary waveguides, but with two small holes that are common between them. These holes are $\lambda_g/4$ distance apart where λ_g is the guide wavelength. A two-hole directional coupler is designed to meet the ideal requirement of directional coupler, which is to avoid back power. Some of the power while travelling between Port 1 and Port 2, escapes through the holes 1 and 2. The magnitude of the power depends upon the dimensions of the holes. This leakage power at both the holes are in phase at hole 2, adding up the power contributing to the forward power P_f . However, it is out of phase at hole 1, cancelling each other and preventing the back power to occur. Hence, the directivity of a directional coupler improves.



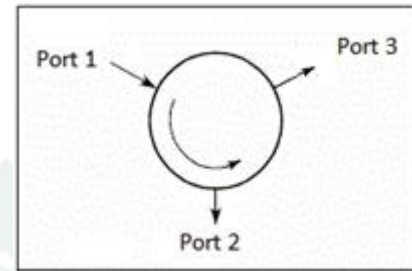
Circulators

RF circulators are used in many RF applications acting as a duplexer, allowing both transmit and receive functions to occur simultaneously, they are widely used in RF design applications including radar systems and a variety of professional radio communications systems. RF circulators are receive their name because they transfer power from one port to the next, circulating it from say,

entering at port one to output at port two, and entering at port two to exit at port three.

How an RF circulator works

The connections to RF circulators are normally called ports, and in addition to this they are normally numbered as 1, 2, 3, etc. The RF circulator gains its name because it circulates the power entering one port only to the next one. A signal applied to port 1 will be passed to port 2: a signal input to port 2 will pass to port 3, but not back to port 1. An input to port 3 will pass to port 1, but not in reverse to port 2.



The ideal circulator will transfer all the power from one port to the required port, and none to any other. However in reality, there is always some attenuation in the transfer path, and some signal always leaks onto the ports that should be isolated. The key RF circuit design challenge for these devices is to ensure the optimum transfer and isolation occurs.

Circulators may use strip-line printed circuit board technology (but normally using very low loss dielectric or PCB materials) and be contained within metal boxes with connectors or other connections to the outside world - some even use surface mount technology. Other circulators may be based within waveguides and these can be used in RF system design applications that incorporate waveguide technology. The type of interface and technology that are need for any given instance will depend upon the RF circuit design for the application. In terms of their operation there most RF circulators are based around the use of ferromagnetic material. There are two main types:

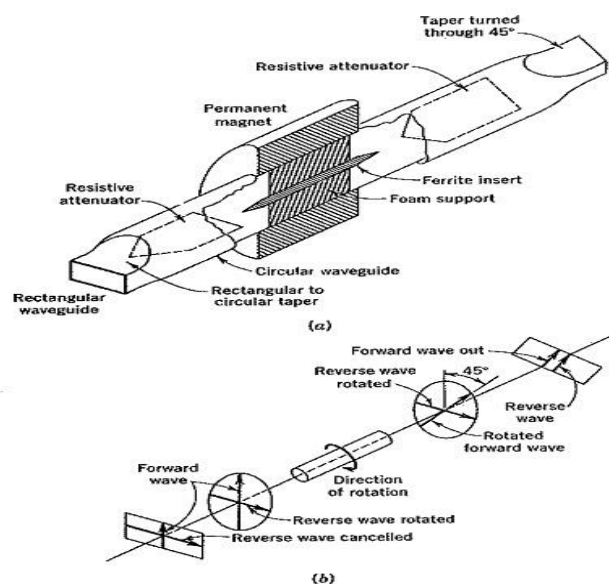
- **Three port circulators:** 3-port "Y-junction" circulators based on cancellation of waves propagating over two different paths near a magnetised material. Waveguide circulators may be of either type, while more compact devices based on strip-line are of the 3-port type.
- **Four port circulators:** 4-port waveguide circulators based on Faraday rotation of waves propagating in a magnetised material. Using this technology, they are able to route the RF signals to four ports.

Isolators

An *isolator* is a nonreciprocal transmission device that is used to isolate one component from reflections of other components in the transmission line. An ideal isolator completely absorbs the power for propagation in one direction and provides lossless transmission in the opposite direction.

The isolator consists of a piece of circular waveguide carrying the $TE_{1,1}$ mode, with transitions to a standard rectangular guide and $TE_{1,0}$ mode at both ends (the output end transition being twisted through 45°). A thin “pencil” of ferrite is located inside the circular guide, supported by polyfoam, and the waveguide is surrounded by a permanent magnet which generates a magnetic field in the ferrite that is generally about 160 A/m. A typical practical X-band (8.0 to 12.4 GHz) device may have a length of 25 mm and a weight of 100 g without the transitions.

Because the dc magnetic field (well below that required for resonance) is applied, a wave passing through the ferrite in the forward direction will have its plane of polarization shifted clockwise (through 45° in practical Waveguide Isolator and Circulators) by the time it reaches the output end. This wave is then passed through the suitably rotated output transition, and it emerges with an **insertion loss** (attenuation in the forward direction) between 0.5 and 1 dB in practice. It has not been affected by either of the resistive vanes because they are at right angles to the plane of its electric field; this is shown in Figure 10-40b.



A wave that tries to propagate through the isolator in the reverse direction is also rotated clockwise, because the direction of the Faraday rotation depends only on the dc magnetic field. Thus, when the wave emerges into the input transition, not only is it absorbed by the resistive vane, but also it cannot propagate in the input rectangular waveguide because of its dimensions. This situation is shown in Figure 10-40b. It results in the returned wave being attenuated by 20 to 30 dB in practice (this reverse attenuation of an isolator is called its isolation). Such a practical Waveguide Isolator and Circulators will have an SWR not exceeding 1.4, with values as low as 1.1, which is sometimes obtainable, and a bandwidth between 5 and 30 percent of the center frequency.

This type of isolator is limited in its peak power-handling ability to about 2 kW, , because of nonlinearities in the ferrite resulting in the phase shift departing from the ideal 45° . However, it has a very wide range of applications in the low-power field, since most microwave amplifiers and oscillators have output powers considerably lower than 2 kW.

Two Cavity Klystron

Klystrons are a special type of vacuum tubes that find applications as amplifiers and oscillators at microwave frequencies. Its principle of operation is velocity modulation. Thus the device used for amplifying microwave signals is known as Two-cavity Klystron.

Operating Principle of Two-cavity Klystron

As we have already discussed in the introduction that Klystron is based on the principle of velocity modulation. Thus two-cavity klystron amplifier utilizes the kinetic energy of moving electron beam for signal amplification. The variation in the velocity of electrons while moving inside the tube is known as velocity modulation. This velocity modulation permits bunching of electrons while propagation. So, the combined energy of bunched electrons is transferred at the output thereby providing an amplified signal.

The RF signal to be amplified is provided at the buncher cavity. The electron gun comprises cathode, heating element and anode. The electron beam is produced by the cathode by making use of a heating element and the high positive potential at the anode provides the required acceleration to the electron beam initially. The region between two cavities is known as drift space. To allow focussed propagation of electron beam inside the tube an external electromagnetic winding is used that generates a longitudinal magnetic field. This is done in order to prevent the spreading of the beam inside the tube. The amplified RF signal is achieved at the catcher cavity. Also, a collector is present near the second cavity that collects the electron bunch.

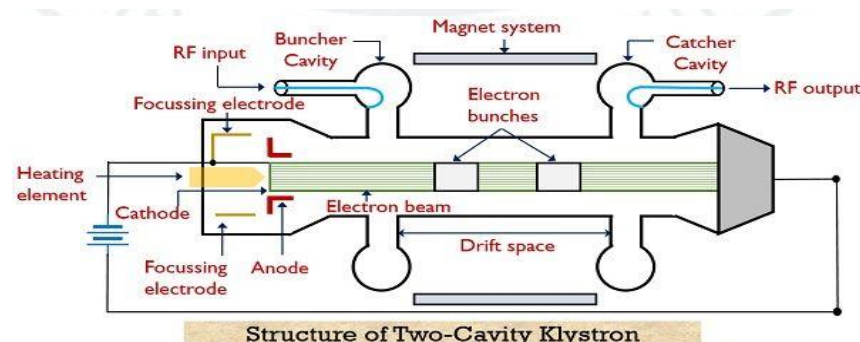
Working of Two-cavity Klystron Amplifier

Initially, electrons are emitted from the electron gun and the anode present in the structure provides the desired acceleration to the beam.

- In the absence of any RF input, the electron will tend to move with their respective uniform velocities to reach the catcher cavity and gets collected at the collector.
- But when external RF signal is applied at the input of the buncher cavity then this causes the generation of a local electric field inside the tube.

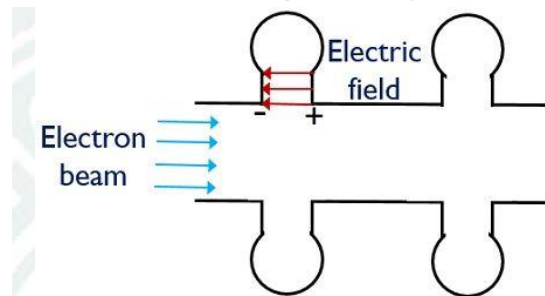
This electric field causes the bunching of electrons as the field applies acceleration and deceleration to the moving electron, according to the polarity of the signal by which the field is generated.

Basically, the reason for causing acceleration and deceleration is that when the direction of movement of an electron is opposite to the direction of the field, then, in this case, the electrons experience a decrease in their moving velocity. However, if the generated electric field and the direction of movement of the electron are the same then, in this case, the electrons experience increase in the velocity of their movement.

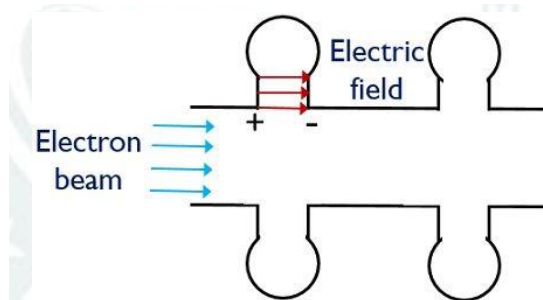


Let us now understand in detail how this increase and decrease in velocity causes bunching of electrons:

When the negative half of the RF signal is provided as input to the buncher cavity then the moving electrons experience a repulsive force due to the presence of a negative charge at the entering plate of the buncher cavity. Or we can say, that due to the negative half of the input the generated field will be in a direction opposite to the direction of the movement of electrons. So, because of the opposition offered by the field, the moving velocity of electrons gets reduced.



Further when the positive half of the RF signal is provided then the positive potential at the first plate of the cavity applies attractive force to the moving electrons. More simply, for the positive half cycle of input, the generated electric field will be in a direction similar to the direction of electron movement.



Thus, combinely when we consider both the cases then the electrons that were emitted earlier by the gun will be decelerated. While the electrons emitted later will be accelerated. Thus all the electrons while moving with different velocities get bunched in the drift space. This change in the velocity of electrons while moving due to RF input is known as velocity modulation. Once the electron bunching is done then the catcher cavity present at another end of the tube absorbs the beam energy.

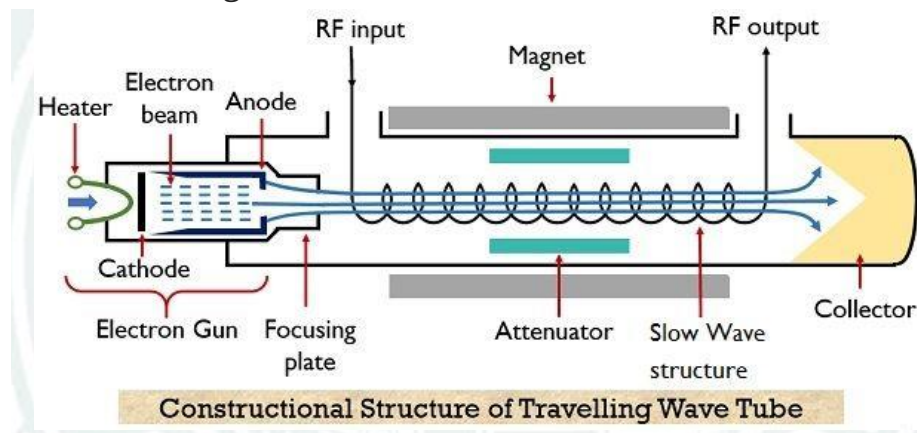
It is to be noted that to find the position of the catcher cavity transit time of the bunches must be considered. This is so because the catcher cavity must be present at a sufficient distance from the buncher cavity so that bunching can be

attained in the drift space. Further, once the energy is transferred to the catcher cavity then electrons (now with low energy) gets collected at the collector.

Travelling WaveTube

Travelling wave tubes are abbreviated as **TWT**. It is majorly used in the amplification of RF signals. Basically a travelling wave tube is nothing but an elongated vacuum tube that allows the movement of electron beam inside it by the action of applied RF input.

The movement of an electron inside the tube permits the amplification of applied RF input. As it offers amplification to a wide range of frequency thus is considered more advantageous for microwave applications than other tubes. It offers average power gain of around **60 dB**. The output power lies in the range of few watts to several megawatts.



As we can see that the helical travelling wave tube consists of an **electron gun** and a **slow-wave structure**. The electron gun produces a narrow beam of the electron. A focusing plate is used that focuses the electron beam inside the tube. A positive potential is provided to the coil (helix) with respect to the cathode terminal. While the collector is more positive than the coil (helix). In order to restrict beam spreading inside the tube. A dc magnetic field is applied between the travelling path by the help of magnets.

The signal which is needed to be amplified is provided at one of the ends of the helix, present adjacent to the electron gun. While the amplified signal is achieved at the opposite end of the helix.

Working of Travelling Wave Tube

Till now we have discussed the complete constructional structure of TWT. Let us now understand how the signal gets amplified while travelling inside the tube. The applied RF signal produces an electric field inside the tube. Due to the applied positive half, the moving electron beam experiences accelerative force. However, the negative half of the input applies a de-accelerative force on the moving electrons.

This is said to be velocity modulation because the electrons of the beam are experiencing different velocity inside the tube. However, the slowly travelling wave inside the tube exhibits continuous interaction with the electron beam.

Due to the continuous interaction, the electrons moving with high velocity transfer their energy to the wave inside the tube and thus slow down. So with the rise in the amplitude of the wave, the velocity of electrons reduces and this causes bunching of electrons inside the tube.

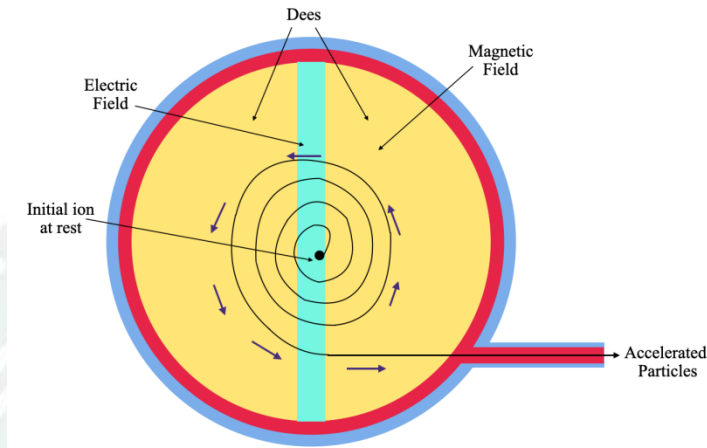
The growing amplitude of the wave resultantly causes more bunching of electrons while reaching the end from the beginning. Thereby causing further amplification of the RF wave inside the tube. More specifically we can say that forward progression of the field along the axis of the tube gives rise to amplification of the RF wave. Thus at the end of the tube an amplified signal is achieved. The positive potential provided at the other end causes collection of electron bunch at the collector. The magnetic field inside the tube restricts the spreading of the beam as the electrons possess repulsive nature.

However, as the TWT is a bidirectional device. Therefore, the reflected signal causes oscillations inside the tube. But as we have already discussed earlier that the presence of attenuators reduces the generation of oscillations due to reflected backwave. Sometimes despite using attenuators, internal impedance terminals are used that puts less lossy effects on the forward signal.

Cyclotron

A cyclotron is a machine that accelerates charged particles or ions to high energies. It was invented to investigate the nuclear structure by E.O Lawrence and M.S Livingston in 1934. Both electric and magnetic fields are used in the cyclotron to increase the energy of the charged particles. As both the fields are perpendicular to each other, they are called cross fields.

In a cyclotron, charged particles accelerate outwards from the centre along a spiral path. These particles are held to a spiral trajectory by a static magnetic field and accelerated by a rapidly varying electric field.



Working Principle of Cyclotron

- A cyclotron accelerates a charged particle beam using a high frequency alternating voltage which is applied between two hollow “D”-shaped sheet metal electrodes known as the “dees” inside a vacuum chamber.
- The dees are placed face to face with a narrow gap between them, creating a cylindrical space within them for particles to move. Particles are injected into the centre of this space.
- Dees are located between the poles of electromagnet which applies a static magnetic field B perpendicular to the electrode plane.
- The magnetic field causes the path of the particle to bend in a circle due to the Lorentz force perpendicular to their direction of motion.
- An alternating voltage of several thousand volts is applied between the dees. The voltage creates an oscillating electric field in the gap between the dees that accelerates the particles.
- The frequency of the voltage is set so that particles make one circuit during a single cycle of the voltage. To achieve this condition, the frequency must be set to the particle’s cyclotron frequency.

Gunn Diode

A Gunn Diode is considered as a type of diode even though it does not contain any typical PN diode junction like the other diodes, but it consists of two electrodes. This diode is also called as a Transferred Electronic Device. This diode

is a negative differential resistance device, which is frequently used as a low-power oscillator to generate microwaves. It consists of only N-type semiconductor in which electrons are the majority charge carriers. To generate short radio waves such as microwaves, it utilizes the Gunn Effect. The Gunn Effect can be defined as generation of microwave power (power with microwave frequencies of around a few GHz) whenever the voltage applied to a semiconductor device exceeds the critical voltage value or threshold voltage value.

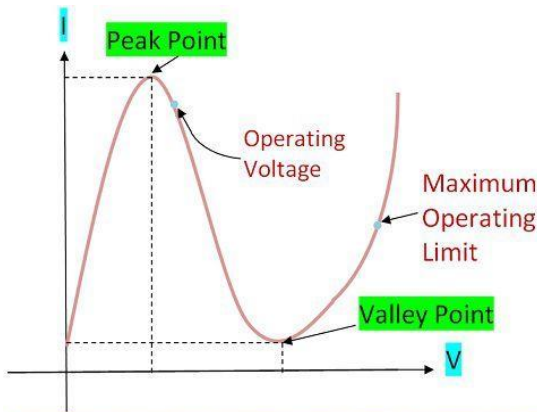
Gunn Diode's Working

This diode is made of a single piece of N-type semiconductor such as Gallium Arsenide and InP (Indium Phosphide). GaAs and some other semiconductor materials have one extra-energy band in their electronic band structure instead of having only two energy bands, viz. valence band and conduction band like normal semiconductor materials. These GaAs and some other semiconductor materials consist of three energy bands, and this extra third band is empty at initial stage.

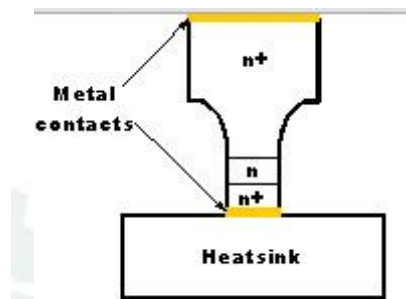
If a voltage is applied to this device, then most of the applied voltage appears across the active region. The electrons from the conduction band having negligible electrical resistivity are transferred into the third band because these electrons are scattered by the applied voltage. The third band of GaAs has mobility which is less than that of the conduction band.

Because of this, an increase in the forward voltage increases the field strength (for field strengths where applied voltage is greater than the threshold voltage value), then the number of electrons reaching the state at which the effective mass increases by decreasing their velocity, and thus, the current will decrease.

Thus, if the field strength is increased, then the drift velocity will decrease; this creates a negative incremental resistance region in V-I relationship. Thus, increase in the voltage will increase the resistance by creating a slice at the cathode and reaches the anode. But, to maintain a constant voltage, a new slice is created at the cathode. Similarly, if the voltage decreases, then the resistance will decrease by extinguishing any existing slice.



Characteristics of Gunn Diode



The current-voltage relationship characteristics of a Gunn diode are shown in the above graph with its negative resistance region. These characteristics are similar to the characteristics of the tunnel diode. As shown in the above graph, initially the current starts increasing in this diode, but after reaching a certain voltage level (at a specified voltage value called as threshold voltage value), the current decreases before increasing again. The region where the current falls is termed as a negative resistance region, and due to this it oscillates. In this negative resistance region, this diode acts as both oscillator and amplifier, as in this region, the diode is enabled to amplify signals.

Tunnel Diode

Tunnel Diode is the P-N junction device that exhibits negative resistance. When the voltage is increased than the current flowing through it decreases. It works on the principle of the Tunneling effect. Metal-Insulator-Metal (MIM) diode is another type of Tunnel diode, but its present application appears to be limited to research environments due to inherit sensitivities, its applications considered to be very limited to research environments. There is one more diode called Metal-Insulator-Insulator-Metal (MIIM) diode which includes an additional insulator layer. The tunnel diode is a two-terminal device with n-type semiconductor as the cathode and p-type semiconductor as an anode.

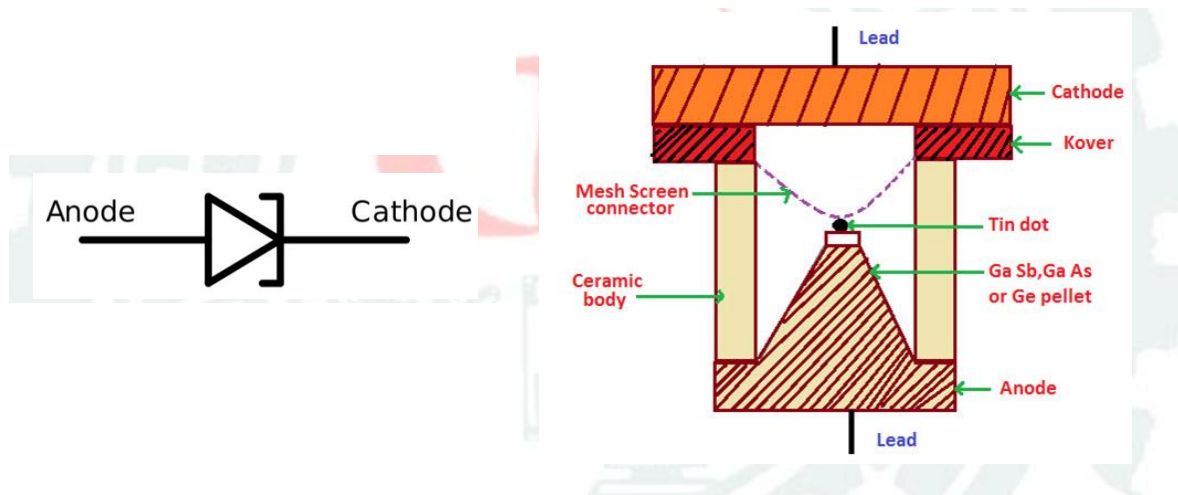
Tunnel Diode Working Phenomenon

Based on the classical mechanics' theory, a particle must acquire energy which is equal to the potential energy barrier height, if it has to move from one side of the barrier to the other. Otherwise, energy has to be supplied from some external

source, so the N-sided electrons of the junction can jump over the junction barrier to reach the P-side of the junction. If the barrier is thin such as in tunnel diode, according to the Schrodinger equation implies that there is a large amount of probability and then an electron will penetrate through the barrier. This process will happen without any energy loss on the part of the electron. The behavior of the quantum mechanical indicates tunneling. The high-impurity P-N junction devices are called as tunnel-diodes. The tunneling phenomenon provides a majority carrier effect.

Construction of Tunnel Diode

The diode has a ceramic body and a hermetically sealing lid on top. A small tin dot is alloyed or soldered to a heavily doped pellet of n-type Ge. The pellet is soldered to anode contact which is used for heat dissipation. The tin-dot is connected to the cathode contact via a mesh screen is used to reduce the inductance.



Operation and its Characteristics

The operation of the tunnel diode mainly includes two biasing methods such as forward and reverse

Forward Bias Condition

Under the forward bias condition, as voltage increases, then-current decreases and thus become increasingly misaligned, known as negative resistance. An increase in voltage will lead to operating as a normal diode where the conduction of electrons travels across the P-N junction diode. The negative resistance region is the most important operating region for a Tunnel diode. The Tunnel diode and normal P-N junction diode characteristics are different from each other.

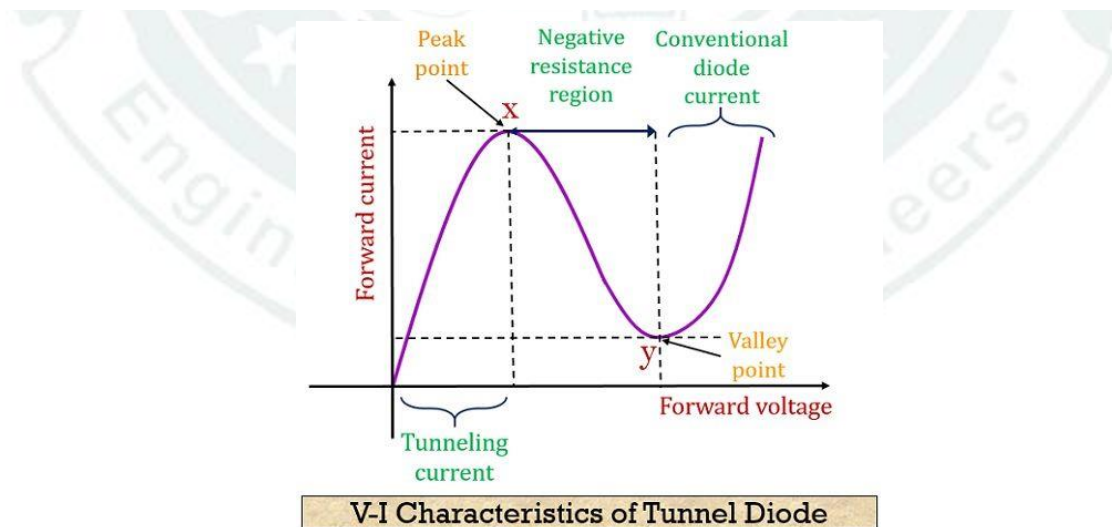
Reverse Bias Condition

Under the reverse condition, the tunnel diode acts as a back diode or backward diode. With zero offset voltage, it can act as a fast rectifier. In reverse bias condition, the empty states on the n-side aligned with the filled states on the p-side. In the reverse direction, the electrons will tunnel through a potential barrier. Because of its high doping concentrations, tunnel diode acts as an excellent conductor.

The forward resistance is very small because of its tunneling effect. An increase in voltage will lead to an increase in the current until it reaches peak current. But if the voltage increased beyond the peak voltage then current will decrease automatically. This negative resistance region prevails till the valley point. The current through the diode is minimum at valley point. The tunnel diode acts as a normal diode if it is beyond the valley point.

Tunnel Diode Applications

- Due to the tunneling mechanism, it is used as an ultra high speed switch.
- The switching time is of the order of nanoseconds or even picoseconds.
- Due to the triple valued feature of its curve from current, it is used as a logic memory storage device.
- Due to extremely small capacitance, inductance and negative resistance, it is used as a microwave oscillator at a frequency of about 10 GHz.
- Due to its negative resistance, it is used as a relaxation oscillator circuit.



Unit 5

BROADBAND COMMUNICATION

Broadband networks, such as asynchronous transfer mode (ATM), frame relay, and leased lines, allow us to easily access multimedia services (data, voice, and video) in one package. Exploring why broadband networks are important in modern-day telecommunications, *Introduction to Broadband Communication Systems* covers the concepts and components of both standard and emerging broadband communication network systems. After introducing the fundamental concepts of broadband communication systems, the book discusses Internet-based networks, such as intranets and extranets. It then addresses the networking technologies of X.25 and frame relay, fiber channels, a synchronous optical network (SONET), a virtual private network (VPN), an integrated service digital network (ISDN), broadband ISDN (B-ISDN), and ATM. The authors also cover access networks, including digital subscriber lines (DSL), cable modems, and passive optical networks, as well as explore wireless networks, such as wireless data services, personal communications services (PCS), and satellite communications. The book concludes with chapters on network management, network security, and network testing, fault tolerance, and analysis. With up-to-date, detailed information on the state-of-the-art technology in broadband communication systems, this resource illustrates how some networks have the potential of eventually replacing traditional dial-up Internet. Requiring only a general knowledge of communication systems theory, the text is suitable for a one- or two-semester course for advanced undergraduate and beginning graduate students in engineering as well as for short seminars on broadband communication systems.

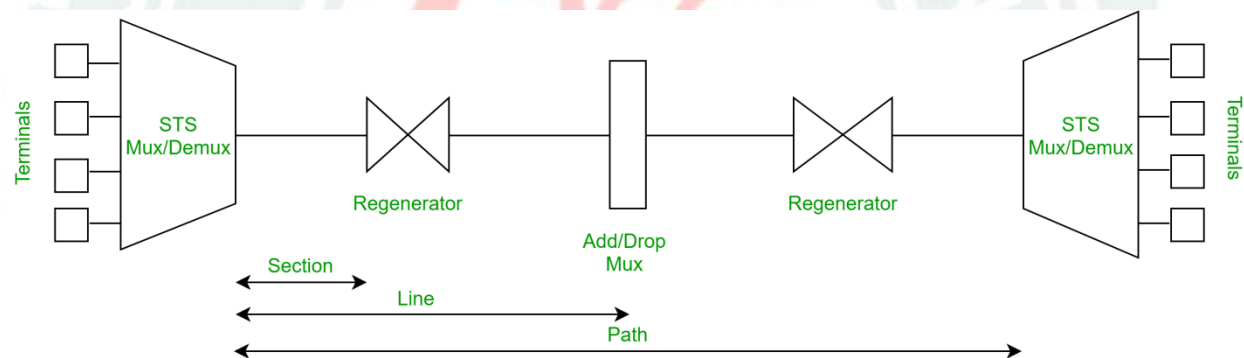
Cable Broadband Data Network Architecture

Cable networks deliver Internet access through a shared architecture that is distinct from DSL or fiber to the premise (FTTP) networks. An overview of cable architecture is, therefore, essential for developing the WFFs of cable networks and understanding how they differ from WFFs pertaining to networks of other technologies. [5] Cable broadband networks utilize statistical multiplexing to share a fixed amount of network capacity across a group of users. The network's architecture is a hybrid of fiber and coaxial cable, utilizing frequency division duplexing to divide upstream and downstream transmissions. Approximately 750

MHz to 1 GHz of spectrum is typically available on cable networks to be shared across all services, including television, broadband, and voice.² Upstream traffic uses the lower portion of the frequency duplex cable system, between 5 MHz and 42 MHz usually, and downstream traffic uses the remaining upper portion of the available frequencies. The precise amounts of upstream, downstream, and total spectrum available may vary by network, but the values noted here are typical for North American systems.

SONET

SONET stands for Synchronous Optical Network. SONET is a communication protocol, developed by Bellcore – that is used to transmit a large amount of data over relatively large distances using optical fibre. With SONET, multiple digital data streams are transferred at the same time over the optical fibre.



SONET Network Elements:

1. STS Multiplexer:

- Performs multiplexing of signals
- Converts electrical signal to optical signal

2. STS Demultiplexer:

- Performs demultiplexing of signals
- Converts optical signal to electrical signal

3. Regenerator:

It is a repeater, that takes an optical signal and regenerates (increases the strength) it.

4. Add/Drop Multiplexer:

It allows to add signals coming from different sources into a given path or remove a signal.

SONET Connections:

- **Section:** Portion of network connecting two neighbouring devices.
- **Line:** Portion of network connecting two neighbouring multiplexers.
- **Path:** End-to-end portion of the network.

SONET Layers:

1. Path Layer:

- It is responsible for the movement of signals from its optical source to its optical destination.
- STS Mux/Demux provides path layer functions.

2. Line Layer:

- It is responsible for the movement of signal across a physical line.
- STS Mux/Demux and Add/Drop Mux provides Line layer functions.

3. Section Layer:

- It is responsible for the movement of signal across a physical section.
- Each device of network provides section layer functions.

4. Photonic Layer:

- It corresponds to the physical layer of the OSI model.
- It includes physical specifications for the optical fibre channel (presence of light = 1 and absence of light = 0).

Advantages of SONET:

- Transmits data to large distances
- Low electromagnetic interference
- High data rates
- Large Bandwidth

ISDN

ISDN is a circuit-switched telephone network system, but it also provides access to packet-switched networks that allows digital transmission of voice and data. This results in potentially better voice or data quality than an analog phone can provide. It provides a packet-switched connection for data in increments of 64 kilobit/s.

ISDN INTERFACE

Basic Rate Interface

There are two data-bearing channels ('B' channels) and one signaling channel ('D' channel) in BRI to initiate connections. The B channels operate at a maximum of 64 Kbps while the D channel operates at a maximum of 16 Kbps. The two channels are independent of each other. For example, one channel is used as a TCP/IP connection to a location while the other channel is used to send a fax to a remote location. In iSeries ISDN supports a basic rate interface (BRI). The basic rate interface (BRI) specifies a digital pipe consisting of two B channels of 64 Kbps each and one D channel of 16 Kbps. This equals a speed of 144 Kbps. In addition, the BRI service itself requires an operating overhead of 48 Kbps. Therefore a digital pipe of 192 Kbps is required.

Primary Rate Interface

Primary Rate Interface service consists of a D channel and either 23 or 30 B channels depending on the country you are in. PRI is not supported on the iSeries. A digital pipe with 23 B channels and one 64 Kbps D channel is present in the usual Primary Rate Interface (PRI). Twenty-three B channels of 64 Kbps each and one D channel of 64 Kbps equals 1.536 Mbps. The PRI service uses 8 Kbps of overhead also. Therefore PRI requires a digital pipe of 1.544 Mbps.

Broadband ISDN

Narrowband ISDN has been designed to operate over the current communications infrastructure, which is heavily dependent on the copper cable however B-ISDN relies mainly on the evolution of fiber optics. According to CCITT B-ISDN is best described as 'a service requiring transmission channels capable of supporting rates greater than the primary rate.'

ISDN Services

ISDN provides a fully integrated digital service to users. These services fall into 3 categories- bearer services, teleservices, and supplementary services.

Bearer Services

Transfer of information (voice, data, and video) between users without the network manipulating the content of that information is provided by the bearer

network. There is no need for the network to process the information and therefore does not change the content. Bearer services belong to the first three layers of the OSI model. They are well defined in the ISDN standard. They can be provided using circuit-switched, packet-switched, frame-switched, or cell-switched networks.

Tele Services

In this, the network may change or process the contents of the data. These services correspond to layers 4-7 of the OSI model. Teleservices rely on the facilities of the bearer services and are designed to accommodate complex user needs. The user need not be aware of the details of the process. Teleservices include telephony, teletex, telefax, videotex, telex, and teleconferencing. Though the ISDN defines these services by name yet they have not yet become standards.

Supplimentary Services

Additional functionality to the bearer services and teleservices are provided by supplementary services. Reverse charging, call waiting, and message handling are examples of supplementary services which are all familiar with today's telephone company services.

Advantages of ISDN:

- ISDN channels have a reliable connection.
- ISDN is used to facilitate the user with multiple digital channels.
- It has faster data transfer rate.

Disadvantages of ISDN:

- ISDN lines costlier than the other telephone system.
- It requires specialized digital devices.
- It is less flexible.

BISDN

- Recommendations to support video services as well as normal ISDN services
- Exploits optical fiber transmission technology

BISDN

There are two types of B-ISDN services which are as follows –

- **Interactive Services** – Two-way exchange of information (other than control signalling information) between two subscribers or between a subscriber and a service provider.
- **Distribution Services** – Primarily one way transfer of information, from service provider to B-ISDN subscriber.

BISDN Architecture

The architecture of the B-ISDN includes low Layer capabilities and high Layer capabilities. These capabilities support the services within the B-ISDN and other networks by means of interworking B-ISDN with those networks.

Low Layer capabilities

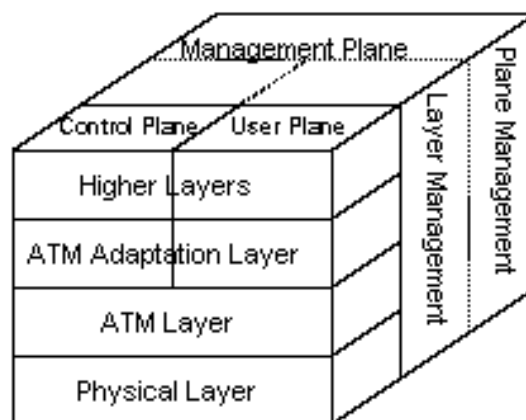
The low layer capabilities of B-ISDN architecture are explained below –

- From the functional capabilities of the B-ISDN, as shown in Figure, the information transfer capabilities require further description.
- Broadband information transfer is provided by an ATM at the B-ISDN user-network interface (UNI) and at switching entities inside the network.

High Layer capabilities

The high layer capabilities of B-ISDN architecture are explained below –

- Normally, the high Layer functional capabilities are involved only in the terminal equipment.
- The support of some services, provision of high layer functions could be made through special nodes in the B-ISDN belonging to the public network or to centres operated by other organizations and accessed via B-ISDN user-network or network node interfaces (NNIs).



User plane

The user plane, with its layered structure, provides for user information flow transfer, along with associated controls (e.g. flow control, and recovery from errors, etc.).

Control plane

This plane has a layered structure and performs the call control and connection control functions

Management Plane

Plane Management function

It provides coordination between all the planes. It also performs management functions related to a system as a whole.

Layer Management function

It handles the operation and maintenance information flows specific to the layer concerned.

Function of each layer

Physical layer

It describes the physical transmission media. We can use shielded twisted pair, coaxial cable etc

ATM layer

The ATM Layer is independent of the physical medium. It translates the cell identifier for switching, provides the user with one QoS class, to mention a few. The ATM Adaptation Layer adapts the ATM Layer service along the requirements imposed by user services as well as control and management functions.

ATM Adaptation layer

It converts the submitted information into streams of 48 octate segments and transports these in the payload field of multiple ATM cells relating to the same cell. It converts the 48 octate information field into required form for delivery to the particular higher protocol layer.