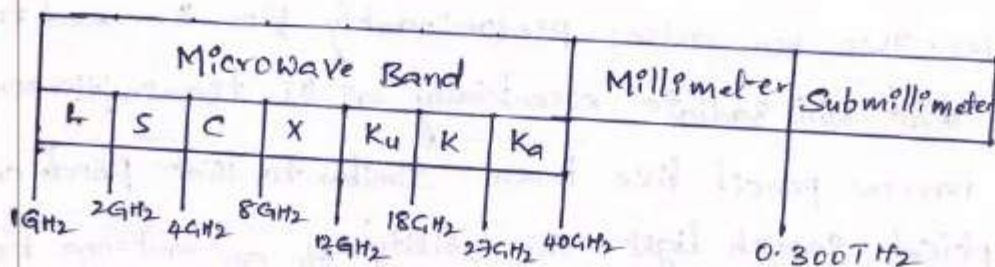
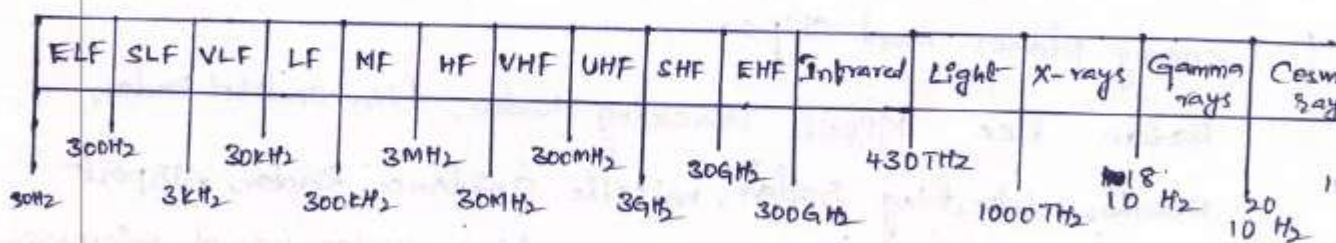


Microwave Transmission Lines

- Microwave frequencies generally used for those waves measured in centimeters, roughly from 30cm to 1mm (i.e. 300 GHz).
- Microwave frequencies are upto infrared and visible regions.
- Micro refers to tinyness of the wavelength and the period a cycle of a cm wave.
- The wavelength (λ) of cm waves at microwave frequencies are very short. In short, a microwave is a signal that has a wavelength of 1 foot or less i.e. $\lambda \leq 30.5 \text{ cm} \approx 1 \text{ foot}$. This converts to a frequency of 984 MHz approximately $= 1 \text{ GHz}$. So all frequencies above approximately 1 GHz are microwave frequencies.

Microwave spectrum and Bands:



- Electromagnetic waves are oscillations that propagate through space with a velocity of light $c = 3 \times 10^8 \text{ m/sec}$.
- Electromagnetic waves are transverse in nature i.e. E-field & H-field are perpendicular to each other and they both are perpendicular to the direction of propagation.

Applications of Microwaves:

→ Microwaves have a broad range of applications in modern technology. Most important among them are in long distance communication systems, Radars, Radio Astronomy, Navigation etc.

(1) Telecommunication:

Intercontinental Telephone and TV, space communication, telemetry communication link for railways etc.

→ Increased Bandwidth for communication channels requires higher carrier frequencies. A microwave link is point-to-point link using the propagation of EM waves at microwave frequencies through free space.

(2) Radars:-

Microwave engineering was almost synonymous with RADAR [Radio detection and Ranging] engineering because of the great stimulus given to the development of microwave systems by the need for high resolution radar capable of detecting and locating enemy planes and ships.

Radars like Missile Tracking Radar, fire control Radar, weather detecting radar, missile guidance Radar, Airspace traffic control radar etc represents a major use of microwave frequencies. This use arises predominantly from the need to have antennas that will radiate essentially all the transmitter power into a narrow pencil like beam similar to that produced by an optical search light. The ability of an antenna to concentrate radiation into a narrow beam is limited by diffraction effects, which in turn are governed by the relative size of the radiating aperture in terms of wavelengths.

(3) Commercial and Industrial Applications

- The domestic microwave oven operates at 2,450 MHz and uses a magnetron tube with a power output of 500 to 1000 W.
- For industrial heating applications, such as drying grain, manufacturing wood and paper products, and material curing, the frequencies of 915 and 2450 MHz have been assigned.
- Biomedical applications (diagnostic/therapeutic) - diathermy for localised superficial heating, deep electromagnetic heating for treatment of cancer, hyperthermia, electromagnetic transmission through human has been used for monitoring of heart beat, lung water detection etc.
- Drying seeds, drying/sterilising grains, sterilising pharmaceuticals, drying textiles, leather, tobacco, power transmission.

(4) Electronic Warfare:-

ECCM (Electronic Counter Measure) / Electronic Counter Counter Measures, spread spectrum systems.

(5) Identifying objects or personnel by non-contact method.

Waveguides:-

- A waveguide consists of a hollow metallic tube of a rectangular or circular shape used to guide an electromagnetic wave.
- Waveguides are used principally at frequencies in the microwave range.
- In waveguides the electric and magnetic fields are confined to the space within the guides. Thus no power is lost through radiation and even the dielectric loss is negligible since the guides are normally air-filled. However, there is some power loss as heat in the walls of the guides, but the loss is very small.

- It is possible to propagate several modes of electromagnetic waves within a waveguide. These modes correspond to the solutions of Maxwell's Equations for particular waveguides.
- If the frequency of the impressed signal is above the cutoff frequency for a given mode, the electromagnetic energy can be transmitted through the guide for that particular mode without attenuation.
- The mode which is having the lowest cut-off frequency is called the "Dominant Mode".
- Waveguides are two types. They are (i) Rectangular (ii) Circular w/g's.
Rectangular Waveguides : —
- A Rectangular waveguide is a hollow metallic tube with a rectangular cross section.
- When the waves travel longitudinally down the guide, because of conducting walls plane waves are reflected from wall to wall. This process results in a component of either electric or magnetic field in the direction of propagation of the resultant wave. Therefore the wave is no longer a transverse electromagnetic (TEM) wave.
- Any uniform plane wave in a lossless guide may be resolved into TE and TM waves.
- A plane wave in a waveguide resolves into two components: one standing wave in the direction normal to the reflecting walls of the guide and one travelling wave in the direction parallel to the reflecting walls.
- In lossless waveguides the modes may be classified as either transverse electric (TE) mode or transverse magnetic (TM) mode.
- In rectangular guides the modes are designated TE_{mn} or TM_{mn} .

Propagation of waves in Rectangular waveguides:

Consider a rectangular waveguide situated in the rectangular coordinate system with its breadth along x-axis, width along y-axis and the wave is assumed to propagate along the z-direction. Waveguide is filled with air. In a waveguide no TEM wave exists.

TEM (Transverse Electromagnetic wave): In TEM both electric and magnetic fields are purely transverse to the direction of propagation and consequently have no z directed E & H-Components.

TE (Transverse Electric) Wave: In TE Wave, only the E-field is purely transverse to the direction of propagation and the magnetic field is not purely transverse. i.e., $E_z = 0, H_z \neq 0$.

TM (Transverse Magnetic) Wave: In this, only the magnetic field is purely transverse to the direction of propagation and the electric field is not purely transverse. $E_z \neq 0, H_z = 0$

HE (Hybrid Wave): In this, neither electric nor magnetic fields are purely transverse to the direction of propagation. $E_z \neq 0, H_z \neq 0$.

Since we assumed that the wave direction is along z-direction then the wave equations are $\nabla^2 E_z = \gamma^2 E_z$ General wave eqn $\nabla^2 E = \mu \epsilon \frac{\partial^2 E}{\partial t^2} + \mu \epsilon \frac{\partial^2 E}{\partial z^2}$ free space $\sigma = 0$.

$\nabla^2 E_z = -\omega^2 \mu \epsilon E_z$ for TM wave

and $\nabla^2 H_z = -\omega^2 \mu \epsilon H_z$ for TE wave

$$\gamma^2 = -\omega^2 \mu \epsilon$$

$$\gamma = j\beta$$

$$\beta = \omega \sqrt{\mu \epsilon}$$

$$\gamma = j\omega \sqrt{\mu \epsilon}$$

where $E_z = E_0 e^{-\gamma z}$, $H_z = H_0 e^{-\gamma z}$

E_0 & $H_0 \rightarrow$ Amplitudes of respective Components at $z=0$.
 $\gamma \rightarrow$ propagation constant, $\gamma = \alpha + j\beta$
 $\downarrow$
 Attenuation constant phase constant = $2\pi/\lambda$

The condition for wave propagation is that γ must be imaginary.

Differentiating eqn (3) w.r.t. 'z', we get

$\frac{\partial E_z}{\partial z} = E_0 e^{-\gamma z} (-\gamma) = -\gamma E_0 e^{-\gamma z} = -\gamma E_z \rightarrow (4)$ Since $E_z = E_0 e^{-\gamma z}$

Hence we can define the operator, $\frac{\partial}{\partial z} = -\gamma \rightarrow (5)$

By differentiating eqn (4) w.r.t. z , we get

$$\frac{\partial^2 E_z}{\partial z^2} = -r \frac{\partial E_z}{\partial z} = -r \frac{\partial}{\partial z} [E_0 e^{-rz}] = -r(-r) E_0 e^{-rz}$$

$$\Rightarrow \boxed{\frac{\partial^2 E_z}{\partial z^2} = r^2 E_z}$$

We can define the operator,

$$\boxed{\frac{\partial^2}{\partial z^2} = r^2} \rightarrow \textcircled{6}$$

$$\nabla^2 E_z = \frac{\partial^2 E_z}{\partial z^2} (\mu_0 \epsilon_0)$$

from eqn (1), we can write

$$\nabla^2 E_z = -\omega^2 \mu \epsilon E_z$$

By expanding $\nabla^2 E_z$ in rectangular coordinate system

$$\frac{\partial^2 E_z}{\partial x^2} + \frac{\partial^2 E_z}{\partial y^2} + \frac{\partial^2 E_z}{\partial z^2} = -\omega^2 \mu \epsilon E_z \rightarrow \textcircled{7}$$

from eqn (6) & (7),

$$\frac{\partial^2 E_z}{\partial x^2} + \frac{\partial^2 E_z}{\partial y^2} + r^2 E_z = -\omega^2 \mu \epsilon E_z \Rightarrow \frac{\partial^2 E_z}{\partial x^2} + \frac{\partial^2 E_z}{\partial y^2} + (r^2 + \omega^2 \mu \epsilon) E_z = 0 \rightarrow \textcircled{8}$$

Let $r^2 + \omega^2 \mu \epsilon = h^2$, be a constant. Then eqn (8) can be written as

$$\frac{\partial^2 E_z}{\partial x^2} + \frac{\partial^2 E_z}{\partial y^2} + h^2 E_z = 0 \text{ for TM wave} \rightarrow \textcircled{9}$$

$$\text{Similarly, } \frac{\partial^2 H_z}{\partial x^2} + \frac{\partial^2 H_z}{\partial y^2} + h^2 H_z = 0 \text{ for TE wave} \rightarrow \textcircled{10}$$

By solving above two partial differential equations, we get solutions for E_z and H_z . Using Maxwell's equation, it is possible to find the various components along x and y -directions.

from Maxwell's 1st equation, we have

$$\nabla \times H = j\omega \epsilon E$$

$$\begin{vmatrix} a_x & a_y & a_z \\ \frac{\partial}{\partial x} & \frac{\partial}{\partial y} & \frac{\partial}{\partial z} \\ H_x & H_y & H_z \end{vmatrix} = j\omega \epsilon [E_x a_x + E_y a_y + E_z a_z]$$

$$\Rightarrow a_x \left[\frac{\partial H_z}{\partial y} - \frac{\partial H_y}{\partial z} \right] - a_y \left[\frac{\partial H_z}{\partial x} - \frac{\partial H_x}{\partial z} \right] + a_z \left[\frac{\partial H_y}{\partial x} - \frac{\partial H_x}{\partial y} \right] = j\omega \epsilon [E_x a_x + E_y a_y + E_z a_z]$$

Since $\frac{\partial}{\partial z} = -r \Rightarrow$

$$\therefore a_x \left[\frac{\partial H_z}{\partial y} + r H_y \right] - a_y \left[\frac{\partial H_z}{\partial x} + r H_x \right] + a_z \left[\frac{\partial H_y}{\partial x} - \frac{\partial H_x}{\partial y} \right] = j\omega \epsilon [E_x a_x + E_y a_y + E_z a_z]$$

By Comparing a_x, a_y and a_z Components

$$a_x \rightarrow r H_y + \frac{\partial H_z}{\partial y} = j\omega \epsilon E_x \longrightarrow (11)$$

$$a_y \rightarrow r H_x + \frac{\partial H_z}{\partial x} = -j\omega \epsilon E_y \longrightarrow (12)$$

$$a_z \rightarrow \frac{\partial H_y}{\partial x} - \frac{\partial H_x}{\partial y} = j\omega \epsilon E_z \longrightarrow (13)$$

Similarly from Maxwell's 2nd equation, we have

$$\nabla \times E = -j\omega \mu H$$

By expanding.

$$\begin{vmatrix} a_x & a_y & a_z \\ \frac{\partial}{\partial x} & \frac{\partial}{\partial y} & \frac{\partial}{\partial z} \\ E_x & E_y & E_z \end{vmatrix} = -j\omega \mu [H_x a_x + H_y a_y + H_z a_z]$$

Since $\frac{\partial}{\partial z} = -r$

$$\begin{vmatrix} a_x & a_y & a_z \\ \frac{\partial}{\partial x} & \frac{\partial}{\partial y} & -r \\ E_x & E_y & E_z \end{vmatrix} = -j\omega \mu [H_x a_x + H_y a_y + H_z a_z]$$

$$\Rightarrow a_x \left[\frac{\partial E_z}{\partial y} + r E_y \right] - a_y \left[\frac{\partial E_z}{\partial x} + r E_x \right] + a_z \left[\frac{\partial E_y}{\partial x} - \frac{\partial E_x}{\partial y} \right] = -j\omega \mu [H_x a_x + H_y a_y + H_z a_z]$$

By Comparing a_x, a_y & a_z components

$$a_x \rightarrow r E_y + \frac{\partial E_z}{\partial y} = -j\omega \mu H_x \longrightarrow (14)$$

$$r E_x + \frac{\partial E_z}{\partial x} = j\omega \mu H_y \longrightarrow (15)$$

$$\frac{\partial E_y}{\partial x} - \frac{\partial E_x}{\partial y} = -j\omega \mu H_z \longrightarrow (16)$$

from eqn (15), $H_y = \frac{1}{j\omega \mu} \left[r E_x + \frac{\partial E_z}{\partial x} \right] \longrightarrow (17)$

By substituting eqn (17) in eqn (11), we get

$$\Rightarrow \frac{r^2 E_x}{j\omega \mu} + \frac{r}{j\omega \mu} \frac{\partial E_z}{\partial x} + \frac{\partial H_z}{\partial y} = j\omega \epsilon E_x$$

$$\Rightarrow \frac{r}{j\omega \mu} \left[\frac{\partial E_z}{\partial x} \right] + \frac{\partial H_z}{\partial y} = \left(j\omega \epsilon - \frac{r^2}{j\omega \mu} \right) E_x = \left[\frac{-\omega^2 \mu \epsilon - r^2}{j\omega \mu} \right] E_x$$

$$\Rightarrow \left(\frac{-r^2 - \omega^2 \mu \epsilon}{j\omega \mu} \right) E_x = \frac{r}{j\omega \mu} \left[\frac{\partial E_z}{\partial x} \right] + \frac{j\omega \mu}{j\omega \mu} \left[\frac{\partial H_z}{\partial y} \right]$$

$$\Rightarrow E_x \left[-(r^2 + \omega^2 \mu \epsilon) \right] = r \frac{\partial E_z}{\partial x} + j\omega \mu \frac{\partial H_z}{\partial y}$$

Since $r^2 + w^2 \mu \epsilon = h^2$

By dividing the above equation with h^2 , we get

$$E_x = \frac{-r}{h^2} \frac{\partial E_z}{\partial x} - \frac{j\omega\mu}{h^2} \frac{\partial H_z}{\partial y} \rightarrow (18)$$

Similarly, $E_y = \frac{-r}{h^2} \frac{\partial E_z}{\partial y} + \frac{j\omega\mu}{h^2} \frac{\partial H_z}{\partial x} \rightarrow (19)$

and $H_x = \frac{-r}{h^2} \frac{\partial H_z}{\partial x} + \frac{j\omega\epsilon}{h^2} \frac{\partial E_z}{\partial y} \rightarrow (20)$

$$H_y = \frac{-r}{h^2} \frac{\partial H_z}{\partial y} - \frac{j\omega\epsilon}{h^2} \frac{\partial E_z}{\partial x} \rightarrow (21)$$

These equations give a general relationship for field components within a waveguide.

Propagation of TEM Waves

for a TEM wave.

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

Substituting these values in eqns (18) to (21) all the field components along X and Y - directions E_x, E_y, H_x, H_y vanish and hence a TEM wave cannot exist inside a waveguide.

Modes

The electromagnetic wave inside a waveguide can have an infinite number of patterns which are called "Modes".

→ The electric field cannot have a component parallel to the surface i.e., the electric field must always be perpendicular to the surface at the conductor.

→ The magnetic field on the other hand always parallel to the surface of the conductor and cannot have a component perpendicular to it at the surface.

→ The mode having the highest cut-off wave length is known as "Dominant mode" of the waveguide and all other modes are called "higher modes".

→ We use subscript for designating a particular mode. TE_{mn} or TM_{mn} where

The integer m denotes the number of half waves of electric or magnetic intensity in the x -direction, and n is the number of half waves in the y -direction if the propagation of the wave is assumed in the positive z -direction.

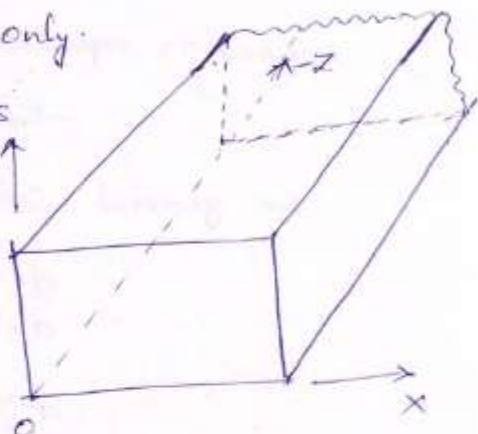
Solutions of Wave Equations in Rectangular Coordinates:

→ There are time domain and frequency domain solutions for each wave equation. For the simplicity of the solution, we are (going for) solving in frequency domain only.

→ The electric and magnetic wave equations in frequency domain are given by

$$\begin{cases} \nabla^2 E = \gamma^2 E & \rightarrow (1) \\ \nabla^2 H = \gamma^2 H & \rightarrow (2) \end{cases}$$

where $\gamma = \sqrt{j\omega\mu(\sigma + j\omega\epsilon)} = \alpha + j\beta$



These are called the 'Vector Wave Equations'.

Rectangular coordinates are the usual right hand system.

The rectangular components of E and H satisfy the Complex scalar wave equation or Helmholtz equation.

$$\nabla^2 \psi = \gamma^2 \psi \quad \rightarrow (3)$$

The Helmholtz equation in rectangular coordinates is

$$\left[\frac{\partial^2 \psi}{\partial x^2} + \frac{\partial^2 \psi}{\partial y^2} + \frac{\partial^2 \psi}{\partial z^2} = \gamma^2 \psi \right] \rightarrow (4)$$

This is a linear and inhomogeneous partial differential equation in three-dimensions.

By the method of separation of variables, the solution is assumed in the form of

$$\psi = X(x) Y(y) Z(z) \quad \rightarrow (5)$$

where $X(x) \rightarrow$ a function of the x -coordinate only
 $Y(y) \rightarrow$ a function of the y -coordinate only
 $Z(z) \rightarrow$ a function of the z -coordinate only

substitution of eqn(5) in eqn(4) and division of the resultant by

eqn(5) gives

$$\frac{1}{x} \frac{d^2 x}{dx^2} + \frac{1}{y} \frac{d^2 y}{dy^2} + \frac{1}{z} \frac{d^2 z}{dz^2} = r^2 \quad \rightarrow (6)$$

Since the sum of the three terms on the left hand side is a constant and each term is independently variable, it follows that each term must be equal to a constant.

Let the three terms be K_x^2 , K_y^2 and K_z^2 respectively; then the separation equation is given by

$$-K_x^2 - K_y^2 - K_z^2 = r^2 \quad \rightarrow (7)$$

The general solution of each differential equation in eqn(6)

$$\frac{d^2 x}{dx^2} = -K_x^2 x \quad \rightarrow (8)$$

$$\frac{d^2 y}{dy^2} = -K_y^2 y \quad \rightarrow (9)$$

$$\frac{d^2 z}{dz^2} = -K_z^2 z \quad \rightarrow (10)$$

will be in the form of

$$x = A \sin(K_x x) + B \cos(K_x x) \quad \rightarrow (11)$$

$$y = C \sin(K_y y) + D \cos(K_y y) \quad \rightarrow (12)$$

$$z = E \sin(K_z z) + F \cos(K_z z) \quad \rightarrow (13)$$

The total solution of the Helmholtz equation in rectangular coordinates is

$$\psi = [A \sin(K_x x) + B \cos(K_x x)] [C \sin(K_y y) + D \cos(K_y y)] * [E \sin(K_z z) + F \cos(K_z z)] \quad \rightarrow (14)$$

The propagation of the wave in the guide is conveniently assumed in the positive z direction. It should be noted that the propagation constant γ_g in the guide differs from the intrinsic propagation constant γ of the dielectric.

Let

$$\gamma^2 = r^2 + k_x^2 + k_y^2 = r^2 + k_c^2$$

where $k_c = \sqrt{k_x^2 + k_y^2}$ is usually called the cut off wave number.

for a lossless dielectric, $r^2 = -\omega^2 \mu \epsilon$ then

$$\gamma = \pm j \sqrt{\omega^2 \mu \epsilon - k_c^2}$$

Since $r = \alpha + j\beta$
 $r = j\beta, r^2 = -\beta^2$
 $k_c^2 = -\beta^2$
 $r^2 = -\omega^2 \mu \epsilon$

There are three cases for the propagation constant γ in the waveguide

Case (i): - There will be no wave propagation in the guide if $\omega_c^2 \mu \epsilon = k_c^2$ and $\gamma = 0$. This is the critical condition for cut off propagation. The cut off frequency is expressed as

$$f_c = \frac{1}{2\pi \sqrt{\mu \epsilon}} \sqrt{k_x^2 + k_y^2}$$

$$\omega^2 \mu \epsilon - k_c^2 = 0 \Rightarrow \omega^2 \mu \epsilon = (\sqrt{k_x^2 + k_y^2})^2$$

$$(2\pi f_c)^2 \mu \epsilon = \omega^2 = \frac{1}{\mu \epsilon} (k_x^2 + k_y^2)$$

$$f_c = \frac{1}{2\pi \sqrt{\mu \epsilon}} \sqrt{k_x^2 + k_y^2} = \frac{k_c}{2\pi \sqrt{\mu \epsilon}}$$

Case (ii): - The wave will be propagating in the guide if $\omega^2 \mu \epsilon > k_c^2$ and

$$\gamma = \pm j\beta_g = \pm j\omega \sqrt{\mu \epsilon} \sqrt{1 - (f_c/f)^2}$$

This means that the operating frequency must be above the cut off frequency in order for a wave to propagate in the guide.

Case (iii): -

The wave will be attenuated if $\omega^2 \mu \epsilon < k_c^2$ and

$$\gamma = \pm \alpha_g = \pm \omega \sqrt{\mu \epsilon} \sqrt{(f_c/f)^2 - 1}$$

$$= \pm j\omega \sqrt{\mu \epsilon} \sqrt{1 - \frac{4\pi^2 f_c^2 \mu \epsilon}{\omega^2}}$$

$$= \pm j\omega \sqrt{\mu \epsilon} \sqrt{1 - (f_c/f)^2}$$

This means that if the operating frequency is below the cut off frequency the wave will decay exponentially with respect to a factor of $-\alpha_g z$ and there will be no wave propagation because the propagation constant is a real quantity. Therefore the solution to the Helmholtz equation in rectangular coordinates is given by

$$\psi = [A \sin(k_x x) + B \cos(k_x x)] [C \sin(k_y y) + D \cos(k_y y)] e^{-\alpha_g z}$$

TE mode Analysis:

The TE_{mn} modes in a rectangular waveguide are characterised by $E_z = 0$. The z component of the magnetic field, H_z , must exist in order to have energy transmission in the guide.

The wave equation [Helmholtz Equation] for TE wave is given by

$$\boxed{\nabla^2 H_z = -\omega^2 \mu \epsilon H_z} \longrightarrow \textcircled{1}$$

$$\text{i.e. } \frac{\partial^2 H_z}{\partial x^2} + \frac{\partial^2 H_z}{\partial y^2} + \frac{\partial^2 H_z}{\partial z^2} = -\omega^2 \mu \epsilon H_z$$

$$\left[\because \frac{\partial^2}{\partial z^2} = -\gamma^2 \right]$$

$$\Rightarrow \frac{\partial^2 H_z}{\partial x^2} + \frac{\partial^2 H_z}{\partial y^2} + \gamma^2 H_z + \omega^2 \mu \epsilon H_z = 0$$

$$\Rightarrow \frac{\partial^2 H_z}{\partial x^2} + \frac{\partial^2 H_z}{\partial y^2} + (\gamma^2 + \omega^2 \mu \epsilon) H_z = 0 \Rightarrow \boxed{\frac{\partial^2 H_z}{\partial x^2} + \frac{\partial^2 H_z}{\partial y^2} + h^2 H_z = 0} \longrightarrow \textcircled{2}$$

This is a partial differential equation whose solution can be assumed.

Assume a solution $\boxed{H_z = XY}$ where

$X \rightarrow$ pure function of x only

$Y \rightarrow$ pure function of y only

from eqn(2),

$$\frac{\partial^2 [XY]}{\partial x^2} + \frac{\partial^2 [XY]}{\partial y^2} + h^2 XY = 0$$

$$\Rightarrow Y \frac{\partial^2 X}{\partial x^2} + X \frac{\partial^2 Y}{\partial y^2} + h^2 XY = 0$$

Dividing above equation with XY on both sides

$$\Rightarrow \boxed{\frac{1}{X} \frac{\partial^2 X}{\partial x^2} + \frac{1}{Y} \frac{\partial^2 Y}{\partial y^2} + h^2 = 0} \longrightarrow \textcircled{3}$$

Here $\frac{1}{X} \frac{\partial^2 X}{\partial x^2}$ is purely a function of x ; and $\frac{1}{Y} \frac{\partial^2 Y}{\partial y^2}$ is purely a function of y .

$$\text{Let } \frac{1}{X} \frac{\partial^2 X}{\partial x^2} = -B^2 \quad \& \quad \frac{1}{Y} \frac{\partial^2 Y}{\partial y^2} = -A^2 \quad \text{where } A^2 \& B^2 \rightarrow \text{constants.}$$

Therefore from eqn(3),

$$-B^2 - A^2 + h^2 = 0 \Rightarrow \boxed{h^2 = A^2 + B^2} \longrightarrow \textcircled{4}$$

By separation of variable method,

$$X = C_1 \cos Bx + C_2 \sin Bx$$

$$Y = C_3 \cos Ay + C_4 \sin Ay$$

Therefore the complete solution for $H_z = XY$ is

$$H_z = (C_1 \cos Bx + C_2 \sin Bx)(C_3 \cos Ay + C_4 \sin Ay) \longrightarrow (5)$$

where C_1, C_2, C_3 and C_4 are constants, which can be evaluated by applying Boundary Conditions.

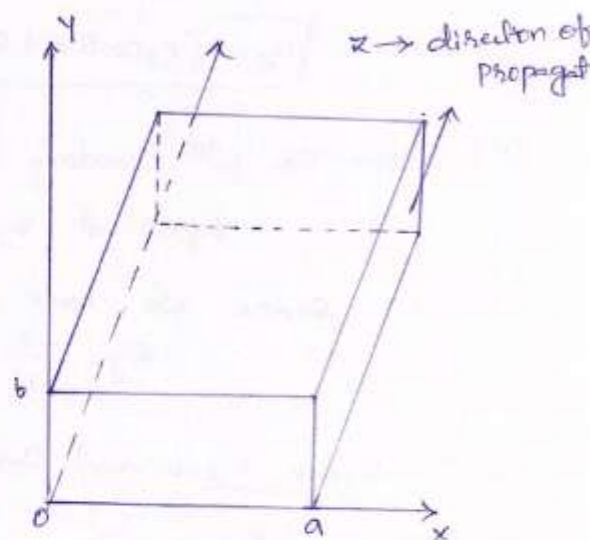
Boundary conditions:

Since we consider a TE wave, propagating along z -direction, so

$E_z = 0$ but we have components along x and y direction.

$E_x = 0 \rightarrow$ waves along bottom and top walls of the waveguide.

$E_y = 0 \rightarrow$ waves along left and right walls of the waveguide.



1st Boundary condition:

$E_x = 0$ at $y = 0 \forall x \rightarrow 0$ to a (bottom wall)

2nd Boundary condition:

$E_x = 0$ at $y = b \forall x \rightarrow 0$ to a (top wall)

3rd Boundary condition:

$E_y = 0$ at $x = 0 \forall y \rightarrow 0$ to b (left side wall)

4th Boundary condition:

$E_y = 0$ at $x = a \forall y \rightarrow 0$ to b (right side wall)

(i) Substituting 1st Boundary Condition in eqn (5)

Since we have

$$E_x = -\frac{\gamma}{h^2} \frac{\partial H_z}{\partial x} = -\frac{j\omega\mu}{h^2} \frac{\partial H_z}{\partial x} \longrightarrow (6)$$

$$\text{Since } E_z = 0 \Rightarrow E_x = -\frac{j\omega\mu}{h^2} \frac{\partial}{\partial x} [(C_1 \cos Bx + C_2 \sin Bx)(C_3 \cos Ay + C_4 \sin Ay)]$$

$$\therefore E_y = -\frac{j\omega\mu}{h^2} (C_1 \cos Bx + C_2 \sin Bx) (-AC_3 \sin Ay + AC_4 \cos Ay)$$

from the first Boundary condition we get

$$0 = -\frac{j\omega\mu}{h^2} (C_1 \cos Bx + C_2 \sin Bx) (0 + AC_4)$$

$$\text{Since } (C_1 \cos Bx + C_2 \sin Bx) \neq 0, A \neq 0.$$

$$\Rightarrow \boxed{C_4 = 0}$$

Substituting the value of C_4 in eqn(5), the solution reduces to,

$$\boxed{H_z = (C_1 \cos Bx + C_2 \sin Bx) (C_3 \cos Ay)} \longrightarrow \textcircled{7}$$

(ii) from the 3rd Boundary condition:

$$E_y = 0 \text{ at } y = 0 \text{ \& } y \rightarrow 0 \text{ to } b.$$

Since we have

$$E_y = -\frac{\gamma}{h^2} \frac{\partial E_z}{\partial y} + \frac{j\omega\mu}{h^2} \frac{\partial H_z}{\partial x} \longrightarrow \textcircled{8}$$

Since $E_z = 0$ and substituting the value of H_z in eqn(7), we get

$$E_y = \frac{j\omega\mu}{h^2} \frac{\partial}{\partial x} [(C_1 \cos Bx + C_2 \sin Bx) C_3 \cos Ay]$$

$$\Rightarrow E_y = \frac{j\omega\mu}{h^2} [C - BC_1 \sin Bx + BC_2 \cos Bx] C_3 \cos Ay$$

from 3rd Boundary condition,

$$0 = \frac{j\omega\mu}{h^2} (0 + BC_2) C_3 \cos Ay$$

$$\text{Since } \cos Ay \neq 0, B \neq 0, C_3 \neq 0$$

$$\therefore C_2 = 0$$

from eqn(7),

$$H_z = (C_1 \cos Bx) (C_3 \cos Ay)$$

$$\Rightarrow \boxed{H_z = C_1 C_3 \cos Bx \cos Ay} \longrightarrow \textcircled{9}$$

(iii) 2nd Boundary condition:

Since we have

$$E_x = -\frac{\gamma}{h^2} \frac{\partial E_z}{\partial x} - \frac{j\omega\mu}{h^2} \frac{\partial H_z}{\partial y}$$

$$= -\frac{j\omega\mu}{h^2} \frac{\partial}{\partial y} [C_1 C_3 \cos Bx \cos Ay] \quad [\because E_z = 0]$$

$$E_x = \frac{j\omega\mu}{h^2} C_1 C_3 A \cos Bx \sin Ay$$

from the second Boundary condition,

$$E_x = 0 \text{ at } y = b \quad \forall x \rightarrow 0 \text{ to } a$$

$$\therefore 0 = \frac{j\omega\mu}{h^2} C_1 C_3 A \cos Bx \sin Ab$$

$$\cos Bx \neq 0, \quad C_1, C_3 \neq 0$$

$$\sin Ab = 0 \text{ as } Ab = n\pi \quad \text{where } n = 0, 1, 2, \dots, \infty$$

$$\therefore \boxed{A = \frac{n\pi}{b}} \longrightarrow (10)$$

(iv) 4th Boundary Condition:

Since

$$E_y = \frac{-r}{h^2} \frac{\partial E_z}{\partial y} + \frac{j\omega\mu}{h^2} \frac{\partial H_z}{\partial x}$$

$$\Rightarrow E_y = \frac{j\omega\mu}{h^2} \frac{\partial}{\partial x} [C_1 C_3 \cos Bx \cos Ay] \quad \left(\text{Since } E_z = 0 \text{ \& } H_z = C_1 C_3 \cos Bx \cos Ay \right)$$

$$\Rightarrow E_y = \frac{-j\omega\mu}{h^2} C_1 C_3 \sin Bx \cdot B \cos Ay$$

from the 4th Boundary condition,

$$E_y = 0 \text{ at } x = a \quad \forall y \rightarrow 0 \text{ to } b$$

$$\Rightarrow 0 = \frac{-j\omega\mu}{h^2} B C_1 C_3 \sin Ba \cos Ay \quad \forall y \rightarrow 0 \text{ to } b$$

$$\cos Ay \neq 0, \quad C_1, C_3 \neq 0$$

$$\therefore \sin Ba = 0 \Rightarrow Ba = m\pi \quad \text{where } m = 0, 1, 2, \dots, \infty$$

$$\therefore \boxed{B = \frac{m\pi}{a}} \longrightarrow (11)$$

from eqn (9),

$$\boxed{H_z = C_1 C_3 \cos\left(\frac{m\pi}{a}\right)x \cos\left(\frac{n\pi}{b}\right)y}$$

Let $C_1 C_3 = C$ (another constant)

$$\boxed{H_z = C \cos\left(\frac{m\pi}{a}\right)x \cos\left(\frac{n\pi}{b}\right)y \cdot e^{j\omega t - \gamma z}} \longrightarrow (12)$$

Field Components:-

$$E_x = \frac{-\gamma}{h^2} \frac{\partial E_z}{\partial x} - \frac{j\omega\mu}{h^2} \frac{\partial H_z}{\partial y}$$

Since $E_x = 0$ for TE wave,

$$\Rightarrow E_x = \frac{j\omega\mu}{h^2} c \cdot \left(\frac{n\pi}{b}\right) \cos\left(\frac{n\pi}{a}\right) x \sin\left(\frac{n\pi}{b}\right) y \cdot e^{(j\omega t - \gamma z)} \longrightarrow (13)$$

$$E_y = \frac{-\gamma}{h^2} \frac{\partial E_z}{\partial y} + \frac{j\omega\mu}{h^2} \frac{\partial H_z}{\partial x}$$

Since $E_x = 0$ for TE wave

$$\Rightarrow E_y = \frac{j\omega\mu}{h^2} \frac{\partial H_z}{\partial x}$$

$$E_y = -\frac{j\omega\mu}{h^2} c \left[\frac{m\pi}{a}\right] \sin\left(\frac{m\pi}{a}\right) x \cos\left(\frac{n\pi}{b}\right) y \cdot e^{(j\omega t - \gamma z)} \longrightarrow (14)$$

Similarly,

$$H_x = \frac{-\gamma}{h^2} \frac{\partial H_z}{\partial x} + \frac{j\omega\epsilon}{h^2} \frac{\partial E_z}{\partial y}$$

$$\therefore H_x = \frac{\gamma}{h^2} c \left(\frac{m\pi}{a}\right) \sin\left(\frac{m\pi}{a}\right) x \cos\left(\frac{n\pi}{b}\right) y \cdot e^{(j\omega t - \gamma z)} \longrightarrow (15)$$

and

$$H_y = \frac{-\gamma}{h^2} \frac{\partial H_z}{\partial y} - \frac{j\omega\epsilon}{h^2} \frac{\partial E_z}{\partial x}$$

$$\therefore H_y = \frac{-\gamma}{h^2} c \cdot \left(\frac{n\pi}{b}\right)^2 \cos\left(\frac{m\pi}{a}\right) x \cdot \sin\left(\frac{n\pi}{b}\right) y \cdot e^{j\omega t - \gamma z} \longrightarrow (16)$$

TM mode Analysis:-

for TM wave, $H_z = 0$; $E_z \neq 0$

The wave equation of a TM wave is

$$\frac{\partial^2 E_z}{\partial x^2} + \frac{\partial^2 E_z}{\partial y^2} + h^2 E_z = 0 \longrightarrow (1)$$

This is a partial differential equation which can be solved to get the different field components E_x , E_y , H_x and H_y by variable separable method.

Let us assume a solution

$$E_z = xy \longrightarrow (2)$$

where x is a pure function of x only

y is a pure function of y only

Since x and y are independent variables,

$$\frac{\partial^2 E_z}{\partial x^2} = \frac{\partial^2 (xy)}{\partial x^2} = y \frac{\partial^2 x}{\partial x^2}$$

$$\frac{\partial^2 E_z}{\partial y^2} = \frac{\partial^2 (xy)}{\partial y^2} = x \frac{\partial^2 y}{\partial y^2}$$

Using these two equations, from eqn (1) we get

$$y \frac{\partial^2 x}{\partial x^2} + x \frac{\partial^2 y}{\partial y^2} + h^2 xy = 0 \longrightarrow (3)$$

By dividing the above equation with xy , we get

$$\frac{1}{x} \frac{d^2 x}{dx^2} + \frac{1}{y} \frac{d^2 y}{dy^2} + h^2 = 0 \longrightarrow (4)$$

$\frac{1}{x} \frac{d^2 x}{dx^2}$ is a pure function of x only

$\frac{1}{y} \frac{d^2 y}{dy^2}$ is a pure function of y only.

The sum of these is a constant. Hence each term must be equal to a constant separately since x and y are independent variables.

By using variable separable method,

Let $\frac{1}{x} \frac{d^2 x}{dx^2} = -B^2$ and $\frac{1}{y} \frac{d^2 y}{dy^2} = -A^2$ where $-A^2$ & $-B^2$ are constants

$$\longrightarrow (5) \qquad \longrightarrow (6)$$

from eqns (4), (5) and (6),

$$-B^2 - A^2 + h^2 = 0 \Rightarrow h^2 = A^2 + B^2 \longrightarrow (7)$$

The solutions of eqn (5) & eqn (6) are

$$x = C_1 \cos Bx + C_2 \sin Bx \longrightarrow (8)$$

$$y = C_3 \cos Ay + C_4 \sin Ay \longrightarrow (9)$$

where C_1, C_2, C_3 and C_4 are constants which can be evaluated by applying the boundary conditions.

from eqn (1), $E_z = x \cdot y$

Substitute eqns (8) & (9) in eqn (10).

$$E_z = [C_1 \cos Bx + C_2 \sin Bx] [C_3 \cos Ay + C_4 \sin Ay] \longrightarrow (10)$$

Boundary conditions:

Since the entire surface of the rectangular waveguide acts as a short-circuit or ground for electric field, $E_z = 0$ all along the boundary walls of the waveguide.

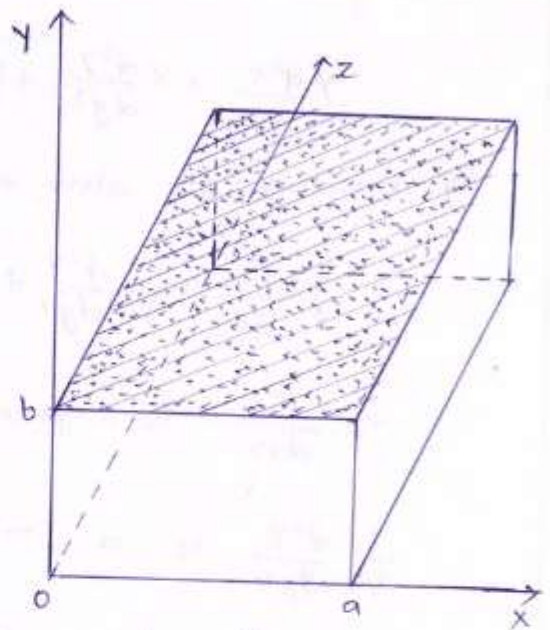
1st Boundary condition: [Bottom plane or Bottom wall]

Since $E_z = 0$ all along the bottom wall
i.e., $E_z = 0$ at $y = 0$ & $x \rightarrow 0$ to a .

2nd Boundary condition: [Left side plane or left side wall]
 $E_z = 0$ at $x = 0$ for all $y \rightarrow 0$ to b .

3rd Boundary condition: [Top plane or top wall]
 $E_z = 0$ at $y = b$ for all $x \rightarrow 0$ to a .

4th Boundary condition: [Right side plane or right side wall]
 $E_z = 0$ at $x = a$ for all $y \rightarrow 0$ to b .



(i) By substituting 1st Boundary condition in eqn (10), we get

$$0 = E_z = [C_1 \cos Bx + C_2 \sin Bx] [C_3 \cos A \cdot 0 + C_4 \sin A \cdot 0]$$

$$\Rightarrow [C_1 \cos Bx + C_2 \sin Bx] C_3 = 0 \quad [\because \cos 0 = 1, \sin 0 = 0]$$

$$C_1 \cos Bx + C_2 \sin Bx \neq 0, \therefore \boxed{C_3 = 0}$$

Therefore E_z becomes

$$E_z = [C_1 \cos Bx + C_2 \sin Bx] C_4 \sin Ay \longrightarrow (11)$$

(ii) By substituting 2nd Boundary condition in eqn (11), we get

$$E_z = 0 = [C_1 \cos B \cdot 0 + C_2 \sin B \cdot 0] C_4 \sin Ay$$

$$C_1 C_4 \sin Ay = 0$$

$$C_1 \neq 0, C_4 \neq 0 \quad C_4 \neq 0, \sin Ay \neq 0$$

$$\therefore \boxed{C_1 = 0}$$

$$\therefore E_z = C_2 C_4 \sin Bx \sin Ay \longrightarrow (12)$$

(iii) Substituting 3rd boundary condition in (12), we get

$$E_z = 0 = C_2 C_4 \sin Bx \sin Ay$$

$\sin Bx \neq 0$, $C_4 \neq 0$, $C_2 \neq 0$, otherwise there would be no solution.

$$\sin Ab = 0 \Rightarrow Ab = n\pi \quad n=0,1,2,\dots$$

$$\boxed{A = \frac{n\pi}{b}} \longrightarrow (13)$$

(iv) Substituting 4th boundary condition in Eqn (12), we get

$$E_z = 0 = C_2 C_4 \sin Ba \sin Ay$$

Since $C_2 \neq 0$, $C_4 \neq 0$, $\sin Ay \neq 0$.

$$\therefore \sin Ba = 0$$

$$\Rightarrow Ba = m\pi \quad \text{where } m=0,1,2,\dots$$

$$\boxed{B = \frac{m\pi}{a}} \longrightarrow (14)$$

from (12), (13) and (14)

$$E_z = C_2 C_4 \sin\left(\frac{m\pi}{a}x\right) \sin\left(\frac{n\pi}{b}y\right) e^{-\gamma x} e^{j\omega t}$$

$e^{-\gamma x} \rightarrow$ propagation along z direction.

$e^{j\omega t} \rightarrow$ sinusoidal variation w.r.t. t

Let $C = C_2 C_4$.

$$\therefore \boxed{E_z = C \sin\left(\frac{m\pi}{a}x\right) \sin\left(\frac{n\pi}{b}y\right) e^{j(\omega t - \gamma z)}} \longrightarrow (15)$$

$$E_x = -\frac{\gamma}{h^2} \frac{\partial E_z}{\partial x} - \frac{j\omega\mu}{h^2} \frac{\partial H_z}{\partial y}$$

$$\therefore \boxed{E_x = -\frac{\gamma}{h^2} \frac{\partial E_z}{\partial x}}$$

$$E_x = -\frac{\gamma}{h^2} C \left[\frac{m\pi}{a}\right] \cos\left(\frac{m\pi}{a}x\right) \sin\left(\frac{n\pi}{b}y\right) e^{j\omega t - \gamma z} \longrightarrow (16)$$

and

$$E_y = -\frac{\gamma}{h^2} \frac{\partial E_z}{\partial y} + \frac{j\omega\mu}{h^2} \frac{\partial H_z}{\partial x}$$

$$\therefore \boxed{E_y = -\frac{\gamma}{h^2} C \left[\frac{n\pi}{b}\right] \sin\left(\frac{m\pi}{a}x\right) \cos\left(\frac{n\pi}{b}y\right) e^{j\omega t - \gamma z}} \longrightarrow (17)$$

$$H_x = -\gamma/h^2 \frac{\partial H_z}{\partial x} + \frac{j\omega\epsilon}{h^2} \frac{\partial E_z}{\partial y}$$

$$H_x = j\omega\epsilon/h^2 C \left(\frac{n\pi}{b}\right) \sin\left(\frac{m\pi}{a}x\right) \cos\left(\frac{n\pi}{b}y\right) e^{j\omega t - \gamma z} \longrightarrow (18)$$

$$H_y = -\frac{r}{h^2} \frac{\partial H_z}{\partial y} + \frac{j\omega\epsilon}{h^2} \frac{\partial E_z}{\partial x}$$

$$\therefore H_y = \frac{j\omega\epsilon}{h^2} c \left[\frac{m\pi}{a} \right] \cos\left(\frac{m\pi}{a}x\right) \sin\left(\frac{n\pi}{b}y\right) e^{j\omega t - rz} \rightarrow (19)$$

Cut-off frequency of a Waveguide: —

Since we have,

$$r^2 + \omega^2\mu\epsilon = h^2 = A^2 + B^2$$

$$A = \frac{m\pi}{b}, \quad B = \frac{n\pi}{a}$$

$$r^2 = \left(\frac{m\pi}{a}\right)^2 + \left(\frac{n\pi}{b}\right)^2 - \omega^2\mu\epsilon$$

$$\Rightarrow \boxed{r = \sqrt{\left(\frac{m\pi}{a}\right)^2 + \left(\frac{n\pi}{b}\right)^2 - \omega^2\mu\epsilon} = \alpha + j\beta}$$

At lower frequencies, ~~also~~ $r > 0 \Rightarrow \sqrt{\left(\frac{m\pi}{a}\right)^2 + \left(\frac{n\pi}{b}\right)^2 - \omega^2\mu\epsilon} > 0$

$$\omega^2\mu\epsilon < \left(\frac{m\pi}{a}\right)^2 + \left(\frac{n\pi}{b}\right)^2$$

'r' then becomes real and positive and equal to the attenuation constant α i.e. the wave is completely attenuated and there is no phase change. Hence the wave cannot propagate.

However, at higher frequencies, $r < 0$

$$\sqrt{\left(\frac{m\pi}{a}\right)^2 + \left(\frac{n\pi}{b}\right)^2 - \omega^2\mu\epsilon} < 0 \Rightarrow \left(\frac{m\pi}{a}\right)^2 + \left(\frac{n\pi}{b}\right)^2 < \omega^2\mu\epsilon$$

'r' becomes imaginary, there will be phase change β and hence the wave propagates.

At the transition, 'r' becomes zero and the propagation starts. The frequency at which 'r' just becomes zero is defined as the cut-off frequency (a threshold frequency) ' f_c '.

i.e. At $f = f_c$, $r = 0$ or $\omega = 2\pi f = 2\pi f_c = \omega_c$

$$\therefore 0 = \left(\frac{m\pi}{a}\right)^2 + \left(\frac{n\pi}{b}\right)^2 - \omega_c^2\mu\epsilon \quad \text{or} \quad \omega_c = \frac{1}{\sqrt{\mu\epsilon}} \left[\left(\frac{m\pi}{a}\right)^2 + \left(\frac{n\pi}{b}\right)^2 \right]^{\frac{1}{2}}$$

$$f_c = \frac{1}{2\pi\sqrt{\mu\epsilon}} \left[\left(\frac{m\pi}{a} \right)^2 + \left(\frac{n\pi}{b} \right)^2 \right]^{1/2}$$

$$f_c = \frac{c}{2} \left[\left(\frac{m}{a} \right)^2 + \left(\frac{n}{b} \right)^2 \right]^{1/2}$$

The Cut-off wavelength (λ_c) is

$$\lambda_c = \frac{c}{f_c} = \frac{c}{\frac{c}{2} \left[\left(\frac{m}{a} \right)^2 + \left(\frac{n}{b} \right)^2 \right]^{1/2}}$$

$$\lambda_{c,m,n} = \frac{2ab}{\sqrt{m^2b^2 + n^2a^2}}$$

$$\therefore c = \frac{1}{\sqrt{\mu\epsilon}}$$

$$\therefore v_g = \lambda_g f = \frac{\lambda_0 f}{\sqrt{1-(f/f_c)^2}} = \frac{v_p}{\sqrt{1-(f/f_c)^2}}$$

$$\Rightarrow \lambda_g = \frac{\lambda_0}{\sqrt{1-(f/f_c)^2}} \quad \therefore f_c = \frac{v}{\lambda_c} \quad \& \quad f = \frac{v}{\lambda_0}$$

$$\therefore \lambda_g = \frac{\lambda_0}{\sqrt{1-(\lambda_0/\lambda_c)^2}} \quad \left\{ \begin{array}{l} \frac{\lambda_g^2 \lambda_c^2 - \lambda_0^2 \lambda_g^2}{\lambda_0^2 \lambda_c^2 \lambda_g^2} = \frac{\lambda_0^2 \lambda_c^2 - \lambda_0^2 \lambda_g^2}{\lambda_0^2 \lambda_c^2 \lambda_g^2} \\ \lambda_g^2 \left[1 - \left(\frac{\lambda_0}{\lambda_c} \right)^2 \right] = \lambda_0^2 \end{array} \right.$$

$$\Rightarrow \lambda_g^2 \lambda_c^2 - \lambda_0^2 \lambda_g^2 = \lambda_0^2 \lambda_c^2$$

By dividing with $\lambda_0^2 \lambda_c^2 \lambda_g^2$ both sides $\Rightarrow \frac{1}{\lambda_g^2} = \frac{1}{\lambda_0^2} - \frac{1}{\lambda_c^2}$

All wavelengths greater than λ_c are attenuated and those less than λ_c are allowed to propagate inside the waveguide.

Guide Wavelength (λ_g): —

It is defined as the distance travelled by the wave in order to undergo a phase shift of 2π radians.

It is related to phase constant by the relation

$$\lambda_g = \frac{2\pi}{\beta}$$

$$v_g = \frac{v_p}{\sqrt{1-(f/f_c)^2}}$$

$$\therefore v_g = \beta \lambda_g = \pm j \omega \sqrt{\mu\epsilon} [1-(f/f_c)^2]$$

$$\therefore \beta \lambda_g = \omega \sqrt{\mu\epsilon} \sqrt{1-(f/f_c)^2}$$

$$\therefore v_g = \frac{\omega}{\beta} \Rightarrow \beta \lambda_g = \frac{\omega}{v_g}$$

$$\frac{\omega}{v_g} = \omega \sqrt{\mu\epsilon} \sqrt{1-(f/f_c)^2} = \frac{\omega \sqrt{1-(f/f_c)^2}}{v_p}$$

The wavelength in the waveguide is different from the wavelength in free space. Guide wavelength is related to free space wavelength λ_0 and cut-off wavelength λ_c by

$$\frac{1}{\lambda_g^2} = \frac{1}{\lambda_0^2} - \frac{1}{\lambda_c^2}$$

$$\text{or } \lambda_g = \frac{\lambda_0}{\sqrt{1-(\lambda_0/\lambda_c)^2}}$$

The above equation is true for any mode in a waveguide of any cross section

→ It is clear that if $\lambda_0 \ll \lambda_c$, then $\lambda_g = \lambda_0$

→ when $\lambda_0 = \lambda_c \Rightarrow \lambda_g = \infty$

→ when $\lambda_0 > \lambda_c \Rightarrow \lambda_g$ is imaginary which is nothing but no propagation in the waveguide.

Phase velocity (v_p):

Wave propagates in the waveguide when guide wavelength λ_g is greater than the free space wavelength λ_0 .

In a waveguide, $v_p = \lambda_g f$ where v_p is the phase velocity. But the speed of light is equal to product of λ_0 and f . This v_p is greater than the speed of light since $\lambda_g > \lambda_0$.

→ The wavelength in the guide is the length of the cycle and v_p represents the velocity of the phase.

→ It is defined as the "rate at which the wave changes its phase in terms of the guide wavelength".

$$\text{i.e., } v_p = \frac{\lambda_g}{\text{unit time}} = \lambda_g \cdot f = \frac{2\pi f \cdot \lambda_g}{2\pi} = \frac{2\pi f}{2\pi/\lambda_g}$$

$$\text{i.e., } v_p = \frac{\omega}{\beta}$$

$$\text{where } \omega = 2\pi f, \quad \beta = \frac{2\pi}{\lambda_g}$$

$$v_p = \frac{\omega}{\omega \sqrt{\mu\epsilon} [1 - (fc/f)^2]^{1/2}} = \frac{c}{\sqrt{[1 - (fc/f)^2]}}$$

$$v_p = \frac{c}{\sqrt{1 - [\lambda_0/\lambda_c]^2}}$$

Group velocity (v_g):

The rate at which the wave propagates through the waveguide and is given by

$$v_g = \frac{d\omega}{d\beta}$$

$$\text{Since } \beta = \sqrt{\mu\epsilon(\omega^2 - \omega_c^2)}$$

now differentiating β w.r.t. ω , we get

$$\frac{d\beta}{d\omega} = \frac{1}{2\sqrt{\mu\epsilon(\omega^2 - \omega_c^2)}} \cdot 2\omega\mu\epsilon = \frac{\omega\mu\epsilon}{\omega\sqrt{\mu\epsilon[1 - \omega_c^2/\omega^2]}} = \frac{\sqrt{\mu\epsilon}}{\sqrt{1 - (fc/f)^2}}$$

$$\therefore v_g = \frac{\sqrt{1 - (fc/f)^2}}{\sqrt{\mu\epsilon}} = c \sqrt{1 - (\lambda_0/\lambda_c)^2}$$

$$\text{Since } c = \frac{1}{\sqrt{\mu\epsilon}}$$

$$v_g = c \sqrt{1 - \left(\frac{\lambda_0}{\lambda_c}\right)^2}$$

1-1-12

consider the product of v_p and v_g

$$\therefore v_p \cdot v_g = \frac{c}{\sqrt{1-(\lambda_0/\lambda_c)^2}} \cdot c \cdot \sqrt{1-(\lambda_0/\lambda_c)^2}$$

$$\therefore \boxed{v_p \cdot v_g = c^2}$$

Since guide wavelength, $\lambda_g = \frac{\lambda_0}{\sqrt{1-(\lambda_0/\lambda_c)^2}} \Rightarrow \sqrt{1-(\lambda_0/\lambda_c)^2} = \frac{\lambda_0}{\lambda_g}$

$$\therefore \boxed{v_p = \frac{\lambda_g}{\lambda_0} \cdot c}$$

$$\therefore v_p = \frac{c}{\sqrt{1-(\lambda_0/\lambda_c)^2}} = \frac{c}{\lambda_0/\lambda_g}$$

$$\boxed{v_p = \frac{\lambda_g}{\lambda_0} \cdot c}$$

The group velocity $v_g = c \sqrt{1-(\lambda_0/\lambda_c)^2}$

$$\Rightarrow \boxed{v_g = \frac{\lambda_0}{\lambda_g} \cdot c}$$

$$\boxed{v_g = c \cdot \frac{\lambda_0}{\lambda_g}}$$

$$\therefore \boxed{v_p \cdot v_g = c^2} = v_c^2$$

Relation between λ_g , λ_0 and λ_c :

we know that $v_p = \lambda_g \cdot f = \frac{\lambda_g}{\lambda_0} \cdot c$

$$v_p = \frac{c}{\sqrt{1-(\lambda_0/\lambda_c)^2}}$$

$$\therefore \frac{c}{\sqrt{1-(\lambda_0/\lambda_c)^2}} = \frac{\lambda_g}{\lambda_0} \cdot c$$

$$\therefore \boxed{\lambda_g = \frac{\lambda_0}{\sqrt{1-(\lambda_0/\lambda_c)^2}}}$$

and

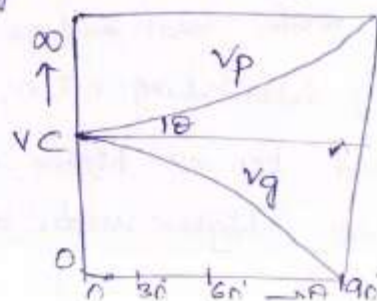
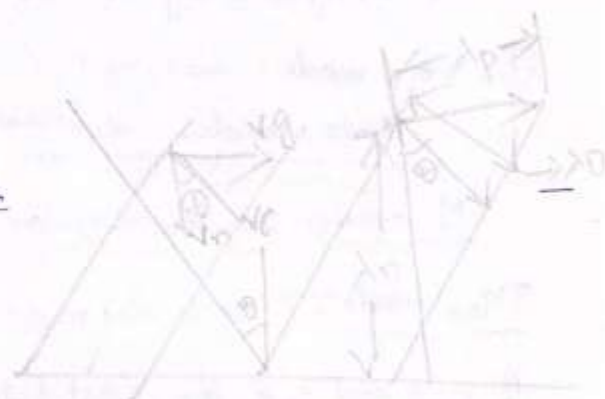
$$v_g = v_c \sin \theta$$

$$v_p = \frac{v_c}{\sin \theta}$$

$$\therefore v_p \cdot v_g = v_c^2 = c^2$$

→ phase velocity is the velocity with which the wave changes phase at a plane boundary and not the velocity with which it travels along the boundary.

→ At cutoff frequencies, v_p becomes infinite and v_g becomes zero.



TE modes in Rectangular Waveguides:

TE_{mn} is the general mode and the specific modes are given by various values of m and n . They are

(a) TE_{00} mode: $m=0, n=0$

By substituting m & n values all field components vanishes therefore it cannot exist.

(b) TE_{01} mode: $m=0, n=1$

$E_y=0, H_x=0, E_x$ & H_y are exists.

(c) TE_{10} mode: $m=1, n=0$

$\Rightarrow E_x=0, H_y=0, E_y$ and H_x are exists. Therefore TE_{10} mode exists.

(d) TE_{11} mode: $m=1, n=1$

This mode also exists.

TE field components:

$$E_z = 0$$
$$E_y = \frac{j\omega\mu}{h^2} C \cdot \left(\frac{n\pi}{b}\right) \cos\left(\frac{m\pi}{a}\right)x \sin\left(\frac{n\pi}{b}\right)y e^{j(\omega t - \gamma z)}$$

$$E_x = -\frac{j\omega\mu}{h^2} C \left[\frac{m\pi}{a}\right] \sin\left(\frac{m\pi}{a}\right)x \cos\left(\frac{n\pi}{b}\right)y e^{j(\omega t - \gamma z)}$$

$$H_x = \frac{\gamma}{h^2} C \left[\frac{m\pi}{a}\right] \sin\left(\frac{m\pi}{a}\right)x \cos\left(\frac{n\pi}{b}\right)y e^{j(\omega t - \gamma z)}$$

$$H_y = -\frac{\gamma}{h^2} C \left[\frac{n\pi}{b}\right] \cos\left(\frac{m\pi}{a}\right)x \sin\left(\frac{n\pi}{b}\right)y e^{j\omega t - \gamma z}$$

$$H_z = \frac{C}{h^2} \cos\left(\frac{m\pi}{a}\right)x \cos\left(\frac{n\pi}{b}\right)y e^{j(\omega t - \gamma z)}$$

TM modes in Rectangular Waveguides:

TM_{00} mode: $m=0$ and $n=0$

If $m=0$ and $n=0$ are substituted then all the field components vanishes hence TM_{00} mode cannot exist.

TM_{10} mode: $m=1, n=0$

By substituting m & n , all field components vanishes. So TM_{10} mode also cannot exist.

TM_{01} mode: $m=0, n=1$

By substituting m & n , all field components vanishes so TM_{01} mode cannot exist.

TM_{11} mode: $m=1$ and $n=1$

By substituting m & n , all E_x, E_y, H_x and H_y i.e. TM_{11} mode exists

and for all higher values of m and n , the components exist i.e. all higher modes do exist.

TM field components:

$$E_x = -\frac{\gamma}{h^2} C \left[\frac{m\pi}{a}\right] \cos\left(\frac{m\pi}{a}\right)x \sin\left(\frac{n\pi}{b}\right)y e^{j\omega t - \gamma z}$$

$$E_y = \frac{\gamma}{h^2} C \left[\frac{n\pi}{b}\right] \sin\left(\frac{m\pi}{a}\right)x \cos\left(\frac{n\pi}{b}\right)y e^{j\omega t - \gamma z}$$

$$H_x = \frac{j\omega\epsilon}{h^2} C \left[\frac{n\pi}{b}\right] \sin\left(\frac{m\pi}{a}\right)x \cos\left(\frac{n\pi}{b}\right)y e^{j\omega t - \gamma z}$$

$$H_y = \frac{j\omega\epsilon}{h^2} C \left[\frac{m\pi}{a}\right] \cos\left(\frac{m\pi}{a}\right)x \sin\left(\frac{n\pi}{b}\right)y e^{j\omega t - \gamma z}$$

$$E_z = 0 \sin\left[\frac{m\pi}{a}\right]x \sin\left[\frac{n\pi}{b}\right]y e^{j\omega t - \gamma z}$$

$$H_z = 0$$

(13)

Dominant Mode:

The mode for which the cut-off wavelength (λ_c) assumes a maximum value.

$$\text{Since } \lambda_{c_{mn}} = \frac{2ab}{\sqrt{m^2b^2 + n^2a^2}}$$

Dominant mode in TE: —

$$\text{for TE}_{01} \text{ mode, } \lambda_{c01} = \frac{2ab}{\sqrt{0^2 + a^2}} = \underline{2b}$$

$$\text{for TE}_{10} \text{ mode, } \lambda_{c10} = \frac{2ab}{\sqrt{b^2 + 0}} = \underline{2a}$$

$$\text{for TE}_{11} \text{ mode, } \lambda_{c11} = \frac{2ab}{\sqrt{a^2 + b^2}}$$

$$\text{for TE}_{21} \text{ mode, } \lambda_{c21} = \frac{2ab}{\sqrt{4b^2 + a^2}}$$

Among all λ_{c10} has the maximum value since 'a' is the longer dimension than 'b'. Hence TE₁₀ mode is the dominant mode in rectangular waveguides.

$$\beta = \frac{2\pi}{\lambda_g} = \sqrt{k^2 \mu \epsilon - \omega^2 \mu \epsilon}$$

$$v_g = c \sqrt{1 - (\lambda_0/\lambda_c)^2}$$

$$v_p = \frac{c}{\sqrt{1 - (\lambda_0/\lambda_c)^2}}$$

$$\lambda_g = \frac{\lambda_0}{\sqrt{1 - (\lambda_0/\lambda_c)^2}}$$

Dominant mode in TM: —

Minimum possible mode is TM₁₁. Higher modes than this also exist.

$$\text{for TM}_{11} \text{ mode, } \lambda_{c11} = \frac{2ab}{\sqrt{m^2b^2 + n^2a^2}} = \frac{2ab}{\sqrt{a^2 + b^2}}$$

$$\text{for TM}_{12} \text{ mode, } \lambda_{c12} = \frac{2ab}{\sqrt{b^2 + 4a^2}}$$

$$\text{for TM}_{21} \text{ mode, } \lambda_c = \frac{2ab}{\sqrt{a^2 + 4b^2}}$$

By observation it is clear that $\lambda_{c11} > \lambda_{c12}$ and $\lambda_{c11} > \lambda_{c21}$ and so on.

→ The minimum cut-off frequency for a rectangular waveguide is obtained for a dimension $a \times b$ for $m=1$ and $n=0$ i.e., TE_{10} mode is the dominant mode for a rectangular waveguide. Some of the higher order modes, having the same cut-off frequency/wavelength are called 'Degenerate Modes'.

→ Degenerate Modes: Two or more modes having the same cut-off frequency/wavelength are called 'Degenerate Modes'.

→ for a rectangular waveguide TE_{mn}/TM_{mn} modes for which both $m \neq 0$, $n \neq 0$ will always be degenerate modes.

→ for a square guide in which $a=b$, all the TE_{mn} , TE_{nm} and TM_{mn} , TM_{nm} modes are together degenerate modes.

TE and TM mode

Wavelengths and Impedance Relations: — (in TE & TM Waves)

→ Guide Wavelength (λ_g):

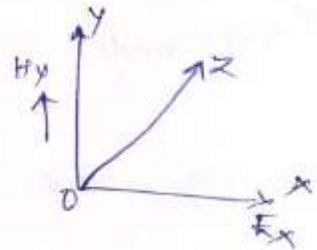
It is defined as the distance travelled by the wave in order to undergo a phase shift of 2π radians.

$$\lambda_g = \frac{\lambda_0}{\sqrt{1 - (\lambda_0/\lambda_c)^2}} \Rightarrow \frac{1}{\lambda_g^2} = \frac{1}{\lambda_0^2} - \frac{1}{\lambda_c^2}$$

→ Wave impedance is defined as the ratio of the strength of electric field in one transverse direction to the strength of the magnetic field along the other transverse direction.

$$\text{i.e., } Z_z = \frac{E_x}{H_y} = -\frac{E_y}{H_x}$$

$$\text{or } Z_z = \frac{\sqrt{E_x^2 + E_y^2}}{\sqrt{H_x^2 + H_y^2}}$$



(1) Wave impedance for a TM wave in rectangular waveguide:

$$Z_z = Z_{TM} = \frac{E_x}{H_y} = \frac{-\frac{\gamma}{h^2} \frac{\partial E_z}{\partial x} - \frac{j\omega\mu}{h^2} \frac{\partial H_z}{\partial y}}{-\frac{\gamma}{h^2} \frac{\partial H_z}{\partial y} - \frac{j\omega\epsilon}{h^2} \frac{\partial E_z}{\partial x}}$$

for a TM wave, $H_z = 0$ and $\gamma = j\beta$

$$Z_{TM} = \frac{-\gamma/h^2 \frac{\partial E_z}{\partial x}}{\frac{j\omega\epsilon}{h^2} \frac{\partial E_z}{\partial x}} = \frac{\gamma}{j\omega\epsilon} = \frac{j\beta}{j\omega\epsilon}$$

$$\Rightarrow \boxed{Z_{TM} = \frac{\beta}{\omega\epsilon}}$$

Since we have $\beta = \sqrt{\omega^2\mu\epsilon - \omega_c^2\mu\epsilon}$

$$\therefore Z_{TM} = \frac{\sqrt{\omega^2\mu\epsilon - \omega_c^2\mu\epsilon}}{\omega\epsilon} = \frac{\sqrt{\mu\epsilon} \cdot \sqrt{\omega^2 - \omega_c^2}}{\omega\epsilon} = \sqrt{\frac{\mu}{\epsilon}} \sqrt{\frac{\omega^2 - \omega_c^2}{\omega^2}}$$

$$= \sqrt{\frac{\mu}{\epsilon}} \sqrt{1 - (\omega_c/\omega)^2} = \sqrt{\frac{\mu}{\epsilon}} \sqrt{1 - (\lambda_0/\lambda)^2}$$

$$\boxed{Z_{TM} = \sqrt{\frac{\mu}{\epsilon}} \sqrt{1 - (\lambda_0/\lambda)^2}}$$

for air, $\sqrt{\frac{\mu}{\epsilon}} = \sqrt{\frac{\mu_0\mu_r}{\epsilon_0\epsilon_r}} = \sqrt{\frac{4\pi \times 10^{-7}}{\frac{1}{36\pi} \times 10^9}} = \sqrt{4\pi \times 36\pi \times 100} = 2 \times 6\pi \times 10$

$$= 120\pi = 377 \Omega = \eta = \text{intrinsic impedance of free space}$$

$$\boxed{Z_{TM} = \eta \sqrt{1 - (\lambda_0/\lambda)^2}}$$

Since λ_0 is always less than λ for wave propagation $Z_{TM} < \eta$.
This shows that wave impedance for a TM wave is always less than free space impedance.

② Wave impedance of TE waves in Rectangular waveguide

$$Z_x = Z_{TE} = \frac{E_y}{H_x} = \frac{-\gamma/h^2 \frac{\partial E_z}{\partial x} - \frac{j\omega\mu}{h^2} \frac{\partial H_z}{\partial y}}{-\gamma/h^2 \frac{\partial H_z}{\partial y} - \frac{j\omega\epsilon}{h^2} \frac{\partial E_z}{\partial x}}$$

for a TE wave, $E_z = 0$ and $\gamma = j\beta$

$$\therefore Z_{TE} = \frac{-\frac{j\omega\mu}{h^2} \frac{\partial H_z}{\partial y}}{-\gamma/h^2 \frac{\partial H_z}{\partial y}} = \frac{j\omega\mu}{\gamma} = \frac{j\omega\mu}{j\beta} = \frac{\omega\mu}{\beta}$$

$$\boxed{Z_{TE} = \frac{\omega\mu}{\beta}}$$

Since $\beta = \sqrt{\omega^2 \mu \epsilon - \omega_c^2 \mu \epsilon}$

$$Z_{TE} = \frac{\omega \mu}{\sqrt{\omega^2 \mu \epsilon - \omega_c^2 \mu \epsilon}} = \frac{\omega \mu}{\omega \sqrt{\mu \epsilon} \sqrt{1 - (\frac{\omega_c}{\omega})^2}} = \sqrt{\frac{\mu}{\epsilon}} \frac{1}{\sqrt{1 - (f_c/f)^2}}$$

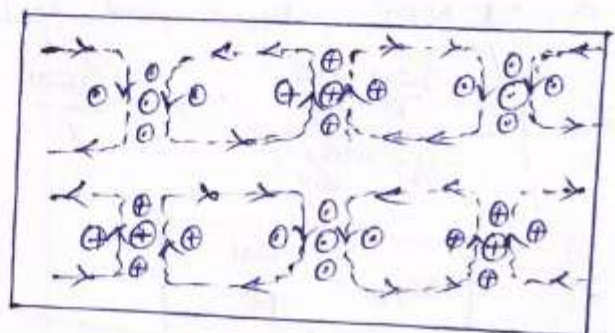
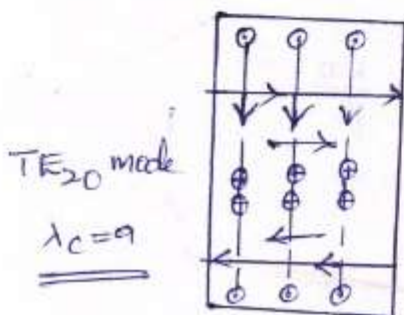
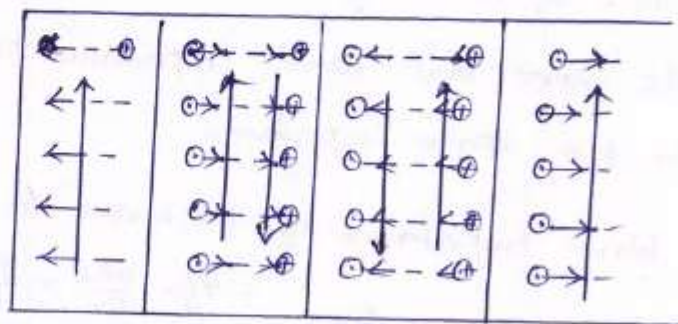
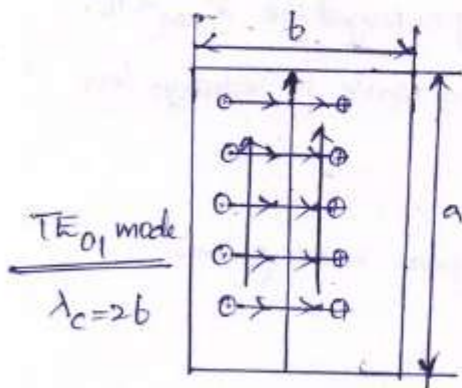
$$\therefore Z_{TE} = \frac{\eta}{\sqrt{1 - (f_c/f)^2}}$$

$$\therefore Z_{TE} = \frac{\eta}{\sqrt{1 - (\lambda_0/\lambda_c)^2}}$$

Therefore $Z_{TE} > \eta$ as $\lambda_0 < \lambda_c$ for wave propagation. This shows that wave impedance for a TE wave is always greater than free space impedance.

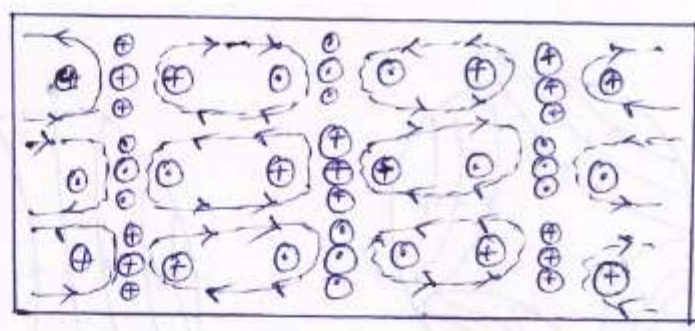
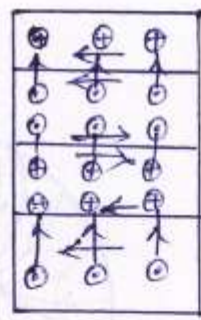
→ for TEM waves between parallel plates or an ordinary parallel wire or co-axial transmission lines the cut-off frequency is zero and wave impedance for TEM wave is the free space impedance itself.
i.e. $Z_0(\text{TEM}) = \eta$

TE/TM Mode fields: —

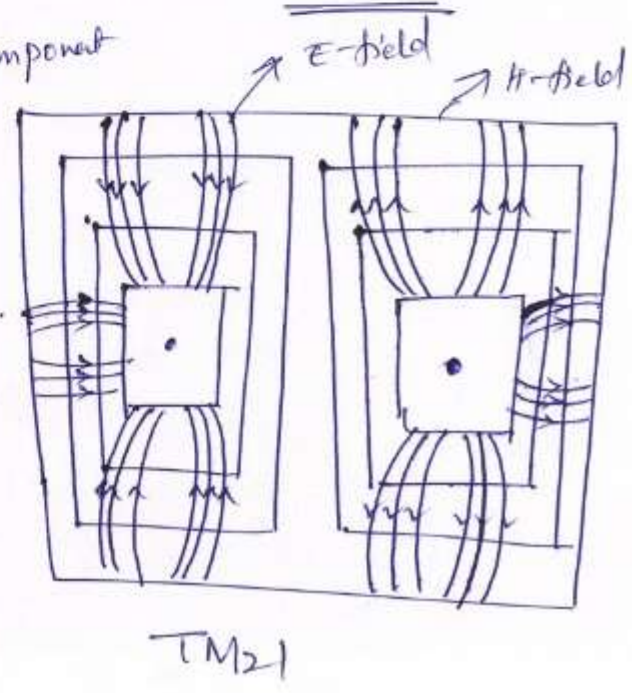
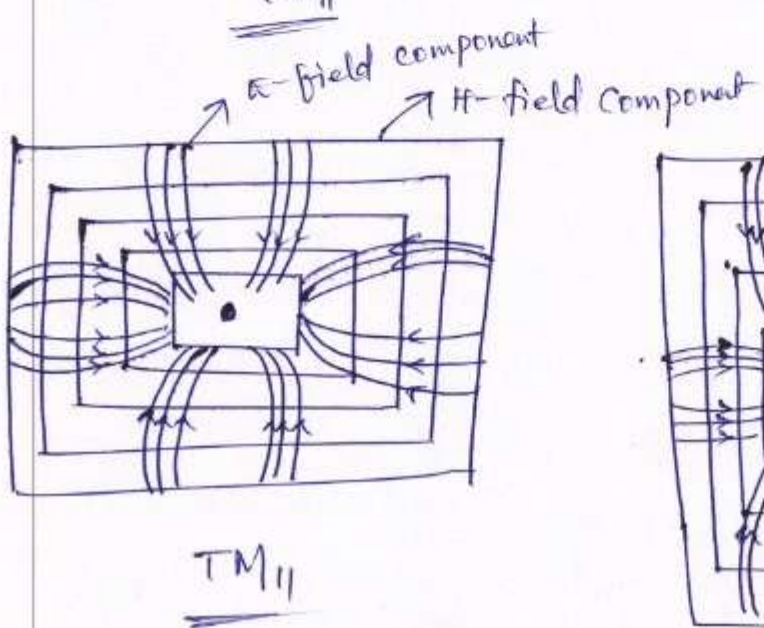
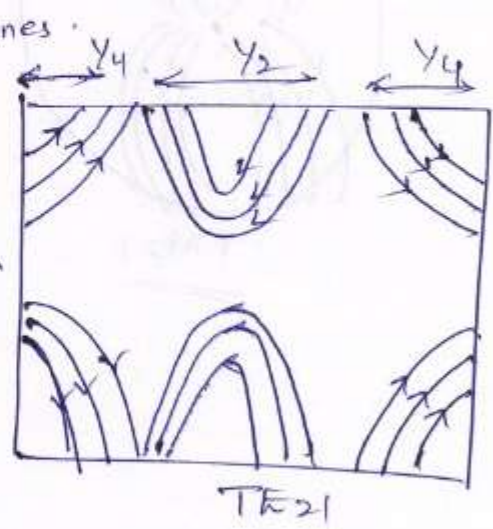
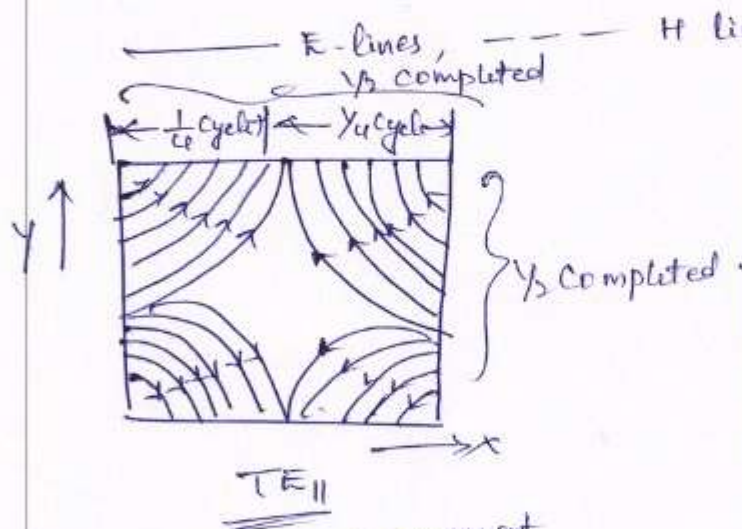
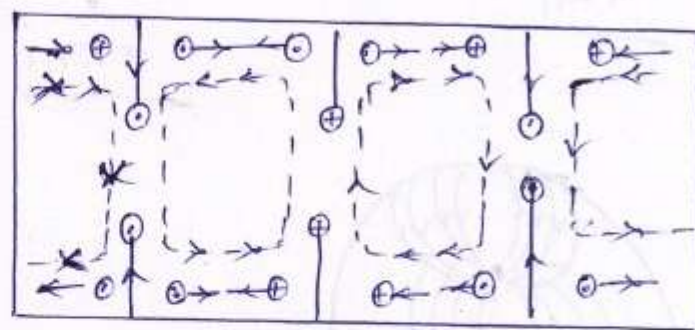
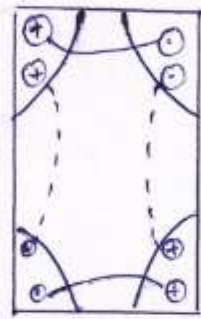


MNR
3-10

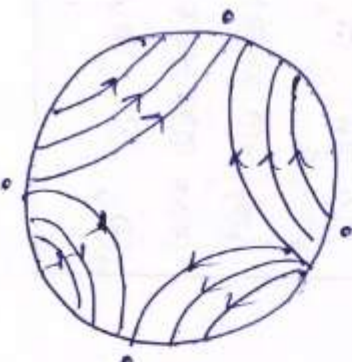
TE₃₀ mode



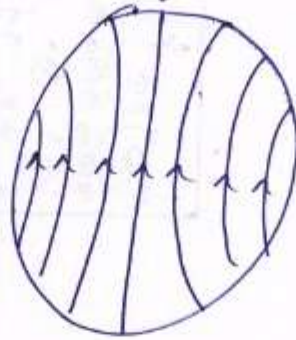
TE₁₁ mode



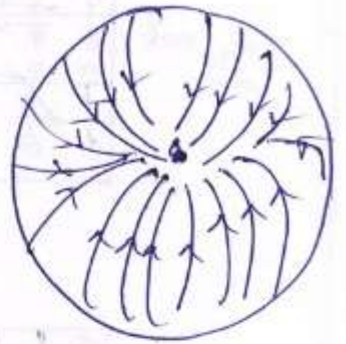
TE/TM in circular w/g:-



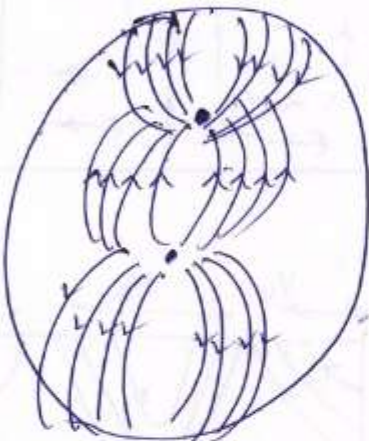
TE₁₁



TE₀₁



TM₀₁



TM₂₁

Rectangular GuidesPower Transmission:

The power transmitted through a waveguide and the power loss in guide walls can be calculated by means of complex Poynting Theorem. We assume that the waveguide is terminated in such a way that there is no reflection from the receiving end or that the waveguide is infinitely long as compared with its wavelength.

The power transmitted P_{ts} through a waveguide given by

$$P_{ts} = \oint P \cdot ds = \oint \frac{1}{2} (\mathbf{E} \times \mathbf{H}^*) \cdot d\mathbf{s}$$

for a lossless dielectric, the time average power flow through a rectangular waveguide is

$$P_{ts} = \frac{1}{2 Z_z} \int_0^a \int_0^b |\mathbf{E}|^2 dx dy = \frac{Z_z}{2} \int_0^a \int_0^b |\mathbf{H}|^2 dx dy$$

where $Z_z = \frac{E_x}{H_y} = -\frac{E_y}{H_x}$

$$|\mathbf{E}|^2 = |E_x|^2 + |E_y|^2, \quad |\mathbf{H}|^2 = |H_x|^2 + |H_y|^2$$

for TM_{mn} mode, the average power transmitted through a rectangular waveguide of dimensions a and b is

$$P_{ts} = \frac{1}{2\eta \sqrt{1 - (\lambda_0/\lambda_c)^2}} \int_0^b \int_0^a |E_x|^2 + |E_y|^2 dx dy$$

for TE_{mn} modes, $Z_z = \frac{\eta}{\sqrt{1 - (\lambda_0/\lambda_c)^2}}$

$$\therefore \boxed{P_{ts} = \frac{\sqrt{1 - (\lambda_0/\lambda_c)^2}}{2\eta} \int_0^b \int_0^a |E_x|^2 + |E_y|^2 dx dy}$$

Power losses in a Waveguide:

losses in a waveguide can be due to attenuation below cut-off and losses associated with attenuation due to dissipation within the waveguide walls and the dielectric within the waveguide.

At frequencies below the cut-off frequency ($f < f_c$), the propagation constant ' γ ' will have only the attenuation term ' α ', that is to say that the phase constant ' β ' itself becomes imaginary implying wave attenuation.

$$\text{Hence } \beta = \frac{j2\pi}{\lambda_g}$$

$$\lambda_g = \frac{\lambda}{\cos \theta} = \frac{\lambda}{\sqrt{1 - (f/f_c)^2}}$$

$$\text{therefore, } \beta = \frac{j2\pi}{\lambda} \sqrt{1 - (f/f_c)^2} = j \frac{2\pi f}{c} \sqrt{1 - (f/f_c)^2} = j\alpha$$

Hence the cut-off attenuation constant ' α ' is given by

$$\alpha = \frac{54.6}{\lambda_0} \sqrt{1 - (f/f_c)^2} \text{ dB/length.}$$

for $f > f_c$, the waveguide exhibits very low loss and for $f < f_c$, the attenuation is high and results in full reflection of the wave i.e. cut-off attenuation is basically the reflection loss.

Attenuation Constant due to an imperfect, non magnetic dielectric in the waveguide is given by

$$\alpha_d = \frac{27.3 \sqrt{\epsilon_r} \tan \delta}{\lambda_0 \sqrt{1 - (f/f_c)^2}} \text{ dB/length}$$

where $\tan \delta$ = dielectric loss tangent of the insulating material (dielectric).

The attenuation constant due to the imperfect conducting walls for TE₁₀ mode is given by

$$\alpha_c = \frac{R_s}{b \eta_0} = \frac{1 + \frac{2b}{a} \left(\frac{f}{f_c} \right)^2}{\sqrt{1 - (f/f_c)^2}} \text{ Np/length}$$

Where $R_s \rightarrow$ Sheet

Resistivity in ohm/m²

$\eta_0 \rightarrow$ intrinsic impedance of

free space (377 ohm)

$$R_s = \frac{1}{\sigma \delta_s}$$

Where $\sigma \rightarrow$ conductivity of the metallic walls in S/m

$$\text{Skin depth } \delta_s = \frac{1}{\sqrt{\pi f \mu_0 \mu_r \sigma}}$$

$f \rightarrow$ frequency

$\mu_r \rightarrow$ relative permeability

$\mu_0 \rightarrow$ permeability of free space ($4\pi \times 10^{-7} \text{ H/m}$)

Impossibility of TEM mode:

Field components exist in a waveguide are

$$E_x = -\frac{\gamma}{h^2} \frac{\partial E_z}{\partial x} - \frac{j\omega\mu}{h^2} \frac{\partial H_z}{\partial y}$$

$$E_y = -\frac{\gamma}{h^2} \frac{\partial E_z}{\partial y} + \frac{j\omega\mu}{h^2} \frac{\partial H_z}{\partial x}$$

$$H_x = -\frac{\gamma}{h^2} \frac{\partial H_z}{\partial x} + \frac{j\omega\epsilon}{h^2} \frac{\partial E_z}{\partial y}$$

$$H_y = -\frac{\gamma}{h^2} \frac{\partial H_z}{\partial y} - \frac{j\omega\epsilon}{h^2} \frac{\partial E_z}{\partial x}$$

In Transverse Electromagnetic [TEM] mode both electric field and magnetic fields exist and they are perpendicular to each other and their direction is perpendicular to the direction of wave. So no field component i.e. neither electric field nor magnetic field are in the direction of a wave.

Since we assume that the direction of wave is along $\hat{z}$. So $E_z = 0$ & $H_z = 0$.

By substituting E_z & H_z in the above expressions

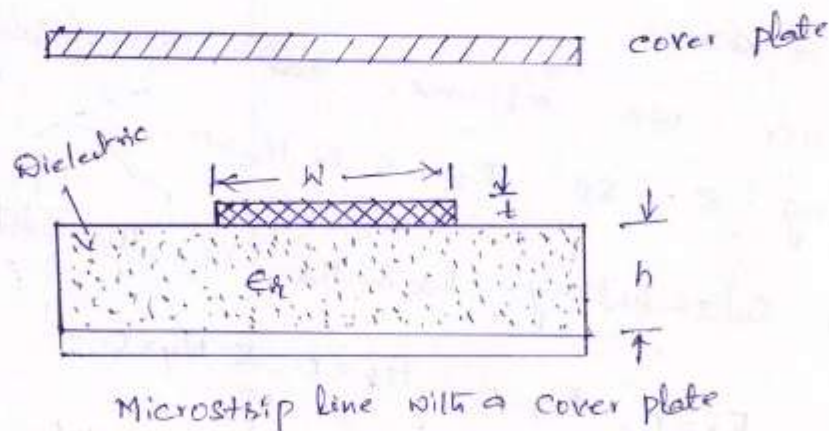
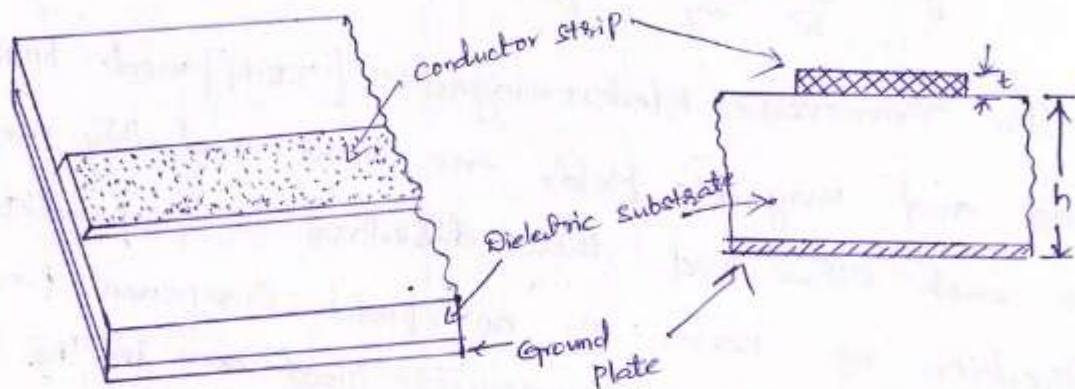
$$E_x = 0, E_y = 0, H_x = 0 \text{ & } H_y = 0$$

All field components vanish when TEM wave is incident into a waveguide.

The impossibility of TEM wave is it cannot exist in a waveguide.

Microstrip Lines:

Microstrip Line is an unsymmetrical stripline that is nothing but a parallel plate transmission line having dielectric substrate, the one face of which is metallised ground and the other face has a thin conducting strip of certain width w and thickness t . Sometimes a coverplate is used for shielding purposes but it is kept much farther away from the ground plane so as not to affect the microstrip field lines.



Advantages:-

- Complete conductor pattern may be deposited and processed on a single dielectric substrate which is supported by a single metal ground plane. Thus fabrication costs would be substantially lower than stripline co-axial or waveguide circuits.

→ Due to the planar nature of the microstrip structure, both packaged and unpackaged semiconductor chips can be conveniently attached to the microstrip element.

→ There is an easy access to the top surface making it easy to mount passive or active discrete devices.

Limitations:-

→ Radiation losses are high or interference due to nearby conductors.

→ Because of the proximity of the air-dielectric interface with the microstrip conductor at the interface, a discontinuity in the electric and magnetic fields is generated. This results in a microstrip configuration that becomes a mixed dielectric transmission structure with impure TEM modes propagating.

→ The electric flux crossing the air-dielectric boundary is small and although a pure TEM mode cannot exist, a small deviation from TEM mode does exist which can be neglected.

The characteristic impedance of a microstrip is a function of the strip line width (W), thickness (t) and the distance between the line and the ground plane (h).

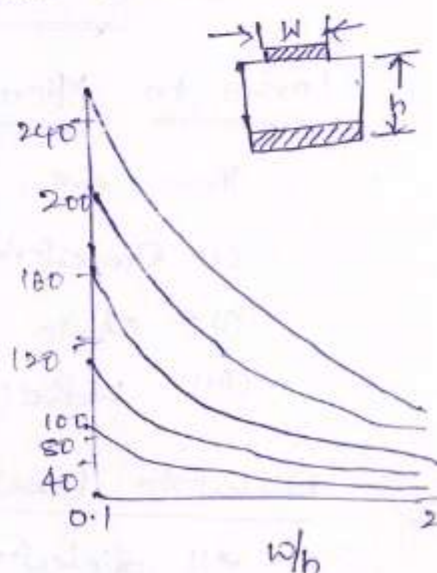
Empirical relation for Z_0 for a microstrip line is given by

$$Z_0 = \frac{60}{\sqrt{\epsilon_r}} \ln\left(\frac{4h}{d}\right) \text{ for } h \gg d.$$

Where, ϵ_r = dielectric constant of the dielectric medium

h = distance between the microstripline and the ground plane.

d = diameter of the wire.



Z_0 vs W/h ratios.

Effective dielectric constant is given by

$$\epsilon_{re} = 0.475 \epsilon_r + 0.67$$

where

$\epsilon_r \rightarrow$ relative dielectric constant of the material.

$\epsilon_{re} \rightarrow$ effective relative dielectric constant for a microstrip line

The diameter of the rectangular cross sectional area is given by $\downarrow$ of microstrip

$$d = 0.67 w \left(0.8 + \frac{t}{w} \right)$$

The phase velocity of a microstrip line is given by

$$v_p = v_c / \sqrt{\epsilon_{re}}$$

where $v_c \rightarrow$ velocity of electro-magnetic waves.

ϵ_{re} is given by an

$$\epsilon_{re} = \frac{\epsilon_r + 1}{2} + \frac{\epsilon_r - 1}{2} \left[1 + \frac{10h}{w} \right]^{-1/2}$$

The relationship for ϵ_{re} and d and Z_0 can be written as,

$$Z_0 = \frac{87}{\sqrt{\epsilon_r + 1.41}} \ln \left[\frac{5.98h}{0.8w + t} \right] \text{ for } h < 0.8w$$

if $w \gg h$ i.e. for a wide microstrip line, Z_0 is given by

$$Z_0 = \frac{377}{\sqrt{\epsilon_r}} \frac{h}{w}$$

Losses in Microstrip lines:

There are basic types of losses taking place in Microstrip.

- (i) Dielectric losses
- (ii) Ohmic losses
- (iii) Radiation losses.

Dielectric losses:

All dielectric materials have some conductivity but it is very small. When the conductivity is considered it cannot be neglected and the electric and magnetic fields in the dielectric will not be in time phase.

Cavity Resonators

→ A Cavity Resonator is a metallic enclosure that confines the Electromagnetic Energy. When one end of the waveguide is terminated in a shorting plate there will be reflections and hence standing waves. When another shorting plate is kept at a distance of a multiple of $\lambda_g/2$ then the hollow space so formed can support a signal which bounces back and forth between the two shorting plates. This results in resonance and hence the hollow space is called "cavity" and the resonator as the 'Cavity Resonator'.

→ The Waveguide section can be rectangular or Circular.

→ The Microwave cavity resonator is similar to a tuned circuit at low frequencies having a resonant frequency

$$f_0 = \frac{1}{2\pi\sqrt{LC}}$$

The Cavity Resonator can resonate at only one particular frequency like a parallel resonant circuit.

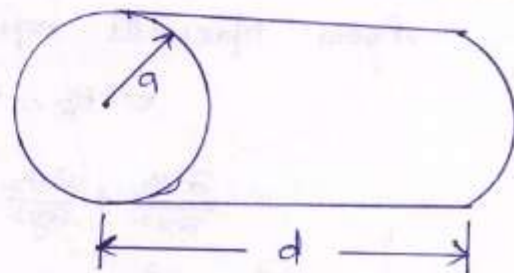
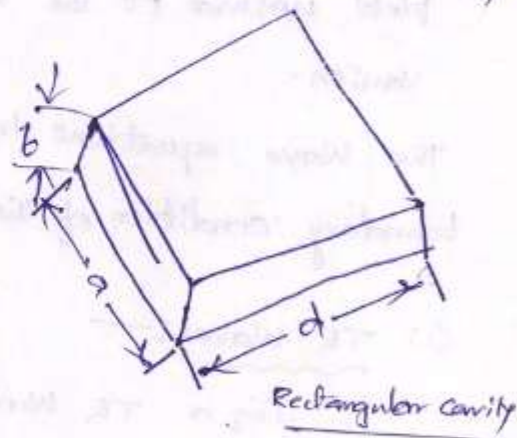
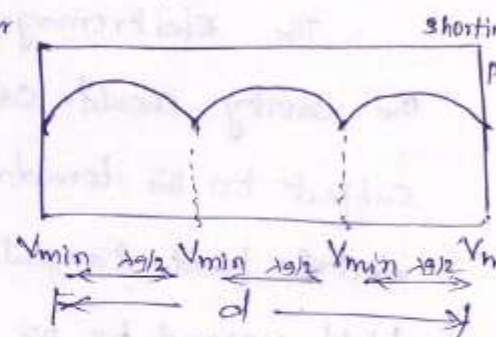
from the figure,

$$d = 3\lambda_g/2$$

→ The stored Electric and magnetic energies inside the cavity determine

its equivalent inductance and capacitance.

→ The energy dissipated by the finite conductivity of the cavity walls determines its equivalent resistance.



Q-factor of Microstrip is given as

$$Q = 0.63 h \sqrt{\epsilon_r f}$$

Where

h = distance between ground plane and microstrip

ϵ_r = conductivity of dielectric

f = operating frequency.

$$\frac{Q_{eff}}{Q} = \frac{1}{1 + \frac{1}{Q^2}}$$

$$Q_{eff} = Q \left(1 - \frac{1}{Q^2} \right)$$

Cavity Resonators

→ A Cavity Resonator is a metallic enclosure that confines the Electromagnetic Energy. When one end of the waveguide is terminated in a shorting plate there will be reflections and hence standing waves. When another shorting plate is kept at a distance of a multiple of $\lambda/2$ then the hollow space so formed can support a signal which bounces back and forth between the two shorting plates. This results in resonance and hence the hollow space is called "cavity" and the resonator as the 'Cavity Resonator'.

→ The Waveguide section can be rectangular or Circular.

→ The Microwave cavity resonator is similar to a tuned circuit at low frequencies having a resonant frequency

$$f_0 = \frac{1}{2\pi\sqrt{LC}}$$

The Cavity Resonator can resonate at only one particular frequency like a parallel resonant circuit.

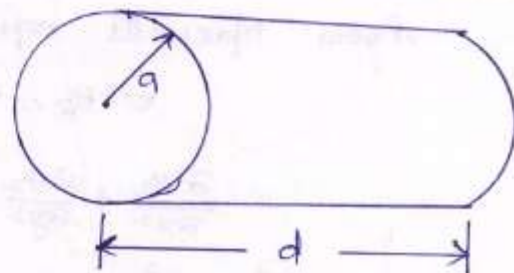
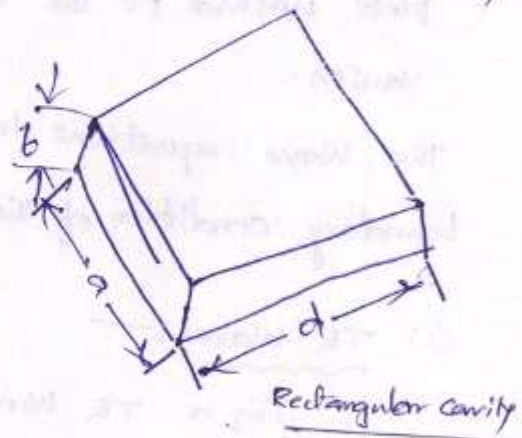
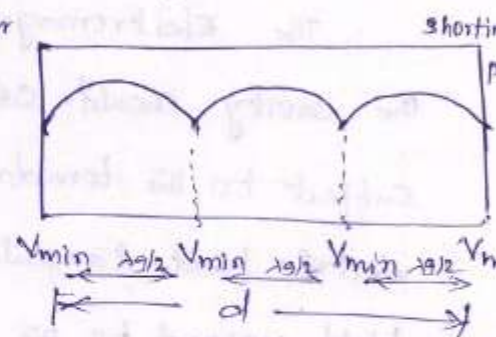
from the figure,

$$d = 3\lambda/2$$

→ The stored Electric and magnetic energies inside the cavity determine

its equivalent inductance and capacitance.

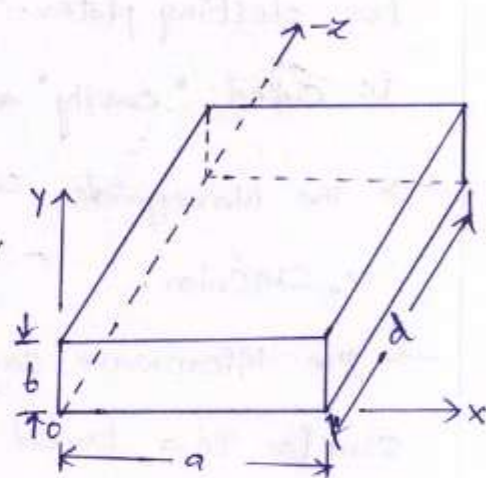
→ The energy dissipated by the finite conductivity of the cavity walls determines its equivalent resistance.



- A given resonator has an infinite number of resonant modes and each mode corresponds to a definite resonant frequency.
- When the frequency of an impressed signal is equal to a resonant frequency, a maximum amplitude of the standing wave occurs and the peak energies stored in the electric and magnetic fields are equal.
- The mode having the lowest resonant frequency is called as the 'Dominant Mode'.

Rectangular Cavity Resonator:

The Electromagnetic field inside the cavity should satisfy Maxwell's equations, subject to the boundary conditions that the electric field tangential to and the magnetic field normal to the metal walls must vanish.



The wave equations in the rectangular resonator should satisfy the boundary condition of the zero tangential E at four of the walls.

(1) TE Waves:

For a TE Wave, $E_z = 0$ & $H_z \neq 0$.

From Maxwell's equation,

$$\nabla^2 H_z = -\omega^2 \mu \epsilon H_z$$

$$\therefore \frac{\partial^2 H_z}{\partial x^2} + \frac{\partial^2 H_z}{\partial y^2} + \frac{\partial^2 H_z}{\partial z^2} = -\omega^2 \mu \epsilon H_z$$

$$\text{Since } \frac{\partial^2}{\partial z^2} = -\gamma^2, \therefore \frac{\partial^2 H_z}{\partial x^2} + \frac{\partial^2 H_z}{\partial y^2} + (\gamma^2 + \omega^2 \mu \epsilon) H_z = 0$$

$$\text{Let } \gamma^2 + \omega^2 \mu \epsilon = h^2$$

$$\therefore \frac{\partial^2 H_z}{\partial x^2} + \frac{\partial^2 H_z}{\partial y^2} + h^2 H_z = 0 \quad \longrightarrow \textcircled{1}$$

This is a partial differential equation of 2nd order.

Let $H_2 = xy \longrightarrow (2)$

where x is a function of x alone, y is a function of y alone.

$$\therefore y \frac{d^2 x}{dx^2} + x \frac{d^2 y}{dy^2} + h^2 xy = 0$$

$$\Rightarrow \frac{1}{x} \frac{d^2 x}{dx^2} + \frac{1}{y} \frac{d^2 y}{dy^2} + h^2 = 0 \longrightarrow (3)$$

where h^2 is a constant since x^2 & y^2 are constants.

so to satisfy the above equation, sum of functions of x and y must be equal to a constant. it is possible when individual one must be a constant.

$$\text{Let } \frac{1}{x} \frac{d^2 x}{dx^2} = -B^2, \quad \frac{1}{y} \frac{d^2 y}{dy^2} = -A^2 \longrightarrow (4)$$

where A^2 & B^2 are constants.

$$\therefore -A^2 - B^2 + h^2 = 0 \Rightarrow \boxed{h^2 = A^2 + B^2} \longrightarrow (5)$$

Solutions of equation (4) are

$$x = C_1 \cos Bx + C_2 \sin Bx \longrightarrow (6)$$

$$y = C_3 \cos Ay + C_4 \sin Ay \longrightarrow (7)$$

where C_1, C_2, C_3 and C_4 are constants,

$$\left. \begin{aligned} \frac{d}{dx} = D, \quad D^2 x &= -B^2 x \\ (D^2 + B^2)x &= 0 \\ D^2 + B^2 = 0 \Rightarrow D &= \pm jB \\ \therefore x &= C_1 \cos Bx + C_2 \sin Bx \\ y &= C_3 \cos Ay + C_4 \sin Ay \end{aligned} \right\}$$

which are determined by applying Boundary Conditions.

I. Boundary Condition (Bottom wall):

$$\therefore \text{from eqn (2), } H_2 = (C_1 \cos Bx + C_2 \sin Bx)(C_3 \cos Ay + C_4 \sin Ay) \longrightarrow (8)$$

$E_x = 0$ for $y=0$ and all values of x varying from 0 to a . We know,

$$E_x = -\frac{y}{h^2} \frac{\partial E_z}{\partial x} - \frac{j\omega\mu}{h^2} \frac{\partial H_z}{\partial y}$$

Since $E_z = 0$

$$\therefore E_x = -\frac{j\omega\mu}{h^2} \frac{\partial H_z}{\partial y}$$

$$\therefore E_x = -\frac{j\omega\mu}{h^2} \frac{\partial}{\partial y} [(C_1 \cos Bx + C_2 \sin Bx)(C_3 \cos Ay + C_4 \sin Ay)]$$

$$= -\frac{j\omega\mu}{h^2} [4 B C_2 \sin Bx + B C_1 \cos Bx] (-A C_3 \sin Ay + A C_4 \cos Ay)$$

from the BC (1),

$$0 = \frac{j\omega\mu}{h^2} [C_1 \cos Bx + C_2 \sin Bx] A C_4$$

$$\Rightarrow A C_4 [C_1 \cos Bx + C_2 \sin Bx] = 0$$

$$A \neq 0, C_1 \cos Bx + C_2 \sin Bx \neq 0$$

$$\therefore \boxed{C_4 = 0}$$

$$\therefore H_z = (C_1 \cos Bx + C_2 \sin Bx) (C_3 \cos Ay)$$

$$\therefore \boxed{H_z = C_3 [C_1 \cos Bx + C_2 \sin Bx] \cos Ay}$$

2nd Boundary condition [Left side wall]:

$E_y = 0$ for $x=0$ and y varying from 0 to b .

$$\text{Since } E_y = \frac{j\omega\mu}{h^2} \frac{\partial H_z}{\partial x}$$

$$\therefore E_y = \frac{j\omega\mu}{h^2} \frac{\partial}{\partial x} [(C_1 \cos Bx + C_2 \sin Bx) C_3 \cos Ay]$$

$$E_y = \frac{j\omega\mu}{h^2} [-BC_1 \sin Bx + BC_2 \cos Bx] C_3 \cos Ay$$

from the Boundary condition,

$$0 = \frac{j\omega\mu}{h^2} [BC_2 C_3 \cos Bx \cos Ay] \quad \because x=0$$

$$0 = \frac{j\omega\mu}{h^2} [BC_2 C_3 \cos Ay]$$

$$C_3 \cos Ay \neq 0, C_2 = 0.$$

$$\therefore \boxed{H_z = C_1 \cos Bx \cdot C_3 \cos Ay}$$

3rd Boundary condition (Top wall):

$E_x = 0$ for $y=b$ and x varies from 0 to a .

$$\therefore E_x = -\frac{j\omega\mu}{h^2} \frac{\partial H_z}{\partial y}$$

$$\therefore E_x = -\frac{j\omega\mu}{h^2} \frac{\partial}{\partial y} [C_1 C_3 \cos Bx \cos Ay] = -\frac{j\omega\mu}{h^2} \cdot C_1 C_3 \cdot A \cos Bx \sin Ay$$

$$0 = \frac{j\omega\mu}{h^2} \cdot C_1 C_3 \cdot A \cdot \cos Bx \cdot \sin Ab \Rightarrow C_1 C_3 \cos Bx \sin Ab = 0$$

MNF
U-11(3)

$$C_1 C_3 \neq 0, \cos Bx \neq 0, \sin Ab = 0$$

$$\Rightarrow Ab = n\pi \quad \text{where } n=0,1,2,3,\dots$$

$$\therefore \boxed{A = \frac{n\pi}{b}}$$

4¹⁵ Boundary Condition (Right Side wall):

$E_y = 0$ for $x=a$ and y varying from 0 to a .

$$E_y = \frac{j\omega\mu}{h^2} \frac{\partial}{\partial x} [C_1 C_3 \cos Bx \cos Ay] = -\frac{j\omega\mu}{h^2} \cdot C_1 C_3 \cdot B \cdot \sin Bx \cos Ay$$

$$E_y = \frac{j\omega\mu}{h^2} C_1 C_3 B \sin Bx \cos Ay$$

from the Boundary Condition, $E_y = 0$ for $x=a$ & $y=0 \rightarrow a$

$$0 = \frac{j\omega\mu}{h^2} \cdot B \cdot C_1 C_3 \cdot \sin Ba \cdot \cos Ay$$

$$C_1 C_3 \sin Ba \cos Ay = 0$$

$$C_1 C_3 \neq 0, \cos Ay \neq 0, \sin Ba = 0 \quad \text{where } m=0,1,2,\dots$$

$$Ba = m\pi$$

$$B = \frac{m\pi}{a}$$

$$\therefore H_z = C_1 C_3 \cos\left(\frac{m\pi}{a}x\right) \cos\left(\frac{n\pi}{b}y\right) \cdot e^{j\omega t - \gamma z}$$

$$\boxed{H_z = C \cos\left(\frac{m\pi}{a}x\right) \cos\left(\frac{n\pi}{b}y\right) e^{j(\omega t - \beta z)}}$$

$$\text{where } C = C_1 C_3$$

The Amplitude constant along the positive z direction is represented by A^+ , and that along the negative z direction by A^- . Adding the two travelling waves to obtain the fields of standing wave we have from above equation.

$$\therefore H_z = (A^+ e^{-j\beta z} + A^- e^{j\beta z}) \cos\left(\frac{m\pi}{a}x\right) \cos\left(\frac{n\pi}{b}y\right) e^{j\omega t}$$

To make E_y vanish at $x=0$ and $x=a$ we must make $A^+ = -A^-$ or $A^- = -A^+$ and also we know that

$$E_y = -\frac{\gamma}{h^2} \frac{\partial E_z}{\partial y} + \frac{j\omega\mu}{h^2} \frac{\partial H_z}{\partial x} \quad \text{Since } E_z = 0$$

$$E_y = \frac{j\omega\mu}{h^2} \frac{\partial H_z}{\partial z} = \frac{j\omega\mu}{h^2} \frac{\partial}{\partial z} \left[(A^+ e^{-j\beta z} + A^- e^{j\beta z}) \cos\left(\frac{m\pi}{a}x\right) \cos\left(\frac{n\pi}{b}y\right) e^{j\omega t} \right]$$

$$0 = \left[(A^+ e^{-j\beta z} + A^- e^{j\beta z}) \cdot -\left(\frac{m\pi}{a}\right) \sin\left(\frac{m\pi}{a}x\right) \cdot \cos\left(\frac{n\pi}{b}y\right) \right] e^{j\omega t}$$

But $\sin\left(\frac{m\pi}{a}x\right)$ and $\cos\left(\frac{n\pi}{b}y\right) \neq 0$

Therefore $A^+ e^{-j\beta z} + A^- e^{j\beta z} = 0$

To make $E_y = 0$, $A^+ = -A^-$.

$$0 = (A^+ e^{-j\beta z} - A^+ e^{j\beta z}) = 0 \Rightarrow A^+ \left[e^{-j\beta z} - e^{j\beta z} \right] = 0$$

$$-2j \sin \beta z \cdot A^+ = 0 \Rightarrow 2A^+ \sin \beta z = 0$$

$A^+ \neq 0$ only $\sin \beta d = 0$ with $z = d$

$$\therefore \beta d = p\pi \quad \text{where } p = 1, 2, 3, \dots$$

$$\therefore \boxed{\beta = \frac{p\pi}{d}}$$

$$\therefore H_z = -2jA^+ \cos\left(\frac{m\pi}{a}x\right) \cos\left(\frac{n\pi}{b}y\right) \sin\left(\frac{p\pi}{d}z\right) e^{j(\omega t - \beta z)}$$

By putting $-2jA^+ = C$

$$\boxed{H_z = C \cos\left(\frac{m\pi}{a}x\right) \cos\left(\frac{n\pi}{b}y\right) \sin\left(\frac{p\pi}{d}z\right) e^{j(\omega t - \beta z)}}$$

② For TM Waves:

For a TM Wave $H_z = 0$ & $E_z \neq 0$

The wave equation is given by

$$\frac{\partial^2 E_z}{\partial x^2} + \frac{\partial^2 E_z}{\partial y^2} + \frac{\partial^2 E_z}{\partial z^2} = -\omega^2 \mu \epsilon E_z$$

$$\Rightarrow \frac{\partial^2 E_z}{\partial x^2} + \frac{\partial^2 E_z}{\partial y^2} + (r^2 + \omega^2 \mu \epsilon) E_z = 0 \quad \text{Since } \frac{\partial^2}{\partial z^2} = -r^2$$

$$\Rightarrow \frac{\partial^2 E_z}{\partial x^2} + \frac{\partial^2 E_z}{\partial y^2} + h^2 E_z = 0 \quad \text{where } h^2 = r^2 + \omega^2 \mu \epsilon$$

This is a partial differential equation of 2nd order. The solution can be obtained by separation of variables method.

Let $E_z = XY = 0$

$$\therefore \frac{\partial^2 XY}{\partial x^2} + \frac{\partial^2 XY}{\partial y^2} + h^2 XY = 0$$

where X is a function of x only

Y is a function of y only

$$y \frac{d^2 x}{dx^2} + x \frac{d^2 y}{dy^2} + h^2 xy = 0$$

$$\Rightarrow \frac{1}{x} \frac{d^2 x}{dx^2} + \frac{1}{y} \frac{d^2 y}{dy^2} + h^2 = 0 \Rightarrow \frac{1}{x} \frac{d^2 x}{dx^2} + \frac{1}{y} \frac{d^2 y}{dy^2} = -h^2 = \text{constant}$$

To get sum of two functions constant individual must be a constant.

$$\text{let } \frac{1}{x} \frac{d^2 x}{dx^2} = -B^2 \text{ and } \frac{1}{y} \frac{d^2 y}{dy^2} = -A^2$$

$$\therefore -A^2 - B^2 = -h^2 \Rightarrow \boxed{A^2 + B^2 = h^2}$$

To find the solution for the following expression

$$\frac{1}{x} \frac{d^2 x}{dx^2} = -B^2$$

$$\frac{d^2 x}{dx^2} = -B^2 x$$

$$\Rightarrow D^2 x + B^2 x = 0 \text{ where } D = \frac{d}{dx} \text{ operator}$$

$$(D^2 + B^2)x = 0$$

$$D = \frac{d}{dx}$$

$$D^2 + B^2 = 0 \Rightarrow D^2 = -B^2$$

$$\boxed{D = \pm jB}$$

$$\frac{1}{y} \frac{d^2 y}{dy^2} = -A^2$$

$$\frac{d^2 y}{dy^2} = -A^2 y$$

$$(D^2 + A^2)y = 0$$

$$D^2 + A^2 = 0 \Rightarrow D^2 = -A^2$$

$$\boxed{D = \pm jA}$$

$$\therefore \boxed{x = C_1 \cos Bx + C_2 \sin Bx}$$

$$\boxed{y = C_3 \cos Ay + C_4 \sin Ay}$$

The constants C_1, C_2, C_3 and C_4 can be determined by the boundary conditions.

$$\text{Since } E_z = xy$$

So the required solution is

$$\boxed{E_z = (C_1 \cos Bx + C_2 \sin Bx)(C_3 \cos Ay + C_4 \sin Ay)}$$

Now we apply Boundary conditions to get constants C_1, C_2, C_3 and C_4 .

By Applying 1st Boundary Condition (Bottom wall)

$$E_z = 0 \text{ for } y=0 \text{ } \forall \text{ } x \text{ varying from } 0 \text{ to } a.$$

$$\therefore 0 = (C_1 \cos Bx + C_2 \sin Bx) (C_3 \cos 0 + C_4 \sin 0)$$

$$0 = C_3 [C_1 \cos Bx + C_2 \sin Bx]$$

for any value of 'x', $C_1 \cos Bx + C_2 \sin Bx \neq 0$.

So the possible condition is $C_3 = 0$

$$\therefore E_z = [C_1 \cos Bx + C_2 \sin Bx] C_4 \sin Ay$$

By Applying 2nd Boundary Condition (Left-side wall)

$$E_z = 0 \text{ for } x=0 \text{ } \forall \text{ } y \rightarrow 0 \text{ to } b.$$

$$\therefore 0 = [C_1 \cos 0 + C_2 \sin 0] [C_4 \sin Ay] \Rightarrow C_1 C_4 \sin Ay = 0$$

C_4 and $\sin Ay$ cannot be zero. if any of them or both are zero E_z component - Vanishes.

The possible condition is

$$C_1 = 0$$

$$\therefore E_z = C_2 \sin Bx C_4 \sin Ay$$

By Applying 3rd Boundary Condition (Top wall)

$$E_z = 0 \text{ for } y=b \text{ } \forall \text{ } x \rightarrow 0 \text{ to } a.$$

$$0 = C_2 \sin Bx \cdot C_4 \sin Ab \Rightarrow C_2 C_4 \sin Bx \sin Ab = 0$$

C_2, C_4 and $\sin Bx$ should not be zero to exist the field component E_z .

The possible condition is $\sin Ab = 0$

$$Ab = n\pi \text{ where } n=0, 1, 2, \dots$$

$$A = \left(\frac{n\pi}{b} \right)$$

MWE
Q-11(5)
By applying 4th Boundary Condition [Right Side Wall]

$$E_z = 0 \text{ at } x=a \text{ } \forall y \rightarrow 0 \text{ to } b$$

$$0 = C_2 C_4 \sin Ba \sin ay$$

C_2, C_4 and $\sin ay$ must not be zero to exist the E_z field component in TM mode,

The possible condition is

$$\sin Ba = 0$$

$$Ba = m\pi \quad \text{where } m=0,1,2,\dots$$

$$B = \frac{m\pi}{a}$$

$\therefore$ The expression for the field component exist in a Rectangular cavity under TM mode is

$$E_z = C_2 C_4 \sin\left(\frac{m\pi}{a}x\right) \sin\left(\frac{n\pi}{b}y\right) e^{j\omega t - \gamma z}$$

This field component is propagating along the direction z in cavity so while propagating amplitude and phase changes.

Therefore the field expression is

$$E_z = C_2 C_4 \sin\left(\frac{m\pi}{a}x\right) \sin\left(\frac{n\pi}{b}y\right) e^{j\omega t - \gamma z}$$

for the wave propagating along $+z$ -direction.

In a cavity, both forward wave and reflected waves are exist. Both will form a standing wave.

field expression for the reflected wave is

$$E_z = C_2 C_4 \sin\left(\frac{m\pi}{a}x\right) \sin\left(\frac{n\pi}{b}y\right) e^{j\omega t + \gamma z}$$

standing wave will be formed in a cavity, so the resultant field expression is

$$E_z = C_2 C_4 \sin\left(\frac{m\pi}{a}x\right) \sin\left(\frac{n\pi}{b}y\right) e^{j\omega t \pm \gamma z}$$

Let A^+ be the amplitude constant for the wave propagating in the positive z -direction and A^- be the amplitude constant for the wave propagating in the negative z -direction.

Then

$$E_z = [A^+ e^{-j\beta z} + A^- e^{j\beta z}] \left[C_2 C_4 \sin\left(\frac{m\pi}{a}x\right) \sin\left(\frac{n\pi}{b}y\right) e^{j(\omega t)} \right]$$

To make E_z vanish at $z=0$, we must choose $A^+ = -A^- = A$ so that

$$E_z = [A^+ e^{-j\beta z} + A^+ e^{j\beta z}] \left[C_2 C_4 \sin\left(\frac{m\pi}{a}x\right) \sin\left(\frac{n\pi}{b}y\right) e^{j\omega t} \right]$$

$$E_z = +A^+ \left[e^{j\beta z} + e^{-j\beta z} \right] \left[C_2 C_4 \sin\left(\frac{m\pi}{a}x\right) \sin\left(\frac{n\pi}{b}y\right) e^{j\omega t} \right]$$

$$\begin{aligned} \text{Since } e^{j\theta} - e^{-j\theta} &= 2j \sin \theta & e^{j\theta} + e^{-j\theta} &= 2 \cos \theta \\ \therefore e^{j\beta z} - e^{-j\beta z} &= 2j \sin \beta z & e^{j\beta z} + e^{-j\beta z} &= 2 \cos \beta z \end{aligned}$$

$$\therefore E_z = +A^+ \cos \beta z$$

$$\therefore E_z = A \cdot C_2 C_4 \sin\left(\frac{m\pi}{a}x\right) \sin\left(\frac{n\pi}{b}y\right) \cos \beta z e^{j\omega t}$$

from the wave equation $\nabla^2 E = -j\omega \mu H$

$$\begin{vmatrix} a_x & a_y & a_z \\ \frac{\partial}{\partial x} & \frac{\partial}{\partial y} & \frac{\partial}{\partial z} \\ E_x & E_y & E_z \end{vmatrix} = -j\omega \mu \begin{bmatrix} H_x & H_y & H_z \end{bmatrix}$$

$$a_x \rightarrow \frac{\partial E_z}{\partial y} - \frac{\partial E_y}{\partial z} = -j\omega \mu H_x \rightarrow \textcircled{a}$$

$$a_y \rightarrow \frac{\partial E_z}{\partial x} - \frac{\partial E_x}{\partial z} = +j\omega \mu H_y \rightarrow \textcircled{b}$$

$$a_z \rightarrow \frac{\partial E_x}{\partial x} - \frac{\partial E_x}{\partial y} = -j\omega \mu H_z \rightarrow \textcircled{c}$$

$$\nabla \times H = +j\omega \epsilon E$$

$$a_x \rightarrow \frac{\partial H_z}{\partial y} - \frac{\partial H_y}{\partial z} = +j\omega \epsilon E_x \rightarrow \textcircled{d}$$

$$\begin{vmatrix} a_x & a_y & a_z \\ \frac{\partial}{\partial x} & \frac{\partial}{\partial y} & \frac{\partial}{\partial z} \\ H_x & H_y & H_z \end{vmatrix} = +j\omega \epsilon \begin{bmatrix} E_x & E_y & E_z \end{bmatrix}$$

$$a_y \rightarrow \frac{\partial H_z}{\partial x} - \frac{\partial H_x}{\partial z} = -j\omega \epsilon E_y \rightarrow \textcircled{e}$$

$$a_z \rightarrow \frac{\partial H_y}{\partial x} - \frac{\partial H_x}{\partial y} = +j\omega \epsilon E_z \rightarrow \textcircled{f}$$

from equation (6)

$$H_y = \frac{1}{j\omega\mu} \left[\frac{\partial E_z}{\partial x} - \frac{\partial E_x}{\partial z} \right]$$

Substitute H_y in eqn (1)

$$j\omega\epsilon E_x = \frac{\partial H_z}{\partial y} - \frac{\partial}{\partial z} \left[\left(\frac{\partial E_z}{\partial x} - \frac{\partial E_x}{\partial z} \right) \frac{1}{j\omega\mu} \right]$$

$$j\omega\epsilon E_x = \frac{\partial H_z}{\partial y} - \frac{1}{j\omega\mu} \frac{\partial^2 E_z}{\partial z^2} + \frac{1}{j\omega\mu} \frac{\partial^2 E_x}{\partial z^2}$$

in TM mode, $H_z = 0$

$$j\omega\epsilon E_x = -\frac{1}{j\omega\mu} \frac{\partial^2 E_x}{\partial z^2} - \frac{1}{j\omega\mu} \frac{\partial^2 E_z}{\partial z^2 \partial x}$$

$$j\omega\epsilon E_x = \frac{1}{j\omega\mu} \left[\frac{\partial^2 E_x}{\partial z^2} - \frac{\partial^2 E_z}{\partial z^2 \partial x} \right]$$

→ (3)

from the Boundary Condition,

$$E_y = 0, E_x = 0 \text{ at } z=0, d.$$

Since $E_z = A C_2 C_4 \sin\left(\frac{m\pi}{a}\right)x \sin\left(\frac{n\pi}{b}\right)y \cos\beta z e^{j\omega t}$

from eqn (3) & (4), Boundary condition is satisfied when

$$\frac{\partial E_z}{\partial z} = 0$$

$$\therefore \frac{\partial}{\partial z} [A C_2 C_4 \sin\left(\frac{m\pi}{a}\right)x \sin\left(\frac{n\pi}{b}\right)y \cos\beta z e^{j\omega t}] = 0$$

$$\frac{d}{dz} [\cos\beta d] = 0 \Rightarrow -\sin\beta d = 0$$

$$\sin\beta d = 0$$

$$\Rightarrow \sin\beta d = \sin p\pi \text{ where } p=0,1,2,\dots$$

$$\beta d = p\pi$$

$$\boxed{\beta = \frac{p\pi}{d}}$$

$$\therefore \boxed{E_z = A C_2 C_4 \sin\left(\frac{m\pi}{a}\right)x \sin\left(\frac{n\pi}{b}\right)y \cos\left(\frac{p\pi}{d}\right)z e^{j\omega t}}$$

from equation (2)

$$H_x = -\frac{1}{j\omega\mu} \left[\frac{\partial E_z}{\partial y} - \frac{\partial E_y}{\partial z} \right]$$

Substitute H_x in eqn (2),

$$-j\omega\epsilon E_y = \frac{\partial H_z}{\partial x} + \frac{1}{j\omega\mu} \frac{\partial}{\partial z} \left[\left(\frac{\partial E_z}{\partial y} - \frac{\partial E_y}{\partial z} \right) \right]$$

$$-j\omega\epsilon E_y = \frac{\partial H_z}{\partial x} + \frac{1}{j\omega\mu} \left[\frac{\partial^2 E_z}{\partial z^2 \partial y} - \frac{\partial^2 E_y}{\partial z^2} \right]$$

in TM mode, $H_z = 0$

$$-j\omega\epsilon E_y = \frac{1}{j\omega\mu} \left[\frac{\partial^2 E_z}{\partial z^2 \partial y} - \frac{\partial^2 E_y}{\partial z^2} \right]$$

$$E_y = +\frac{1}{\omega^2\mu\epsilon} \left[\frac{\partial^2 E_z}{\partial z^2 \partial y} - \frac{\partial^2 E_y}{\partial z^2} \right]$$

→ (4)

In TE we have, $E_z = 0$.

$$H_z = c \cos\left(\frac{m\pi}{a}x\right) \cos\left(\frac{n\pi}{b}y\right) \sin\left(\frac{p\pi}{d}z\right) e^{j\omega t}$$

field components can be determined by substituting E_z & H_z in following expressions.

$$E_x = -\frac{\gamma}{h^2} \frac{\partial E_z}{\partial x} - \frac{j\omega\mu}{h^2} \frac{\partial H_z}{\partial y}$$

$$E_x = -\frac{j\omega\mu}{h^2} \frac{\partial}{\partial y} \left[c \cos\left(\frac{m\pi}{a}x\right) \cos\left(\frac{n\pi}{b}y\right) \sin\left(\frac{p\pi}{d}z\right) e^{j\omega t} \right]$$

$$= \underbrace{\frac{j\omega\mu \cdot c \cdot \left(\frac{n\pi}{b}\right)}{h^2}}_{\text{constant}} \cos\left(\frac{m\pi}{a}x\right) \sin\left(\frac{n\pi}{b}y\right) \sin\left(\frac{p\pi}{d}z\right) e^{j\omega t}$$

$$\underline{E_x = C_1 \cos\left(\frac{m\pi}{a}x\right) \sin\left(\frac{n\pi}{b}y\right) \sin\left(\frac{p\pi}{d}z\right) e^{j\omega t}}$$

$$E_y = -\frac{\gamma}{h^2} \frac{\partial E_z}{\partial y} + \frac{j\omega\mu}{h^2} \frac{\partial H_z}{\partial x}$$

$$\therefore \underline{E_y = C_2 \sin\left(\frac{m\pi}{a}x\right) \cos\left(\frac{n\pi}{b}y\right) \sin\left(\frac{p\pi}{d}z\right) e^{j\omega t}}$$

$$H_x = -\frac{\gamma}{h^2} \frac{\partial H_z}{\partial x} + \frac{j\omega\mu}{h^2} \frac{\partial E_z}{\partial y} = -\frac{\gamma}{h^2} \frac{\partial H_z}{\partial x}$$

$$\underline{H_x = C_3 \sin\left(\frac{m\pi}{a}x\right) \cos\left(\frac{n\pi}{b}y\right) \sin\left(\frac{p\pi}{d}z\right) e^{j\omega t}}$$

$$H_y = -\frac{\gamma}{h^2} \frac{\partial H_z}{\partial y} - \frac{j\omega\mu}{h^2} \frac{\partial E_z}{\partial x} = -\frac{\gamma}{h^2} \frac{\partial H_z}{\partial y}$$

$$\therefore \underline{H_y = C_4 \cos\left(\frac{m\pi}{a}x\right) \sin\left(\frac{n\pi}{b}y\right) \sin\left(\frac{p\pi}{d}z\right) e^{j\omega t}}$$

In TM wave, we have $H_z = 0$.

$$E_z = c \sin\left(\frac{m\pi}{a}x\right) \sin\left(\frac{n\pi}{b}y\right) \cos\left(\frac{p\pi}{d}z\right) e^{j\omega t}$$

field components are

$$E_x = -\frac{\gamma}{h^2} \frac{\partial E_z}{\partial x} - \frac{j\omega\mu}{h^2} \frac{\partial H_z}{\partial y} = -\frac{\gamma}{h^2} \frac{\partial E_z}{\partial x} = -\frac{\gamma}{h^2} \cdot c \cdot \left(\frac{m\pi}{a}\right) \cos\left(\frac{m\pi}{a}x\right) \sin\left(\frac{n\pi}{b}y\right) \cos\left(\frac{p\pi}{d}z\right) e^{j\omega t}$$

$$\underline{E_x = C_1 \cos\left(\frac{m\pi}{a}x\right) \sin\left(\frac{n\pi}{b}y\right) \cos\left(\frac{p\pi}{d}z\right) e^{j\omega t}}$$

MWE
U=11(7)

$$E_y = -\frac{r}{h^2} \frac{\partial E_2}{\partial y} + \frac{j\omega\mu}{h^2} \frac{\partial H_2}{\partial x} = -\frac{r}{h^2} \frac{\partial E_2}{\partial y}$$

$$E_y = c_2 \sin\left(\frac{m\pi}{a}\right)x \cos\left(\frac{n\pi}{b}\right)y \cos\left(\frac{p\pi}{d}\right)z e^{j\omega t}$$

$$H_x = -\frac{r}{h^2} \frac{\partial H_2}{\partial x} + \frac{j\omega\mu}{h^2} \frac{\partial E_2}{\partial y} = \frac{j\omega\mu}{h^2} \frac{\partial E_2}{\partial y}$$

$$\therefore H_x = c_3 \sin\left(\frac{m\pi}{a}\right)x \cos\left(\frac{n\pi}{b}\right)y \cos\left(\frac{p\pi}{d}\right)z e^{j\omega t}$$

$$H_y = -\frac{r}{h^2} \frac{\partial H_2}{\partial y} - \frac{j\omega\mu}{h^2} \frac{\partial E_2}{\partial x} = -\frac{j\omega\mu}{h^2} \frac{\partial E_2}{\partial x}$$

$$\therefore H_y = c_4 \cos\left(\frac{m\pi}{a}\right)x \sin\left(\frac{n\pi}{b}\right)y \cos\left(\frac{p\pi}{d}\right)z e^{j\omega t}$$

Resonant frequency (f_0):

Since we know that for a rectangular waveguide

$$h^2 = r^2 + w^2\mu\epsilon = A^2 + B^2 = \left(\frac{m\pi}{a}\right)^2 + \left(\frac{n\pi}{b}\right)^2$$

$$\therefore w^2\mu\epsilon = \left[\frac{m\pi}{a}\right]^2 + \left[\frac{n\pi}{b}\right]^2 - r^2$$

for a wave propagation, $r = j\beta$

$$\text{or } r^2 = j^2\beta^2 = -\beta^2$$

$$\therefore w^2\mu\epsilon = \left[\frac{m\pi}{a}\right]^2 + \left[\frac{n\pi}{b}\right]^2 + \beta^2$$

If a wave has to exist in a cavity resonator there must be a phase change corresponding to a given guide wavelength ($\beta = 2\pi/\lambda_g$). The condition for the resonator to resonate is $\beta = p\pi/d$.

where $p = \text{a constant} = 1, 2, 3, \dots, \infty$, that indicates half-wave variation of either electric or magnetic field along that z -direction. $d = \text{length of the resonator}$.

When

$$\beta = \frac{p\pi}{d}, \quad f = f_0, \quad \omega = 2\pi f_0 = \omega_0$$

$$\therefore \omega_0^2\mu\epsilon = \left[\frac{m\pi}{a}\right]^2 + \left[\frac{n\pi}{b}\right]^2 + \left[\frac{p\pi}{d}\right]^2 = \omega_0^2\mu\epsilon = \left[\frac{m\pi}{a}\right]^2 + \left[\frac{n\pi}{b}\right]^2 + \left[\frac{p\pi}{d}\right]^2$$

$$f_0 = \frac{1}{2\pi\sqrt{\mu\epsilon}} \left[\left(\frac{m\pi}{a} \right)^2 + \left(\frac{n\pi}{b} \right)^2 + \left(\frac{p\pi}{d} \right)^2 \right]^{1/2}$$

Since $c = \frac{1}{\sqrt{\mu\epsilon}}$

$$f_0 = \frac{c}{2\pi} \left[\left(\frac{m\pi}{a} \right)^2 + \left(\frac{n\pi}{b} \right)^2 + \left(\frac{p\pi}{d} \right)^2 \right]^{1/2}$$

$$f_0 = \frac{c}{2} \left[\left(\frac{m}{a} \right)^2 + \left(\frac{n}{b} \right)^2 + \left(\frac{p}{d} \right)^2 \right]^{1/2}$$

→ General mode of propagation in a cavity resonator is TE_{mnp} or TM_{mnp} .

→ For both TE and TM the resonant frequency is the same in a rectangular cavity resonator.

Dominant Mode: —

→ The mode for which cut-off wavelength is large is called 'Dominant Mode'.

→ For $a > b > d$, the dominant mode is the TE_{101} mode.

Q-factor: —

The quality factor Q is a measure of the frequency selectivity of a resonator or antiresonant circuit, and it is defined as

$$Q = 2\pi \frac{\text{Maximum Energy stored per cycle}}{\text{Energy dissipated per cycle}} = \frac{\omega W}{P}$$

$$\begin{aligned} \text{Since } Q &= 2\pi \frac{W}{P \cdot T} \\ Q &= 2\pi f \frac{W}{P} = \omega \frac{W}{P} \end{aligned}$$

Where $W \rightarrow$ Maximum stored Energy and P is the average power loss.
 $\omega =$ Resonant frequency.

→ The Q of a perfect or ideal cavity resonator is infinite since in a perfect conductor forming the cavity, P would be zero and also once energised it would resonate for ever.

→ If there is more than one resonant frequency, there will be different values of Q for the various values of frequencies.

→ Coupling loops are used to couple the energy in and out of a cavity resonator. This coupling changes the value of Q . This Q that takes into account the coupling between the cavity and coupling part is known as the loaded Q_L . So, Q_L can be given by

$$\boxed{\frac{1}{Q_L} = \frac{1}{Q_0} + \frac{1}{Q_{ext}}}$$

where

$Q_0 = Q$ of an unloaded cavity

$Q_{ext} = Q$ due to external ohmic losses

There are three types of couplings. They are critically coupling, under coupling and over coupling.

An unloaded resonator can be represented by a series or parallel resonant circuit. The unloaded Q is then given by

$$Q_0 = \frac{\omega_0 L}{R}$$

$$Q_L = \frac{Q_0 \cdot Q_{ext}}{Q_0 + Q_{ext}}$$

then Q_{ext} can be written as

$$Q_{ext} = \frac{Q_0}{k} = \frac{\omega_0 L}{kR}$$

$k \rightarrow$ coupling coefficient of the cavity.

$$\therefore Q_L = \frac{Q_0 \cdot Q_0/k}{Q_0 + Q_0/k} = \frac{Q_0^2/k}{Q_0(1+k)/k} = \frac{Q_0}{1+k}$$

→ for Critical Coupled cavity (resonator matched to the generator) $k=1$ and hence loaded Q [Q_L] is given by

$$\boxed{Q_L = \frac{1}{2} Q_{ext} = \frac{1}{2} Q_0}$$

→ for under coupled cavity, $k < 1$, the cavity terminals are at a voltage minimum and the input impedance is the reciprocal of standing wave ratio

$$\text{i.e. } k = \frac{1}{p}$$

Therefore

$$\boxed{Q_L = \frac{p}{1+p} Q_0}$$

$$\frac{1}{Q_L} = \frac{Q_{ext} + Q_0}{Q_{ext} Q_0} \Rightarrow Q_L = \frac{Q_{ext} \cdot Q_0}{Q_{ext} + Q_0}$$

$$Q_L = \frac{1 \cdot Q_0^2}{p Q_0 + Q_0} = \frac{p Q_0}{1+p} \because Q_{ext} = p Q_0$$

→ for over coupled cavity, $k > 1$. cavity terminals are at a voltage maximum and the impedance is standing wave ratio.
i.e. $k = p$

therefore $Q_L = \frac{Q_0}{1+f}$

Unloaded Quality factor, Q_0 of a cavity can also be defined by taking into consideration the volume of cavity that stores the energy and the value of the metal that determines the energy dissipated. Volume of the metal in which currents are flowing is equal to the surface area of the cavity multiplied by the skin depth (δ).

Hence, $Q_0 = \frac{\text{Volume of cavity}}{2 \times \text{Surface area of cavity}}$

$$Q_0 = \frac{\text{Cross section area of cavity}}{2 \times \text{periphery of cavity}}$$

Thus Q factor of cavity can be increased by increasing the size of cavity or conductivity of the walls or by decreasing the coupling into the cavity.

→ Q also increases with an increase in frequency as skin depth decreases with frequency.

Q of a rectangular cavity Resonator:-

The stored energy in a resonator is obtained by integrating over the volume of the resonator, the power loss ^{obtained} by integrating the power density over the surface of the resonator.

$$U_{\text{stored}} = \frac{1}{2} \int_{\text{vol}} E^2 dv = \frac{1}{2} \int_{\text{vol}} \mu H^2 dv \quad \longrightarrow \quad ①$$

$$U_{\text{lost}} = \frac{R_s}{2} \int_{\text{sur}} H_t^2 dA \quad \longrightarrow \quad ②$$

Where E, H, H_t = peak values.

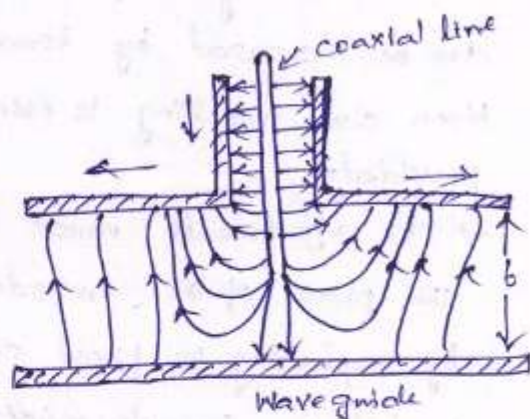
R_s = Surface resistance of the resonator wall

Coupling Mechanisms:

In order to excite a particular mode, the cavity must be properly coupled by an external source. Different coupling methods are used in Microwave fitness and Wave meters are

① Coupling probes:

A coaxial line may be coupled to a Waveguide by means of either a probe parallel to the electric field at (or near) a point where the electric field has its maximum value or a loop at a point where the magnetic field has its max. value.



→ A coupling probe usually consists of an extension of the center conductor of the coaxial line at the mid point of one of the wide walls of the guide, that is, of one of the walls normal to the E -field as shown. Figure shows the progress of two lines of electric flux at a point in the wave as the wave passes through the junction from the coaxial line to the waveguide. Because the electric field in the surroundings of the probe has components normal to the axis of the probe, and because the electric and magnetic fields in the surroundings of the probe differ in other respects from the desired TE_{10} field, higher order modes of propagation are also excited.

→ If the dimensions of the guide are properly chosen, all higher order modes are greatly attenuated within a wavelength.

→ Usually the waveguide is terminated in a short circuit and the probe is placed approximately a quarter wavelength from the termination. Then, when the probe is used to couple out power, it is at an antinode of electric field and when it is used to excite the waveguide, direct waves from the probe are reinforced by waves reflected from the

$$\frac{1}{2} \int_{-\infty}^{\infty} \frac{1}{x^2} dx = \frac{1}{2} \left[-\frac{1}{x} \right]_{-\infty}^{\infty} = \frac{1}{2} \left(0 - 0 \right) = 0$$

$$x^2 + y^2 = r^2$$

If r is sufficiently large, the area of the circle is approximately πr^2 .

$$\frac{1}{2} \int_{-\infty}^{\infty} \frac{1}{x^2} dx = \frac{1}{2} \left[-\frac{1}{x} \right]_{-\infty}^{\infty} = 0$$

$$\left\{ \begin{aligned} & \frac{1}{2} \int_{-\infty}^{\infty} \frac{1}{x^2} dx = 0 \\ & \frac{1}{2} \int_{-\infty}^{\infty} \frac{1}{x^2} dx = 0 \end{aligned} \right.$$

$$\frac{1}{2} \int_{-\infty}^{\infty} \frac{1}{x^2} dx = 0$$

The area of the circle is approximately πr^2 if r is sufficiently large.

$$\frac{1}{2} \int_{-\infty}^{\infty} \frac{1}{x^2} dx = 0$$

$$\frac{1}{2} \int_{-\infty}^{\infty} \frac{1}{x^2} dx = 0$$

The area of the circle is approximately πr^2 if r is sufficiently large.

$$\frac{1}{2} \int_{-\infty}^{\infty} \frac{1}{x^2} dx = 0$$

The area of the circle is approximately πr^2 if r is sufficiently large.

$$\frac{1}{2} \int_{-\infty}^{\infty} \frac{1}{x^2} dx = 0$$

$$\frac{1}{2} \int_{-\infty}^{\infty} \frac{1}{x^2} dx = 0$$

$$\frac{1}{2} \int_{-\infty}^{\infty} \frac{1}{x^2} dx = 0$$

$$\frac{1}{2} \int_{-\infty}^{\infty} \frac{1}{x^2} dx = 0$$

The area of the circle is approximately πr^2 if r is sufficiently large.

11/11/20
V-3 ① To excite a particular mode, the probe (or probes) is placed parallel to the electric field in a position where the field exhibited by that mode has its greatest intensity. A multiplicity of coupling probes is desirable for the excitation of higher order modes.

Where several probes are employed for the excitation, they must not only be placed in the proper positions within the waveguide, but they must also be excited with a relative phasing corresponding to that of the fields of the desired mode. Variable length transmission line sections are useful for the adjustment of this phasing.

⑥ Coupling Loops:

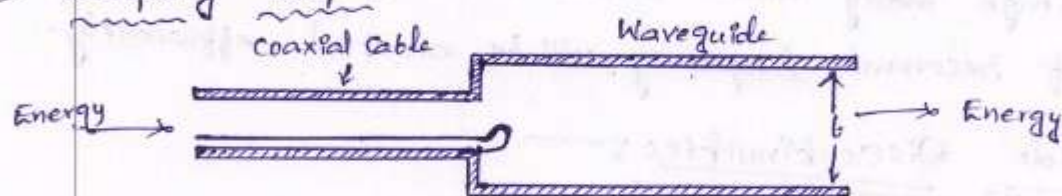


Figure shows the use of a coupling loop in a junction between coaxial line and waveguide. Since loop coupling is principally magnetic, the loop must be placed at or near a point of high magnetic field strength and turned in such a position that its plane is normal to the flux lines. It may be mounted in the end wall of a short circuited waveguide, or in the middle of the top or bottom wall an integral number of half wavelengths from a short circuited end. The plane of the loop should be perpendicular to the H-field for maximum coupling.

The degree of coupling obtained with a loop depends upon its size and shape and in general increases with the area of the loop. The increased self inductance resulting from increased loop size may be objectionable. If the length of the loop conductor is not small in comparison with the wavelength, there is appreciable potential gradient along the conductor, and magnetic coupling may be accompanied by appreciable electric coupling.

closed end of the guide.

→ In order to minimize reflections at the junction, the probe must be matched to the waveguide. Matching is accomplished by proper choice of probe length and position of the probe relative to the closed end of the guide.

→ The frequency range over which satisfactory matching is obtained can be increased by rounding and flaring the end of the probe. When close matching is essential, adjustable matching elements must be provided.

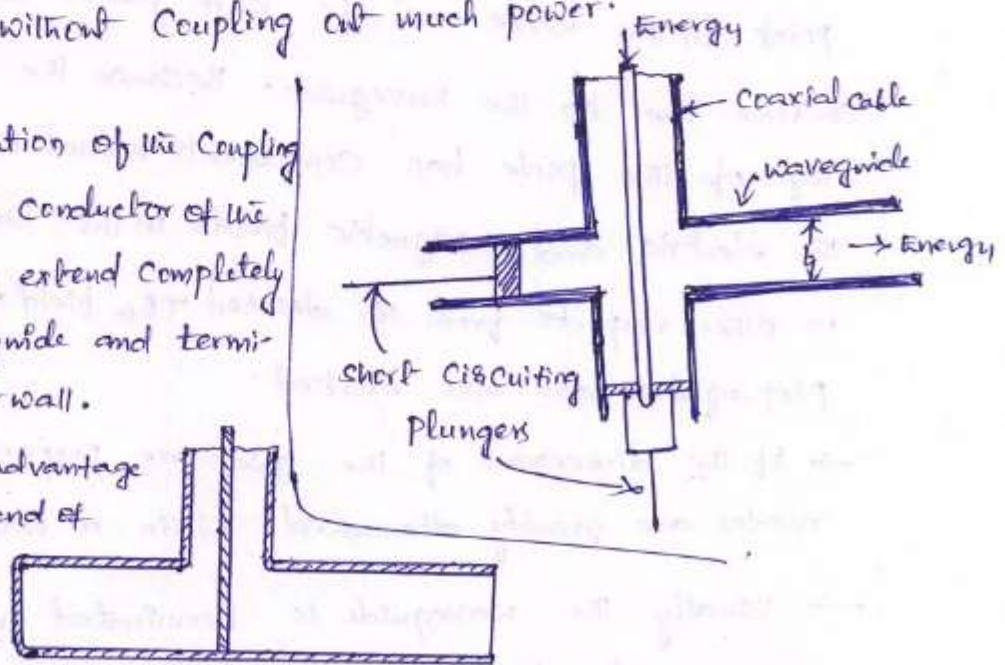
→ Two adjustments must be provided to match both the magnitude and the phase of the impedances. One convenient adjustment is obtained by replacing the fixed closed end of the waveguide by an adjustable short-circuiting plunger. Matching elements must be provided short-circuiting plungers. A second adjustment is readily provided by one or more stubs in the coaxial line, close to the junction as shown below.

→ A less convenient adjustment is variation of probe length. Complete matching may be accomplished by the use of two or more stubs in the coaxial line or waveguide.

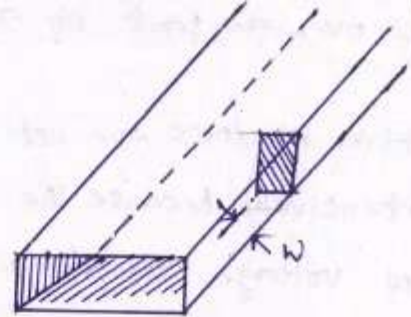
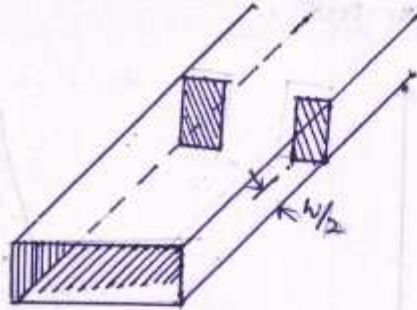
→ Short probes are used to measuring the field strength within the waveguide without coupling out much power.

→ In a modification of the coupling probe, the center conductor of the coaxial line may extend completely across the waveguide and terminate on the far wall.

This design has the advantage that the (waveguide) end of the center conductor is rigidly supported.



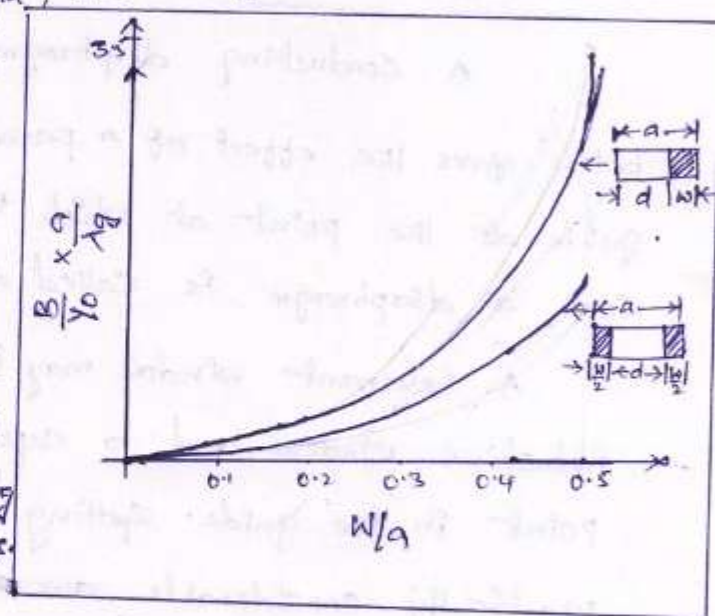
→ When the center conductor extends across the waveguide or projects an appreciable distance into the guide, the magnetic coupling, as well as the electric coupling, is effective.



Inductive Windows (a) Symmetrical type (b) Asymmetrical type.

Typical curves of $\frac{B}{Y_0} \times \frac{a}{\lambda_g}$ vs $\frac{W}{a}$ for symmetrical and asymmetrical inductive windows shown below. The normalized susceptance is given by B/Y_0 , λ_g is Waveguide Wavelength, $a \rightarrow$ Waveguide width, $d = a - w$ is the aperture width and Y is the reciprocal of the waveguide characteristic impedance.

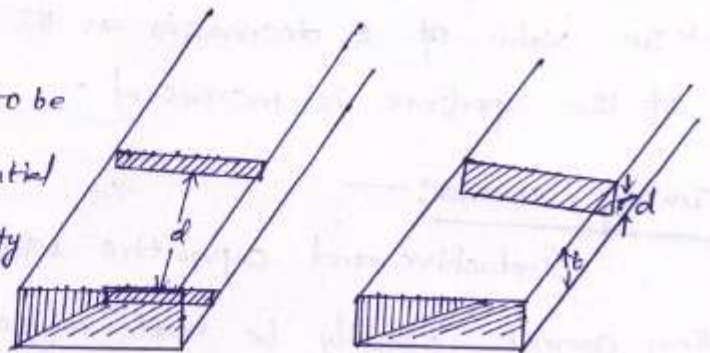
The choice between a symmetrical or an asymmetrical type of inductive window may be governed by mechanical considerations such as ease of machining and installation of pressurizing windows.



(c) Capacitive Windows: —

conducting diaphragms extending into a rectangular waveguide from the top and bottom walls produce the effect of capacitive susceptance shunted across the waveguide at that point. Therefore they are called "capacitive windows".

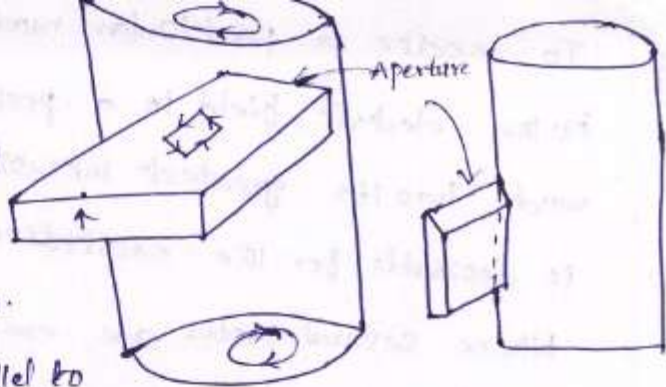
If diaphragms are assumed to be a perfectly conducting equipotential surface, the electric field intensity must be more intense in the aperture than in a plane immediately



in front of the window. This is true since the E -field in the plane of the window must be zero over part of the path between top and bottom walls of the waveguide while the E -field in the adjacent plane need not

③ Aperture Coupling:—

This coupling is between cylindrical cavity and rectangular waveguide operating in dominant TE_{10} mode.



→ Magnetic field component parallel to the long dimension of the slot will be coupled through the aperture.

TE_{10} magnetic field couples into cavity, which allows any of TM resonant modes to be excited.

Even though many modes can be excited in the waveguide, only that mode of resonant frequency will be excited efficiently.

Waveguide Discontinuities:—

Waveguide Windows:—

Impedance matching in waveguides can be accomplished by means of waveguide stubs and adjustable length sections similar to those used in coaxial lines but such components are in general heavy and difficult to use them. The same results can be achieved in a simple manner by the use of fixed or adjustable projections from the walls of the waveguide i.e. by waveguide windows. Different types of waveguide windows are

- ① Inductive Window
- ② Capacitive Window
- ③ Resonant windows

① Inductive Windows:—

conducting diaphragms extending into a waveguide from the side walls (as shown in figure) have the effect. If the thickness of the diaphragm is small in comparison to the wavelength then the inductive susceptance across the waveguide increases at the point at which the diaphragms are placed. Such an element is therefore called an "Inductive Window".

MWk
u II (12)
A better way to obtain variable susceptance is by means of a screw inserted into the top or bottom walls of the waveguide.

→ A screw of length less than approximately a quarter wave length produces an effective (capacitance) capacitive susceptance, the value of which increases with depth of penetration. When the depth of penetration is approximately a quarter wave length the screw is series resonant and further insertion causes the susceptance to become inductive.

→ The sharpness of resonance is an inverse function of the screw diameter.

→ The most direct method of impedance matching with a matching screw is to use a single screw, adjustable as to bolt length and position along the waveguide.

The disadvantage of the use of a single screw in impedance matching is the requirement of a slot in the wall of the waveguide.

An alternative is the use of double or triple screw units, in which two or three adjustable depth screws are spaced an eighth or a quarter of a guide wavelength apart.

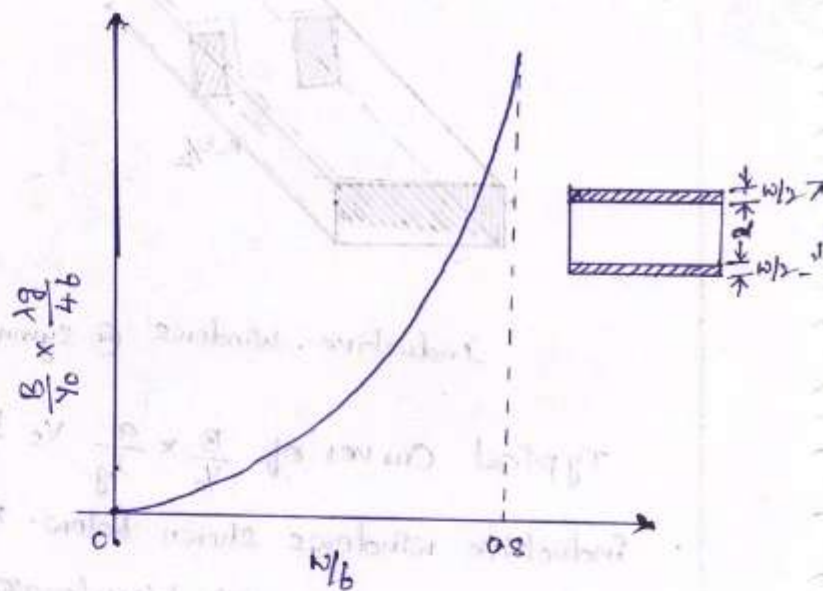
Although the triple screw unit has the advantage that it can be used to correct for any mismatch, it is very difficult to adjust. Consequently the double screw unit is usually used. All possible mismatches can be corrected with a double screw unit designed so that turning it end for end in the waveguide system displaces the screws by a quarter of a guide wavelength.

Posts:

When a metallic cylindrical post is introduced into the broader side of the waveguide, it produces the same effect as a window in providing lumped reactance at that point.

be zero over any part of a similar path.

→ capacitive windows are not used extensively because the danger of voltage breakdown across the window reduces the power that can be transmitted by the waveguide.

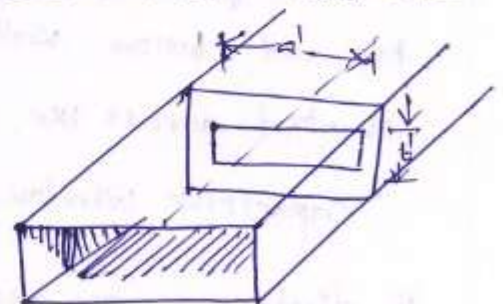


(c) Resonant Windows:—

A conducting diaphragm of the form shown in figure shown below gives the effect of a parallel L-C circuit shunted across the guide at the point at which the diaphragm placed. for this reason such a diaphragm is called a "Resonant window".

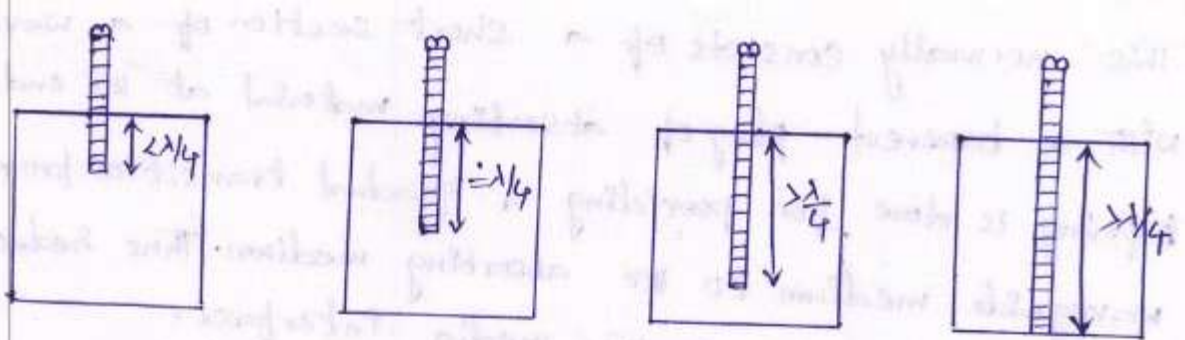
A resonant window may be considered as a combination of an inductive window and a capacitive window inserted at the same point in the guide. Getting accurate results is difficult but a practically considerable amount of information has been collected. practically zero susceptance is obtained at a chosen frequency with a wide variety of diaphragm shapes if the dimensions a' & b' are properly related. However, a minimum size of aperture that can be used.

→ the value of Q decreases as the size of the aperture is increased.



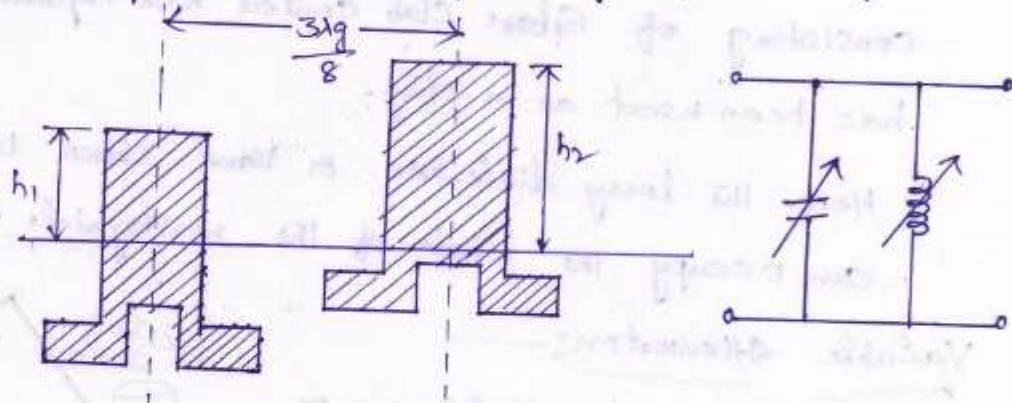
Tuning screws:—

Inductive and capacitive windows have the disadvantage that they cannot readily be made adjustable. Although it should theoretically be possible to vary window insertion, difficulty invariably results from imperfect contact between the diaphragm and the walls of the waveguide.



→ A combination of two screws $\lambda/4$ apart can be used to match a waveguide to its load similar to use of two fixed stubs in a Transmission line.

→ A very effective waveguide matching can be obtained when two tuning screws are placed in close proximity separated by $3\lambda/8$.



Microwave Attenuators :-

Attenuators are commonly used for measuring power gain or loss in dB's, for providing isolation between instruments, for reducing the power input to a particular stage to prevent overload and also for providing the signal generators with a means of calibrating their outputs accurately so that precise measurement could be made.

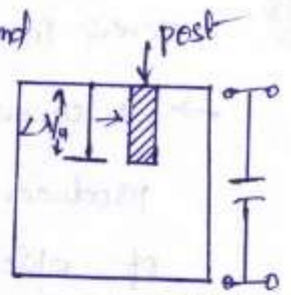
Attenuators are two types. They are (1) fixed Attenuators

(2) Variable (Continuous or step variation) Attenuators

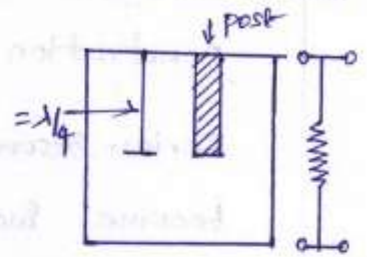
fixed Attenuators :-

These are used where fixed amount of attenuation to be provided. If such a fixed attenuator absorbs all the energy entering into it, it's called as "Waveguide Terminator".

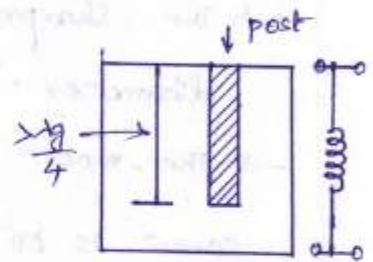
→ If the post extends only a short distance (i.e. $< \lambda_g/4$) into the waveguide it behaves capacitively, and this capacitive susceptance increases with the depth of penetration.



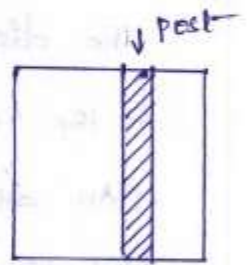
→ When the depth of penetration is equal to $\lambda_g/4$, the post acts as a series Resonant circuit.



→ If the depth of penetration is greater than $\lambda_g/4$, the post behaves inductively and this inductive susceptance decreases when the post is moved further away from the center of the waveguide.



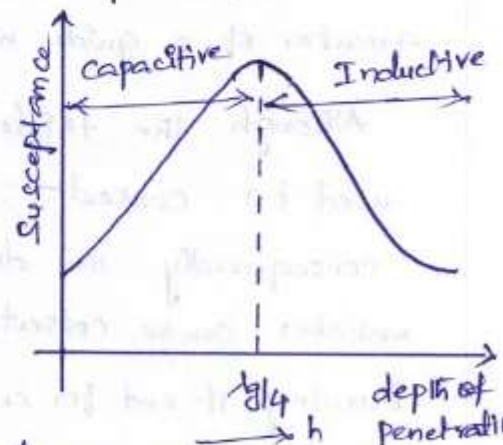
→ When the post is extended completely across the waveguide, the post becomes Inductive.



→ The susceptance vs penetration characteristics as shown below.

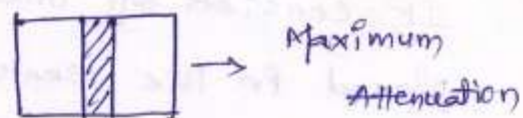
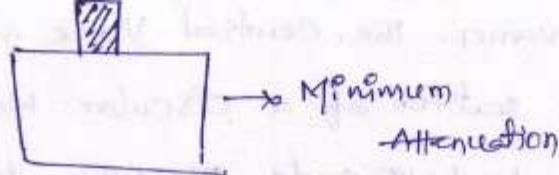
→ The amount of susceptance decreases as the diameter of the post is reduced.

→ If the post is made thicker, effective Q will be lowered and can act as a 'Band Pass Filter'.



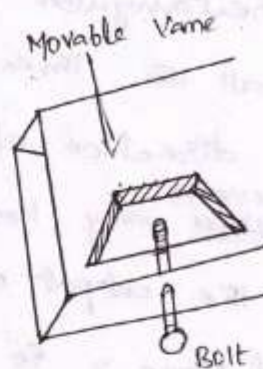
→ The main advantage of the post over the window is, it is readily adjustable. (this) Such an adjustable post is called 'screw or slug'.

→ As in case of posts depending upon the depth of penetration, the tuning screw may introduce inductive or capacitive susceptance.

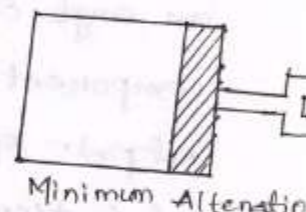
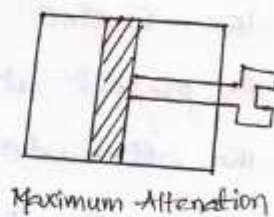


Vane Type Attenuator:

The vane type attenuator, basically consists of a glass vane with a coating of aquadag or carbon similar to a fixed vane attenuator.



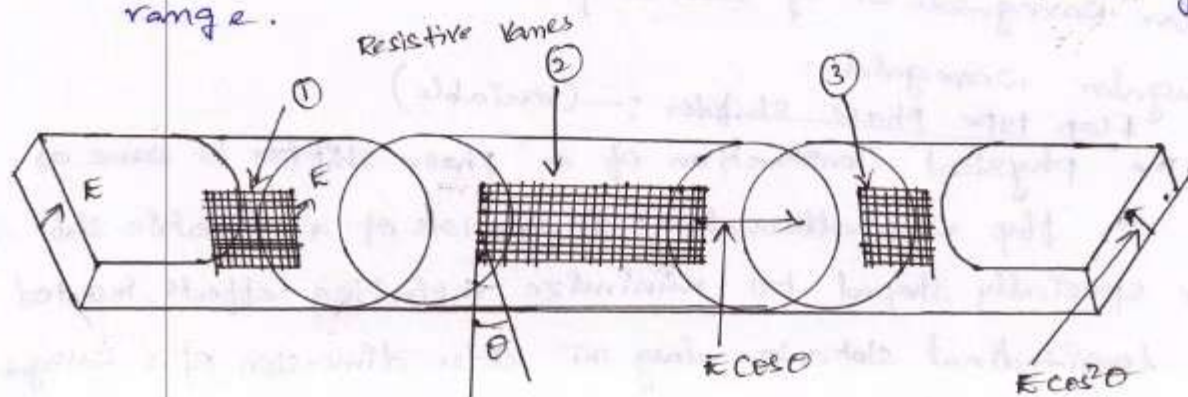
If this vane is used at the centre is made movable, it can be used as a variable attenuator. The vane positioned at the centre of the waveguide can be moved laterally from the centre of the waveguide where it provides maximum attenuation to the edges where the attenuation is considerably reduced since the field lines are always concentrated at the centre of the waveguide. The vane is tapered at both ends for matching the attenuator to the waveguide. An adequate match obtained if the taper length is made equal to $\lambda/2$.



The amount of attenuation is frequency sensitive and also has to be calibrated against a precision attenuator.

Rotary Vane Attenuator:

It provides precision attenuation with an accuracy of $\pm 2.1\%$ of the indicated attenuation over the operating frequency range.



This normally consists of a short section of a waveguide with a tapered plug of absorbing material at the end. The tapering is done for providing a gradual transition from the waveguide medium to the absorbing medium. Thus reducing the reflection occurring at the media interface.

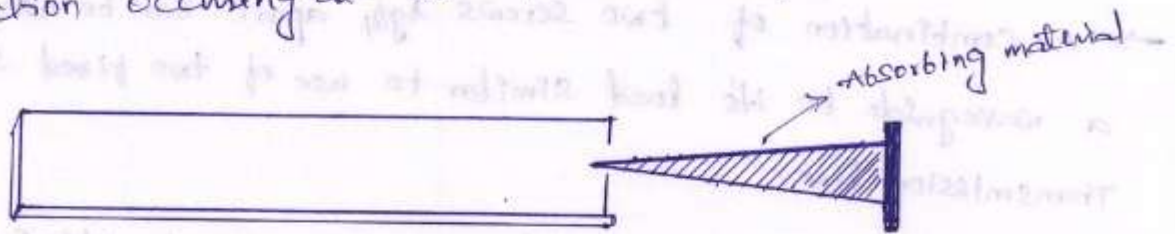


Figure shows fixed attenuator where a dielectric slab consisting of Glass slab coated with aquadag or Carbon film has been used as a plug.

Here the lossy dielectric or Vane shown is V-shaped and can occupy the whole of the waveguide.

Variable Attenuators:

It provides continuous or step wise variable attenuation.

For rectangular waveguides, these attenuators can be flap type or Vane type. For circular waveguides

rotary type is used.

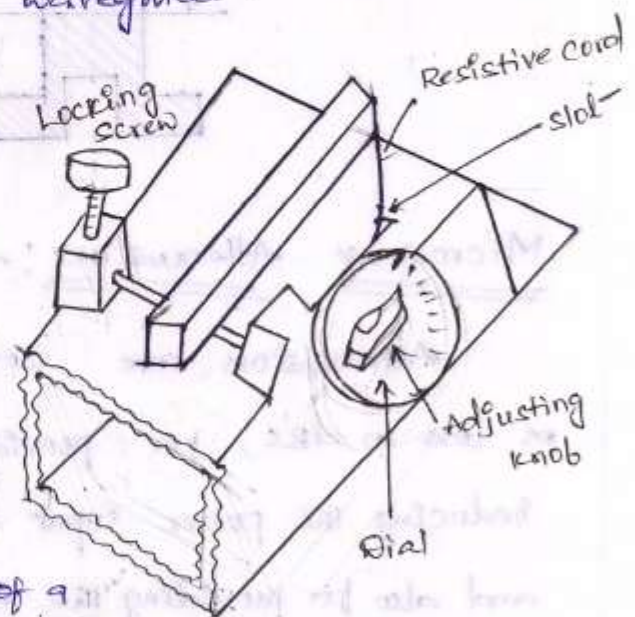
Flap type Attenuator:

→ The flap type attenuator consists of a

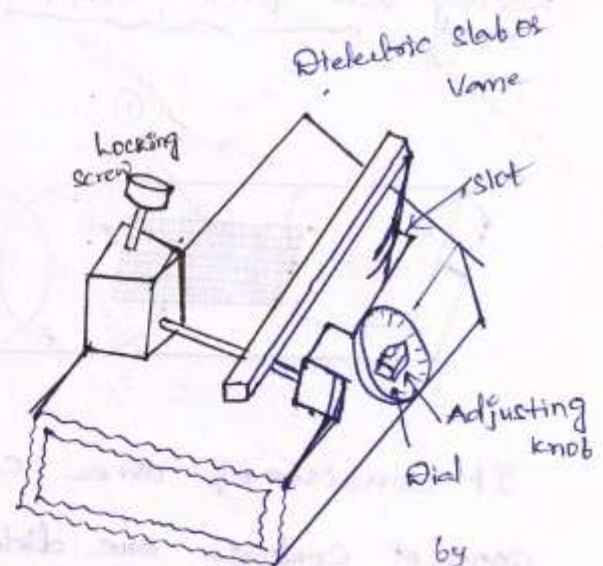
resistive element or disc inserted into

a longitudinal slot cut along the centre of the wider dimension of the guide.

The flap is mounted on the hinged arm allowing it to descend into the centre of the waveguide. The degree of attenuation is determined by the depth of insertion of the flap. However the flap attenuator does need to be calibrated against a standard as it is not a precision attenuator.



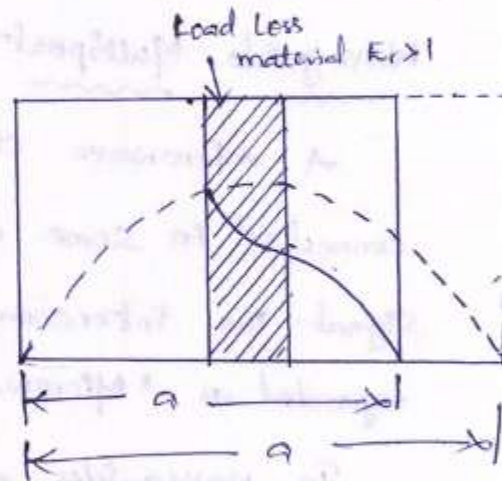
The dielectric slab is made of some low loss material (i.e. polyfoam) with $\epsilon_r > 1$. Higher the dielectric constant of the medium, the more slowly the micro-wave signal ^{travels} in a waveguide. Since most of the microwave signal in a waveguide travels through the centre of the waveguide, movement of a dielectric slab with $\epsilon_r > 1$ towards the centre of the



waveguide means that microwave signal moves more slowly ~~at~~ ^{by} the nearer the slab gets to the centre. The electric field distribution in the broader dimension of the waveguide will be modified by the dielectric slab so that it is distorted from sinusoidal to that indicated in figure below.

If the vane is inserted deeper, there is more change in the medium and there is a great phase shift.

The electric field distribution also shows that the dielectric has the same effect of increasing the broader dimension of the waveguide which reduces the wavelength in the waveguide. The amount of phase shift is maximum when the slab is at the centre and minimum when it is adjacent to the wall of the waveguide if the dielectric vane is placed such that the vane's inside dimension is parallel to the direction of the electric flux lines.



→ In the above dielectric ^{variable} vane phase shifter the phase shift is changed continuously from one value to another and hence they are also termed as "Analog phase shifters".

→ If a fixed phase shift is produced, then they are called "Discrete Digital phase shifters". Which are used in phased array antennas where as analog phase shifters are used in bridges and instruments.

It consists of three vanes. The central vane rotating type placed in the central section of a circular waveguide arrangement tapered at both ends. The other two vanes are in the rectangular sections.

When all the three vanes are aligned their planes are at 90° to the direction of the $\vec{E}$ -field. Hence there is no attenuation. Vane 1 prevents any horizontal polarizations and hence electric field at the output of vane 1 is vertically polarized.

The central vane 2 is rotating type and if it is rotated by an angle θ , the $\vec{E} \sin \theta$ component is attenuated and the $\vec{E} \cos \theta$ component is present at the output of vane 2 and the final output of the attenuator becomes $\vec{E} \cos^2 \theta$ which has the same polarization as the input wave.

The attenuation due to the rotary type vane attenuator is equal to $20 \log \cos^2 \theta = 40 \log \cos \theta$ which is independent of frequency.

Phase Shifters : —

Many applications require phase shifts to be introduced b/w two given positions in a waveguide system. The phase shift required may be fixed or variable.

fixed ^{TYPE} phase shifters :

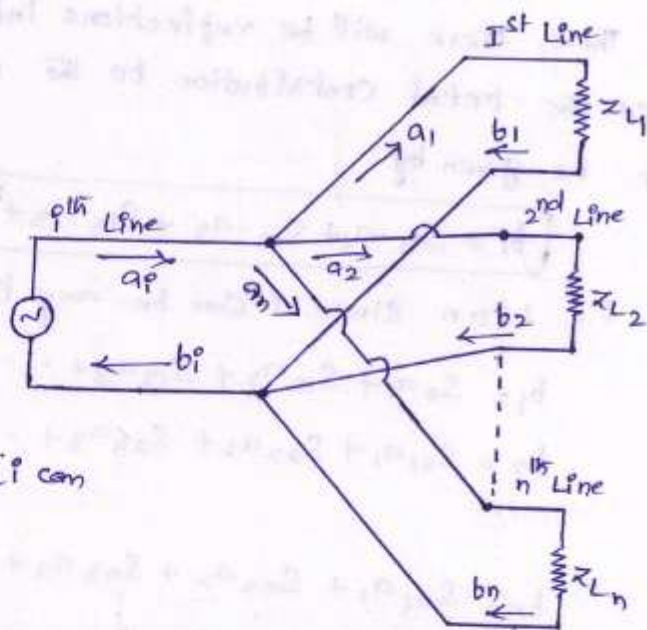
fixed amount of phase shift can be obtained by the use of capacitive / inductive pieces in the waveguide or by inserting dielectric rods across the diameters of a circular waveguide or by reducing the wider dimension of a rectangular waveguide.

flap type phase shifter : — (variable)

The physical construction of a phase shifter is same as that of a flap vane attenuator. It consists of a dielectric slab or vane specially shaped to minimize reflection effects, inserted through longitudinal slot cut along the wider dimension of a waveguide.

Junction the elements of this matrix are also called as 'scattering coefficients' of the scattering matrix.

To obtain the relationship between the scattering matrix and the input/output powers at different ports, consider a junction of n no. of transmission lines wherein the p 'th line terminated in a source [it can be any line from 1 to n].



Case (i) :-

Let the first line be terminated in an impedance other than its characteristic impedance (i.e., $Z_L \neq Z_0$) and all the remaining lines (from 2nd to n 'th line) in an impedance equal to Z_0 [i.e., $Z_L = Z_0$].

If a_p be the incident wave at the junction due to a source at the p 'th line, then it divides itself among $(n-1)$ number of lines as $a_1, a_2, \dots, a_n$. There will be no reflections from 2nd to n 'th line and the incident waves are absorbed since their impedances are equal to characteristic impedance (Z_0). But, there is a mismatch at the 1st line and hence there will be a reflected wave b_1 going back into the junction.

$$b_1 = (\text{reflection coefficient}) a_1 = S_{11} \cdot a_1$$

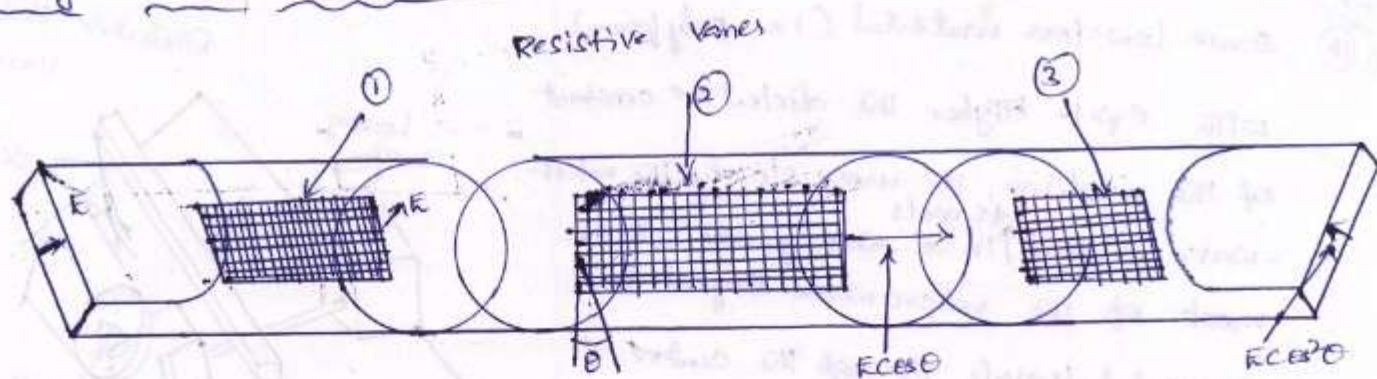
Where, S_{11} = reflection coefficient of 1st line
 Source connected at p 'th line Reflection from the 1st line.

$$b_2 = b_3 = \dots = b_n = 0$$

Hence, the contribution to the outward travelling wave in the p 'th line is given by

$$b_p = S_{11} \cdot a_1$$

Rotary Vane Phase Shifter: —



It consists of three circular waveguide sections all of which consist contain one dielectric vane. Center section is rotatable providing the required phase shift. It works on the principle of converting a linearly polarized TE_{11} mode into a circular polarized mode.

Waveguide Multipoint Junctions: —

A microwave circuit consists of several microwave devices connected in some way to achieve desired transmission of microwave signal. The interconnection of two or more microwave devices may be regarded as "Microwave Junction".

In waveguides, energy is transported by electromagnetic waves, not by the ^{in form of} voltages and currents by special coupling components

called "Waveguide Tees".

→ At low frequencies, circuits can be described by two port networks and their parameters such as Z, Y, h and $ABCD$ etc. Here network parameters relate the total voltages and total currents.

→ At microwave frequencies, we talk of travelling waves with associated powers instead of voltages and currents and the Microwave Junction can be defined by what are called "Scattering parameters or S-parameters". corresponding Matrix is called the "Scattering Matrix".

→ It is a square matrix which gives all the combinations of power relationships between the various input and output ports of Microwave

Properties of [S] Matrix:-

→ [S] is always a square matrix of order (n x n).

→ [S] is a symmetric Matrix i.e. $S_{ij} = S_{ji}$

→ [S] is a unitary Matrix, i.e. $[S][S]^* = [I]$

[I] → Unit Matrix, $S^* \rightarrow$ Complex Conjugate of [S].

→ The sum of the products of each term of any row (or column) multiplied by the complex conjugate of the corresponding terms of any other row (or column) is zero.

$$\text{i.e. } \sum_{i=1}^n S_{ik} S_{ij}^* = 0 \text{ for } k \neq j \quad \left(\begin{array}{l} k=1,2,3,\dots,n \\ j=1,2,3,\dots,n \end{array} \right)$$

→ A T-junction is an intersection of three waveguides in the form of English Alphabet "T". There are several types. They are

→ E-plane Tee Junction

→ H-plane Tee Junction

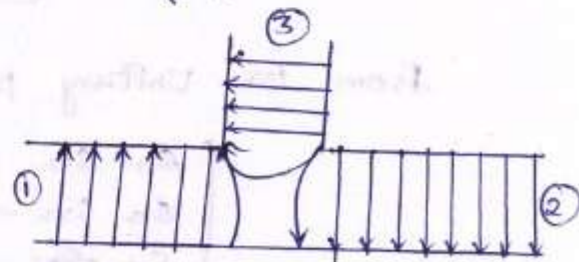
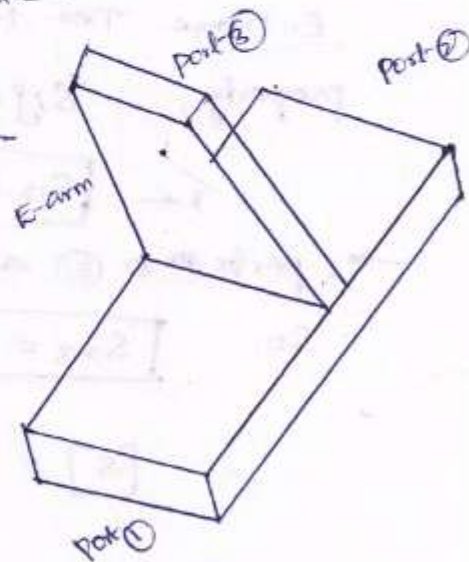
→ EH-plane (Hybrid) Junction

→ Rat Race Junction etc.

E-plane Tee (Series Tee) Junction:-

An E-plane Tee is a waveguide Tee, in which the axis of its side arm is parallel to the E-field of main guide.

A rectangular slot is cut along the broader dimension of a long waveguide and a side arm is attached. ports ① and ② are in the collinear arms and port ③ is the E-arm.



When TE₁₀ mode is made to propagate into port ③, the two outputs at port ① and ② will have a phase shift of 180°. Since the E-field lines change their direction when they come out of port ① and ②, it is called a "E-plane Tee".

Case ②: - Let all the $(n-1)$ lines be terminated in an impedance other than Z_0 . [i.e. $Z_L \neq Z_0$ for all the lines].

Then, there will be reflections into the junction from every line and hence the total contribution to the outward travelling wave in the i^{th} line is given by

$$b_i = S_{i1} \cdot a_1 + S_{i2} \cdot a_2 + S_{i3} \cdot a_3 + \dots + S_{in} \cdot a_n$$

$i = 1$ to n Since i can be any line from 1 to n . Therefore we have

$$b_1 = S_{11} a_1 + S_{12} a_2 + S_{13} a_3 + \dots + S_{1n} a_n$$

$$b_2 = S_{21} a_1 + S_{22} a_2 + S_{23} a_3 + \dots + S_{2n} a_n$$

$$\vdots$$

$$b_n = S_{n1} a_1 + S_{n2} a_2 + S_{n3} a_3 + \dots + S_{nn} a_n$$

In Matrix form,

$$\begin{bmatrix} b_1 \\ b_2 \\ \vdots \\ b_n \end{bmatrix} = \begin{bmatrix} S_{11} & S_{12} & \dots & S_{1n} \\ S_{21} & S_{22} & \dots & S_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ S_{n1} & S_{n2} & \dots & S_{nn} \end{bmatrix} \begin{bmatrix} a_1 \\ a_2 \\ \vdots \\ a_n \end{bmatrix}$$

Column Matrix $[b]$ corresponding to Reflected waves or output

Scattering Matrix $[S]$ of order $n \times n$

Matrix $[a]$ corresponding to Incident waves or Input.

$$\therefore [b] = [S][a]$$

$a_i \rightarrow$ inputs to particular ports

$b_i \rightarrow$ outputs ~~to~~ out of various ports

$S_{ij} \rightarrow$ scattering coefficients resulting due to input at j^{th} port and output taken out of i^{th} port.

$S_{ii} \rightarrow$ power reflected back from the junction into the i^{th} port - when input power is applied at the i^{th} port itself.

from ① & ②,

$$|s_{11}|^2 + |s_{12}|^2 + |s_{13}|^2 = |s_{12}|^2 + |s_{22}|^2 + |s_{13}|^2$$

$$\therefore \boxed{s_{11} = s_{22}}$$

from eqn ③,

$$2s_{13}^2 = 1 \Rightarrow s_{13}^2 = \frac{1}{2} \Rightarrow \boxed{s_{13} = \frac{1}{\sqrt{2}}}$$

from eqn ④,

$$s_{13}^* (s_{11} - s_{12}) = 0$$

$$\text{Since } s_{13} \neq 0, \therefore s_{12} - s_{11} = 0$$

$$\Rightarrow \boxed{s_{11} = s_{12}}$$

from equation ①,

$$|s_{11}|^2 + |s_{11}|^2 + |s_{13}|^2 = 1 \quad \text{Since } |s_{13}|^2 = \frac{1}{2}$$

$$2s_{11}^2 + \frac{1}{2} = 1 \Rightarrow 2s_{11}^2 = \frac{1}{2} \Rightarrow s_{11}^2 = \frac{1}{4}$$

$$\boxed{s_{11} = \frac{1}{2}}$$

$$\text{Since } s_{11} = s_{12} = s_{22}$$

$$\therefore s_{12} = s_{22} = \frac{1}{2}$$

$$[S] = \begin{bmatrix} s_{11} & s_{12} & s_{13} \\ s_{12} & s_{22} & -s_{13} \\ s_{13} & -s_{13} & 0 \end{bmatrix}$$

$$\therefore \boxed{[S] = \begin{bmatrix} \frac{1}{2} & \frac{1}{2} & \frac{1}{\sqrt{2}} \\ \frac{1}{2} & \frac{1}{2} & -\frac{1}{\sqrt{2}} \\ \frac{1}{\sqrt{2}} & -\frac{1}{\sqrt{2}} & 0 \end{bmatrix}}$$

$$\text{Since } [b] = [S][a]$$

$$\begin{bmatrix} b_1 \\ b_2 \\ b_3 \end{bmatrix} = \begin{bmatrix} \frac{1}{2} & \frac{1}{2} & \frac{1}{\sqrt{2}} \\ \frac{1}{2} & \frac{1}{2} & -\frac{1}{\sqrt{2}} \\ \frac{1}{\sqrt{2}} & -\frac{1}{\sqrt{2}} & 0 \end{bmatrix} \begin{bmatrix} a_1 \\ a_2 \\ a_3 \end{bmatrix}$$

$$\therefore b_1 = \frac{1}{2}a_1 + \frac{1}{2}a_2 + \frac{1}{\sqrt{2}}a_3$$

$$b_2 = \frac{1}{2}a_1 + \frac{1}{2}a_2 - \frac{1}{\sqrt{2}}a_3$$

$$b_3 = \frac{1}{\sqrt{2}}a_1 - \frac{1}{\sqrt{2}}a_2 //$$

Since it is symmetrical so any signal that is to be split or any two signals that are to be combined will be fed from the E-arm.

The scattering matrix of an E-plane Tee can be used to describe its properties. In general, the power out of port ③ is proportional to the difference between instantaneous power entering from ports ① & ②.
 → The effective value of the power leaving the E-arm is proportional to the phasor difference between the powers entering ports ① & ②.
 when power entering the main arms (ports ① & ②) are in phase opposition maximum energy comes out of port ③.

Since it is a 3-port junction, the S-matrix can be derived as follows.

$$[S] = \begin{bmatrix} S_{11} & S_{12} & S_{13} \\ S_{21} & S_{22} & S_{23} \\ S_{31} & S_{32} & S_{33} \end{bmatrix}$$

Assume that port ③ is perfectly matched to the junction. so $S_{33} = 0$

→ ports ① and ② are symmetrical with respect to port ③ i.e. E-plane Tee is a symmetrical device. so from the symmetry property, $S_{ij} = S_{ji}$

$$\Rightarrow \boxed{S_{12} = S_{21}, S_{31} = S_{13}, S_{32} = S_{23}}$$

→ ports ① & ② are in phase opposition with respect to port ③.

So $\boxed{S_{23} = -S_{13}}$

$$\therefore [S] = \begin{bmatrix} S_{11} & S_{12} & S_{13} \\ S_{12} & S_{22} & -S_{13} \\ S_{13} & -S_{13} & 0 \end{bmatrix}$$

from the Unitary property, $[S][S]^* = [I]$

$$\begin{bmatrix} S_{11} & S_{12} & S_{13} \\ S_{12} & S_{22} & -S_{13} \\ S_{13} & -S_{13} & 0 \end{bmatrix} \begin{bmatrix} S_{11}^* & S_{12}^* & S_{13}^* \\ S_{12}^* & S_{22}^* & -S_{13}^* \\ S_{13}^* & -S_{13}^* & 0 \end{bmatrix} = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix}$$

$$R_1 C_1 \Rightarrow S_{11} \cdot S_{11}^* + S_{12} \cdot S_{12}^* + S_{13} \cdot S_{13}^* = 1 \Rightarrow |S_{11}|^2 + |S_{12}|^2 + |S_{13}|^2 = 1 \rightarrow \textcircled{1}$$

$$R_2 C_2 \Rightarrow S_{12} \cdot S_{12}^* + S_{22} \cdot S_{22}^* + S_{13} \cdot S_{13}^* = 1 \Rightarrow |S_{12}|^2 + |S_{22}|^2 + |S_{13}|^2 = 1 \rightarrow \textcircled{2}$$

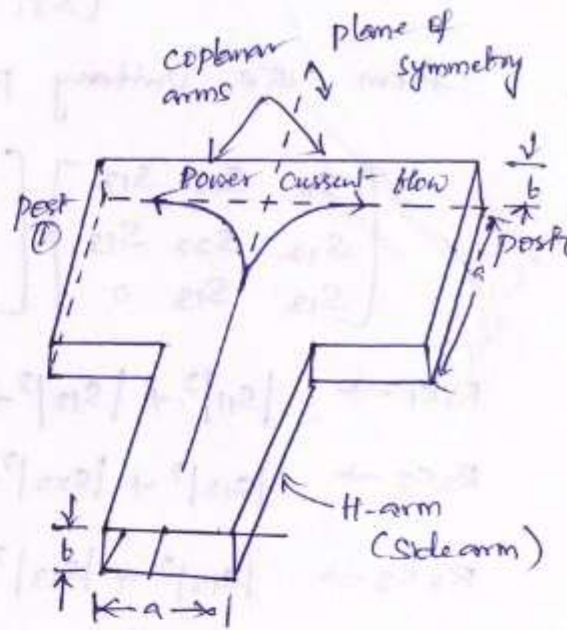
$$R_3 C_3 \Rightarrow S_{13} \cdot S_{13}^* + S_{23} \cdot S_{23}^* = 0 \Rightarrow |S_{13}|^2 + |S_{23}|^2 = 0 \rightarrow \textcircled{3}$$

$$R_1 C_3 \Rightarrow S_{11} \cdot S_{13}^* + S_{12} \cdot (-S_{13})^* = 0 \Rightarrow (S_{11} - S_{12}) S_{13}^* = 0 \rightarrow \textcircled{4}$$

H-plane Tee Junction: [or Shunt-Tee]

H-plane Tee junction is formed by cutting a rectangular slot along the width of a main waveguide and attaching another waveguide. The side arm is called "H-arm".

H-plane Tee is so called because the axis of the side arm is parallel to the plane of the main transmission line. As all the three arms of H-plane Tee lie in the plane of the magnetic field. The magnetic field divides itself into the arms. ports ① & ② are phase inphase with respect to port ③.



This H-plane Tee junction also called "Current Junction".

The port ① and ② of the main waveguide are called "Collinear arms or ports" and port ③ is the "H-arm or side arm".

The properties of a H-plane Tee can be completely defined by its $[S]$ matrix. The order of scattering matrix is 3×3 since there are three possible inputs and 3 possible outputs.

$$[S] = \begin{bmatrix} S_{11} & S_{12} & S_{13} \\ S_{21} & S_{22} & S_{23} \\ S_{31} & S_{32} & S_{33} \end{bmatrix}$$

Assume that port ③ is perfectly matched to the junction i.e., $S_{33} = 0$.

H-plane Tee is a symmetrical device. so from the symmetrical property, $S_{ij} = S_{ji}$

$$\text{i.e., } S_{21} = S_{12}, S_{31} = S_{13}, S_{32} = S_{23}$$

Since ports ① and ② are in inphase with respect to port ③.

$$\text{i.e., } S_{13} = S_{23}$$

Case ①: - Assume that Incidence signals into ① & ② are zero and there is an incidence signal into port ③ i.e., $a_1 = a_2 = 0, a_3 \neq 0$

$$\therefore b_1 = \frac{1}{\sqrt{2}} a_3, b_2 = -\frac{1}{\sqrt{2}} a_3, b_3 = 0.$$

i.e. An input at port ③ equally divides between ① and ② but introduces a phase shift of 180° between the two outputs. Hence E-plane Tee also acts as a 3dB splitter.

Since power is incident at port ③, which is equally divided to both ① & ② ^{but introduces 180° phase shift}. i.e., $P_3 = P_1 + P_2$ & $P_1 = P_2$, $P_3 = 2P_1 = 2P_2$

Power delivered from ① or ② in dB is

$$P_{dB} = 10 \log_{10} \frac{P_1}{P_3} = 10 \log_{10} \frac{P_1}{2P_1} = 10 \log_{10} \frac{1}{2} = \underline{\underline{-3 \text{ dB}}}$$

The power coming out of port ① or port ② is 3dB down with respect to input power at port ③, hence the E-plane Tee is called as "3-dB splitter".

Case ②: -

Assume that power incident into ports ① & ② is same and ^{no} incidence signal into port ③ i.e., $a_1 = a_2 = a, a_3 = 0$.

$$b_1 = \frac{a}{2} + \frac{a}{2} = a // \quad b_2 = \frac{a}{2} + \frac{a}{2} = a // \quad b_3 = 0.$$

i.e., equal inputs at port ① and port ② result in no output at port ③.

Case ③: - Assume that incidence signals into port ② and ③ are zero and there ^{is an} incidence signal into port ① i.e., $a_1 \neq 0, a_2 = a_3 = 0$.

$$\underline{\underline{b_1 = \frac{a_1}{2}, b_2 = \frac{a_1}{2}, b_3 = \frac{a_1}{\sqrt{2}}}}$$

case ①: — if $a_1 = a_2 = 0$ & $a_3 \neq 0$ [i.e. no incidence signals into ports ① & ②]

only the signal is incident into port ③]

$$\therefore b_1 = \frac{a_3}{\sqrt{2}}, \quad b_2 = \frac{a_3}{\sqrt{2}} \quad \text{and} \quad b_3 = 0.$$

Let P_3 be the power input at port ③, then this power divides equally between ports ① and ② in phase i.e., $P_1 = P_2$.

$$P_3 = P_1 + P_2 = P_1 + P_1 = 2P_1 = 2P_2.$$

The amount of power coming out of port ① or port ② due to input at port ③.

$$= 10 \log_{10} \frac{P_1}{P_3} = 10 \log_{10} \frac{P_1}{2P_1} = 10 \log_{10} \frac{1}{2} = -3 \text{ dB}$$

Hence the power coming out of port ① or port ② is 3 dB down with respect to input power at port ③, hence the H-plane Tee is called a "3-dB splitter".

Further when TE₁₀ mode is allowed to propagate into port ③, the electric field lines do not change their direction when they come out of ports ① and ② hence called "H-plane Tee", i.e. the waves that come out of ports ① and ② are equal in magnitude and phase.

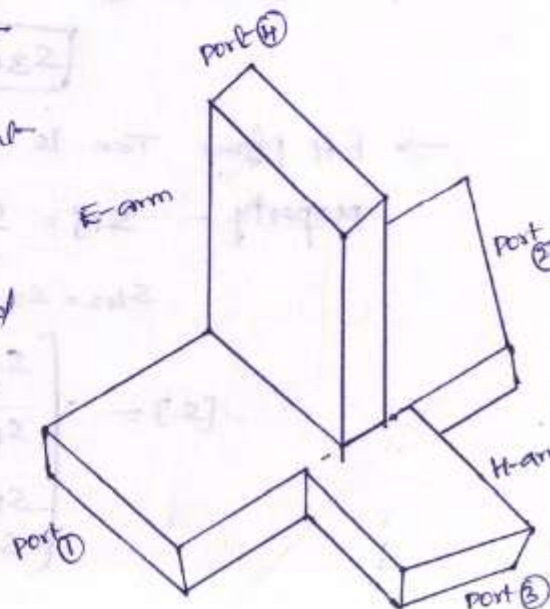
case ②: — if $a_1 = a_2 = a, a_3 = 0$

$$\therefore b_1 = \frac{a}{2} - \frac{a}{2} + \frac{1}{\sqrt{2}} a_3 = \frac{a_3}{\sqrt{2}} = 0 \quad // \quad b_2 = 0, \quad b_3 = \frac{a_1}{\sqrt{2}} + \frac{a_2}{\sqrt{2}} = \frac{a}{\sqrt{2}} + \frac{a}{\sqrt{2}}$$

i.e., The output at port ③ is addition of the two inputs at port ① and port ② and these are added in phase.

E-H Plane (Hybrid or Magic) Tee : —

In E-H plane Tee, rectangular slots are cut both along the width and breadth of a long waveguide and side arms are attached as shown in figure. so it is called as "Hybrid Tee". ports ① and ② are collinear arms, port ③ is the H-arm and port ④ is the E-arm.



$$\therefore S = \begin{bmatrix} S_{11} & S_{12} & S_{13} \\ S_{21} & S_{22} & S_{23} \\ S_{31} & S_{32} & S_{33} \end{bmatrix} = \begin{bmatrix} S_{11} & S_{12} & S_{13} \\ S_{12} & S_{22} & S_{13} \\ S_{13} & S_{13} & 0 \end{bmatrix}$$

from the unitary property, $[S][S]^* = [I]$

$$\begin{bmatrix} S_{11} & S_{12} & S_{13} \\ S_{12} & S_{22} & S_{13} \\ S_{13} & S_{13} & 0 \end{bmatrix} \begin{bmatrix} S_{11}^* & S_{12}^* & S_{13}^* \\ S_{12}^* & S_{22}^* & S_{13}^* \\ S_{13}^* & S_{13}^* & 0 \end{bmatrix} = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix}$$

$$R_1 C_1 \rightarrow |S_{11}|^2 + |S_{12}|^2 + |S_{13}|^2 = 1 \rightarrow \textcircled{1}$$

$$R_2 C_2 \rightarrow |S_{12}|^2 + |S_{22}|^2 + |S_{13}|^2 = 1 \rightarrow \textcircled{2}$$

$$R_3 C_3 \rightarrow |S_{13}|^2 + |S_{13}|^2 = 1 \rightarrow \textcircled{3}$$

$$R_1 C_3 \rightarrow S_{11} \cdot S_{13}^* + S_{12} \cdot S_{13}^* = 0 \rightarrow \textcircled{4}$$

from eqn(3), $2|S_{13}|^2 = 1 \Rightarrow |S_{13}|^2 = \frac{1}{2} \Rightarrow \boxed{S_{13} = \frac{1}{\sqrt{2}}}$

from (1) & (2), $\boxed{S_{11} = S_{22}}$

from equation (4), $S_{13}^* [S_{11} + S_{12}] = 0$

since $S_{13} \neq 0$, $S_{11} + S_{12} = 0 \Rightarrow \boxed{S_{12} = -S_{11}}$

$\therefore S_{11} = S_{22} = -S_{12}$

from equation (1), $|S_{11}|^2 + |S_{11}|^2 + \frac{1}{2} = 1 \Rightarrow 2|S_{11}|^2 = \frac{1}{2} \Rightarrow |S_{11}|^2 = \frac{1}{4}$

$$\boxed{S_{11} = \frac{1}{2} = S_{22}}$$

$$\underline{\underline{S_{12} = -\frac{1}{2}}}$$

$$\therefore [S] = \begin{bmatrix} \frac{1}{2} & -\frac{1}{2} & \frac{1}{\sqrt{2}} \\ -\frac{1}{2} & \frac{1}{2} & \frac{1}{\sqrt{2}} \\ \frac{1}{\sqrt{2}} & \frac{1}{\sqrt{2}} & 0 \end{bmatrix}$$

$$[b] = [S][a] \quad \begin{bmatrix} b_1 \\ b_2 \\ b_3 \end{bmatrix} = \begin{bmatrix} \frac{1}{2} & -\frac{1}{2} & \frac{1}{\sqrt{2}} \\ -\frac{1}{2} & \frac{1}{2} & \frac{1}{\sqrt{2}} \\ \frac{1}{\sqrt{2}} & \frac{1}{\sqrt{2}} & 0 \end{bmatrix} \begin{bmatrix} a_1 \\ a_2 \\ a_3 \end{bmatrix}$$

$$\therefore b_1 = \frac{1}{2}a_1 - \frac{1}{2}a_2 + \frac{1}{\sqrt{2}}a_3$$

$$\therefore b_2 = -\frac{1}{2}a_1 + \frac{1}{2}a_2 + \frac{1}{\sqrt{2}}a_3$$

$$b_3 = \frac{1}{\sqrt{2}}a_1 + \frac{1}{\sqrt{2}}a_2$$

from the unitary property, $[S][S]^* = [I]$

$$\begin{bmatrix} S_{11} & S_{12} & S_{13} & S_{14} \\ S_{12} & S_{22} & S_{13} & -S_{14} \\ S_{13} & S_{13} & 0 & 0 \\ S_{14} & -S_{14} & 0 & 0 \end{bmatrix} \begin{bmatrix} S_{11}^* & S_{12}^* & S_{13}^* & S_{14}^* \\ S_{12}^* & S_{22}^* & S_{13}^* & -S_{14}^* \\ S_{13}^* & S_{13}^* & 0 & 0 \\ S_{14}^* & -S_{14}^* & 0 & 0 \end{bmatrix} = \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{bmatrix}$$

$$R_1 C_1: |S_{11}|^2 + |S_{12}|^2 + |S_{13}|^2 + |S_{14}|^2 = 1 \rightarrow \textcircled{1}$$

$$R_2 C_2: |S_{12}|^2 + |S_{22}|^2 + |S_{13}|^2 + |S_{14}|^2 = 1 \rightarrow \textcircled{2}$$

$$R_3 C_3: |S_{13}|^2 + |S_{13}|^2 = 1 \Rightarrow 2|S_{13}|^2 = 1 \Rightarrow \boxed{S_{13} = \frac{1}{\sqrt{2}}}$$

$$R_4 C_4: |S_{14}|^2 + |S_{14}|^2 = 1 \Rightarrow 2|S_{14}|^2 = 1 \Rightarrow \boxed{S_{14} = \frac{1}{\sqrt{2}}}$$

from eqn ① & ②, $\boxed{S_{11} = S_{22}}$

from eqn ①,

$$|S_{11}|^2 + |S_{12}|^2 + \frac{1}{2} + \frac{1}{2} = 1 \Rightarrow |S_{11}|^2 + |S_{12}|^2 = 0 \rightarrow \textcircled{3}$$

from eqn ②,

$$|S_{12}|^2 + |S_{22}|^2 + \frac{1}{2} + \frac{1}{2} = 1 \Rightarrow |S_{12}|^2 + |S_{22}|^2 = 0 \rightarrow \textcircled{4}$$

$$R_1 C_4: S_{11} S_{14}^* - S_{12} S_{14}^* = 0 \Rightarrow S_{14}^* (S_{11} - S_{12}) = 0$$

$$\text{since } S_{14} \neq 0, \therefore S_{14}^* \neq 0, S_{11} - S_{12} = 0$$

$$\boxed{S_{11} = S_{12}}$$

$$\therefore \boxed{S_{11} = S_{12} = S_{22}}$$

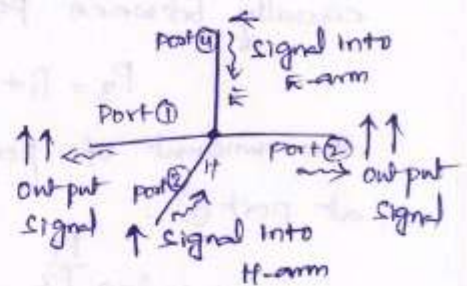
from eqn ③,

$$|S_{11}|^2 + |S_{11}|^2 = 0 \Rightarrow \boxed{S_{11} = 0}$$

$$\therefore \boxed{S_{12} = S_{22} = 0} \quad \text{since } S_{11} = S_{12} = S_{22}$$

This means that ports ① and ② are also perfectly matched to the junction. Hence in any four port junction, if any two ports are perfectly matched to the junction, then the remaining two ports are automatically matched to the junction. Such a junction where in all the four ports are automatically matched to the junction is called

Such a device became necessary because of the difficulty of obtaining a completely matched three port junction. This four port hybrid Tee junction combines the power dividing properties of both H-plane Tee and E-plane Tee and has the advantage of being completely matched at all its ports. Using the properties of E-H Plane Tee, its scattering matrix can be obtained as follows.



Since there are four ports, so

$[S]$ is a 4×4 matrix, i.e.

$$[S] = \begin{bmatrix} S_{11} & S_{12} & S_{13} & S_{14} \\ S_{21} & S_{22} & S_{23} & S_{24} \\ S_{31} & S_{32} & S_{33} & S_{34} \\ S_{41} & S_{42} & S_{43} & S_{44} \end{bmatrix}$$

Assume that ports ② & ④ are perfectly matched to the junction.

so, $S_{33} = S_{44} = 0$.

When E-field lines are reached to the collinear arms from the E-arm, they are in phase opposition at the o/p of ports ① and ②.

$$S_{24} = -S_{14}$$

If we observe from H-arm into ports ① and ②, Magnetic field lines are in phase.

$$\therefore S_{23} = S_{13}$$

ports ③ and ④ are isolated to each other. so

$$S_{34} = S_{43} = 0$$

→ E-H plane Tee is a symmetrical device. from the symmetry property. $S_{ij} = S_{ji}$ i.e., $S_{12} = S_{21}$, $S_{31} = S_{13}$, $S_{41} = S_{14}$, $S_{32} = S_{23}$,

$$S_{42} = S_{24}, S_{43} = S_{34}$$

$$\therefore [S] = \begin{bmatrix} S_{11} & S_{12} & S_{13} & S_{14} \\ S_{12} & S_{22} & S_{13} & -S_{14} \\ S_{13} & S_{13} & 0 & 0 \\ S_{14} & -S_{14} & 0 & 0 \end{bmatrix}$$

Case (4): - when $a_1 = a_2 = 0$ and $a_3 = a_4$

$$\text{then } b_1 = \frac{1}{\sqrt{2}} (2a_3), \quad b_2 = 0, \quad b_3 = b_4 = 0.$$

This is nothing but the additive property. Equal inputs at ports and (4) results in an output at port (1) [in phase and equal in amplitude].

Case (5): - when $a_1 = a_2, a_3 = a_4 = 0$

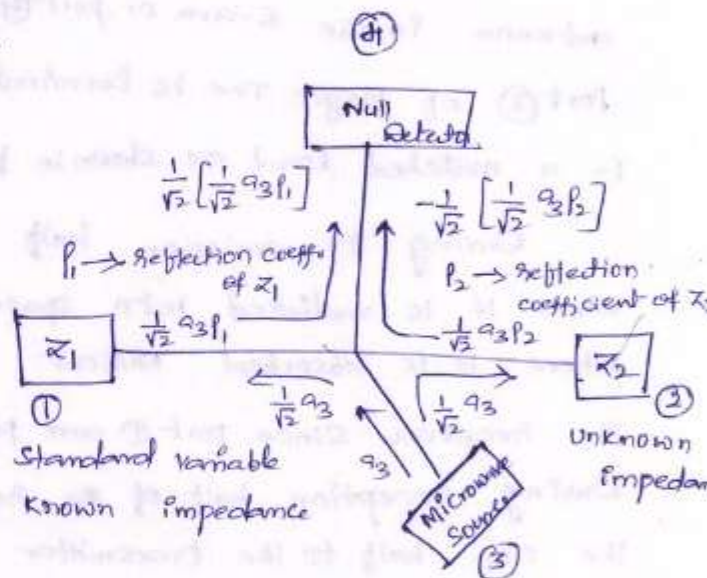
$$\text{then } b_1 = 0 = b_2 = b_4, \quad b_3 = \frac{1}{\sqrt{2}} (2a_1)$$

equal inputs at ports (1) and (2) results in an output at port (3) and no outputs at port (1), (2) and (4).

Applications of Magic Tee:

(a) Measurement of Impedance:

Microwave source is connected in arm (3) and a null detector in arm (4). The unknown impedance is connected in arm (2) and a standard variable known impedance in arm (1). Using the properties of Magic Tee, the power from microwave



source (a_3) gets divided equally between arms (1) and (2) $\left[\frac{a_3}{\sqrt{2}} \right]$. These impedances are not equal to characteristic impedance Z_0 and hence there will be reflections from arms (1) and (2).

If p_1 and p_2 are the reflection coefficients, powers $\frac{p_1 a_3}{\sqrt{2}}$ and $\frac{p_2 a_3}{\sqrt{2}}$ enter the Magic Tee junction from arms (1) and (2) as shown. The resultant wave into arm (4) i.e. the null detector can be calculated as follows:

The net wave reaching the null detector

$$= \frac{1}{\sqrt{2}} \left[\frac{1}{\sqrt{2}} a_3 p_1 \right] - \frac{1}{\sqrt{2}} \left[\frac{1}{\sqrt{2}} a_3 p_2 \right] = \frac{1}{2} a_3 [p_1 - p_2]$$

for perfect balancing of the bridge (null detection) is equated to zero.

a "Magic" tee.

Then $[S]$ matrix becomes

$$[S] = \begin{bmatrix} 0 & 0 & \frac{1}{\sqrt{2}} & \frac{1}{\sqrt{2}} \\ 0 & 0 & \frac{1}{\sqrt{2}} & -\frac{1}{\sqrt{2}} \\ \frac{1}{\sqrt{2}} + \frac{1}{\sqrt{2}} & 0 & 0 & 0 \\ \frac{1}{\sqrt{2}} - \frac{1}{\sqrt{2}} & 0 & 0 & 0 \end{bmatrix}$$

Since $[b] = [S][a]$

$$\text{i.e.} \begin{bmatrix} b_1 \\ b_2 \\ b_3 \\ b_4 \end{bmatrix} = \begin{bmatrix} 0 & 0 & \frac{1}{\sqrt{2}} & \frac{1}{\sqrt{2}} \\ 0 & 0 & \frac{1}{\sqrt{2}} & -\frac{1}{\sqrt{2}} \\ \frac{1}{\sqrt{2}} & \frac{1}{\sqrt{2}} & 0 & 0 \\ \frac{1}{\sqrt{2}} & -\frac{1}{\sqrt{2}} & 0 & 0 \end{bmatrix} \begin{bmatrix} a_1 \\ a_2 \\ a_3 \\ a_4 \end{bmatrix}$$

$$\therefore \boxed{\begin{aligned} b_1 &= \frac{1}{\sqrt{2}} a_3 + \frac{1}{\sqrt{2}} a_4, & b_2 &= \frac{1}{\sqrt{2}} a_3 - \frac{1}{\sqrt{2}} a_4 \\ b_3 &= \frac{1}{\sqrt{2}} a_1 + \frac{1}{\sqrt{2}} a_2, & b_4 &= \frac{1}{\sqrt{2}} a_1 - \frac{1}{\sqrt{2}} a_2. \end{aligned}}$$

Case 1:- if power enter into all ports except into port-③ is zero i.e. no incidence signal into ports ①, ② and ④ then
i.e. $a_1 = a_2 = a_4 = 0, a_3 \neq 0$

$$b_1 = \frac{a_3}{\sqrt{2}}, \quad b_2 = \frac{a_3}{\sqrt{2}}, \quad b_3 = 0, \quad b_4 = 0$$

This is the property of H-plane Tee. ports ③ and ④ are isolated since nothing out of port ④ when power is fed into port ③.

case ②:- If power incident into ports ①, ② and ③ is zero and there is an incidence signal into port-④ i.e. $a_1 = a_2 = a_3 = 0, a_4 \neq 0$.

$$b_1 = \frac{a_4}{\sqrt{2}}, \quad b_2 = -\frac{a_4}{\sqrt{2}}, \quad b_3 = 0, \quad b_4 = 0.$$

This is the property of E-plane Tee.

case ③:- if $a_1 \neq 0, a_2 = a_3 = a_4 = 0$

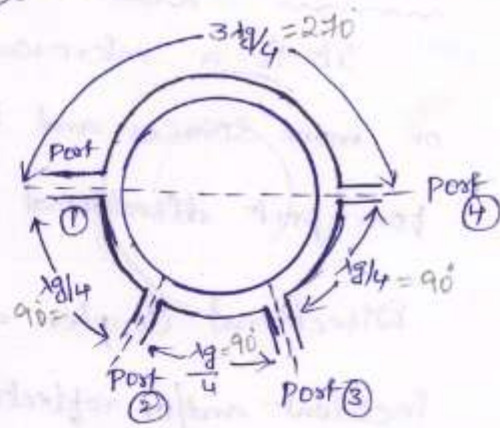
$$\therefore b_1 = 0, \quad b_2 = 0, \quad b_3 = \frac{a_1}{\sqrt{2}}, \quad b_4 = \frac{a_1}{\sqrt{2}}$$

when power is fed into port ①, nothing comes out of port ② even though they are collinear ports. Hence ports ① and ② are called 'Isolated ports'. Similarly an input at port ② cannot come out at port ①. Similarly E and H are isolated ports.

Rat Race Junction : (Hybrid Rings) :

This is a four port Junction,
the fourth port being added to a
normal three port Tee.

The four ports are connected in the
form of an angular ring at proper intervals
by means of series or parallel junctions.



These ports are separated by proper electrical lengths to sustain standing waves. for proper operation, it is necessary that the mean circumference of the total race be $1.5\lambda_g$ and that each of the four ports be separated from its neighbor by a distance of $\lambda_g/4$.

When power is fed into port ① it is equally into ports ② and ④ and nothing enters port ③. At ports ② and ④ the power combine in phase but at port ③ cancellation occurs due to $\lambda_g/2$ path difference. Similarly, input applied at port ③ is equally divided between ports ② and ④ but the output at port ① will be zero.

The rat race can be used for combining two signals or dividing a single signal into two equal halves.

If two unequal signals are applied at port ①, an output proportional to their sum will emerge from ports ② and ④ while a differential output will appear at port ③.

The Scattering Matrix of a rat race Junction can be written as

$$[S] = \begin{bmatrix} 0 & S_{12} & 0 & S_{14} \\ S_{21} & 0 & S_{23} & 0 \\ 0 & S_{32} & 0 & S_{34} \\ S_{41} & 0 & S_{43} & 0 \end{bmatrix}$$

in this

$$\lambda_g = 360^\circ$$

$$\lambda_g/4 = 90^\circ$$

$$\text{i.e., } \frac{1}{2} a_3 (I_1 - I_2) = 0$$

$$\text{or } I_1 - I_2 = 0 \quad \text{or } I_1 = I_2$$

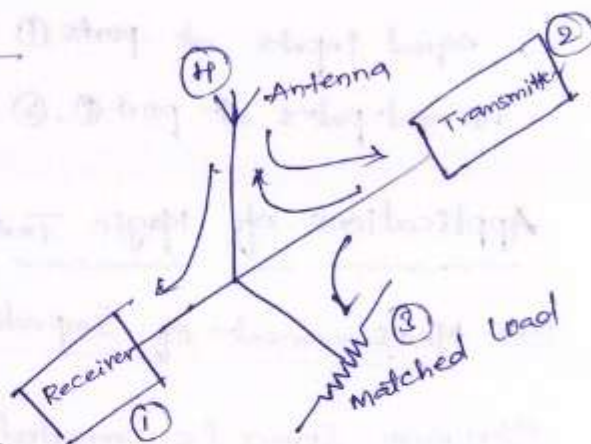
$$\text{or } \frac{Z_1 - Z_2}{Z_1 + Z_2} = \frac{Z_2 - Z_2}{Z_2 + Z_2} \quad \therefore \boxed{Z_1 = Z_2}$$

$$\boxed{R_1 = R_2} \quad \text{and} \quad \boxed{X_1 = X_2} \quad R_1 + jX_1 = R_2 + jX_2$$

Thus the unknown impedance can be measured by adjusting the standard variable impedance till the bridge is balanced and both impedances become equal.

(b) Magic Tee as a Duplexer:

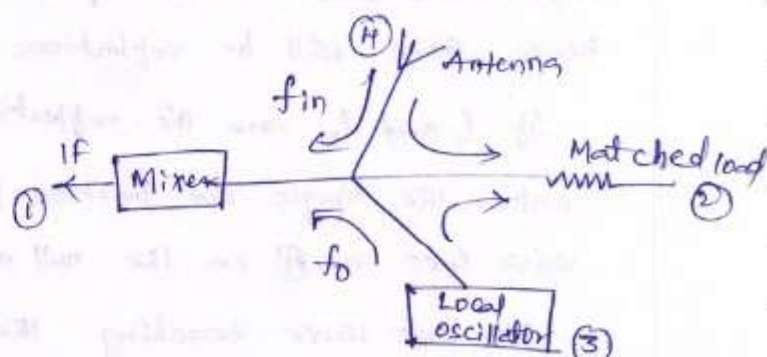
The transmitter and receiver are connected in ports ② and ① respectively, antenna in the E-arm or port ④ and port ③ of Magic Tee is terminated in a matched load as shown in figure.



During transmission, half the power reaches the antenna from where it is radiated into space. other half reaches the matched load where it is absorbed without reflections. No transmitter power reaches the receiver since port ① and port ② are isolated ports in Magic Tee. During reception, half of the received power goes to the receiver and the other half to the transmitter are isolated during reception as well as during transmission.

(c) Magic Tee as a Mixer:

A magic Tee can be used in Microwave receivers as a Mixer where the signal and local oscillator are fed into the E and H-arms.



Half of the Local oscillator power and half of the received power from antenna goes to the mixer where they are mixed to generate the IF frequency.

$$\boxed{IF = f_{in} - f_0}$$

Coupling factor (C):

The Coupling factor is defined as the ratio of incident power $[P_i]$ to the forward power $[P_f]$ measured in dB.

$$C = 10 \log_{10} \left[\frac{P_i}{P_f} \right] = 10 \log_{10} \left[\frac{P_1}{P_4} \right] \text{ dB}$$

Directivity:

It is defined as the ratio of forward power (P_f) to the back power (P_b) expressed in dB.

$$D = 10 \log_{10} \left[\frac{P_f}{P_b} \right] \text{ dB} = 10 \log_{10} \left[\frac{P_4}{P_3} \right] \text{ dB}$$

Isolation (I):

It is defined as the ratio of the incident power $[P_i]$ to the back power $[P_b]$ expressed in dB.

$$I = 10 \log_{10} \left[\frac{P_i}{P_b} \right] \text{ dB} = 10 \log_{10} \left[\frac{P_1}{P_3} \right] \text{ dB}$$

$$I = 10 \log_{10} \left[\frac{P_i}{P_3} \right] = 10 \log_{10} \left[\frac{P_1}{P_4} \cdot \frac{P_4}{P_3} \right]$$

$$I = 10 \log_{10} \left[\frac{P_1}{P_4} \right] + 10 \log_{10} \left[\frac{P_4}{P_3} \right]$$

$$\boxed{I = C + D}$$

Scattering Matrix of a Directional Coupler:

→ It is a four port network. Hence $[S]$ is a 4×4 matrix.

$$[S] = \begin{bmatrix} S_{11} & S_{12} & S_{13} & S_{14} \\ S_{21} & S_{22} & S_{23} & S_{24} \\ S_{31} & S_{32} & S_{33} & S_{34} \\ S_{41} & S_{42} & S_{43} & S_{44} \end{bmatrix}$$

→ In a directional coupler all four ports are matched perfectly to the junction. Hence

$$\boxed{S_{11} = S_{22} = S_{33} = S_{44} = 0}$$

→ It is a symmetrical device, so from the symmetry property, $S_{ij} = S_{ji}$

$$\therefore \boxed{S_{23} = S_{32}, S_{13} = S_{31}, S_{24} = S_{42}, S_{34} = S_{43}, S_{41} = S_{14}}$$

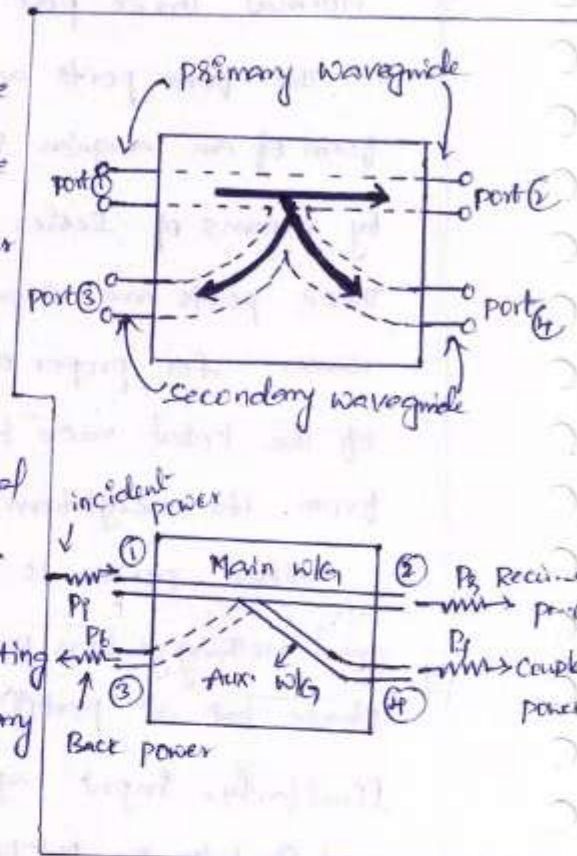
IV Directional Coupler: — [Hybrid Rings]

It is a microwave junction which receive energy from two or more sources and transmit them to two or more outlets. A four port directional coupler is show in figure below.

Directional Couplers are designed to measure incident and/or reflected powers, standing wave Ratio values, provide a signal path to a receiver or perform other desirable operations.

Directional Couplers can be unidirectional (measuring only incident power) or bi-directional (measuring both incident and reflected powers).

It is a four port waveguide junction consisting of a primary main waveguide and a secondary auxiliary waveguide as shown.



With matched terminations at all its ports, the properties of an ideal directional coupler can be summarized as follows:

- When the power is incident into port ①, portion of the incident power is coupled to ports ② and ④ but not to port ③.
- When the power is incident into port ②, portion of the incident power is coupled to ports ① and ③ but not to port ④.
- When the power is incident into port ③, portion of the incident power is coupled to ports ① and ④ but not to port ②. Similarly
- When the power is incident into port ④, portion of the incident power is coupled to ports ② and ③ but not to port ①.

The performance of a directional coupler is usually defined in terms of two parameters which are defined as follows:

Therefore $[S]$ matrix of a directional coupler is reduced

to

$$[S] = \begin{bmatrix} 0 & P & 0 & jQ \\ P & 0 & jQ & 0 \\ 0 & jQ & 0 & P \\ jQ & 0 & P & 0 \end{bmatrix}$$

In this directional coupler, two waveguides (primary and Auxiliary/secondary) share a common wall. This common wall has got hole or holes for coupling the energy flowing into the main waveguide to the side waveguide and hence called a "side hole coupler".

→ If single hole at the common wall is used to couple the energy that is called "single hole directional coupler".

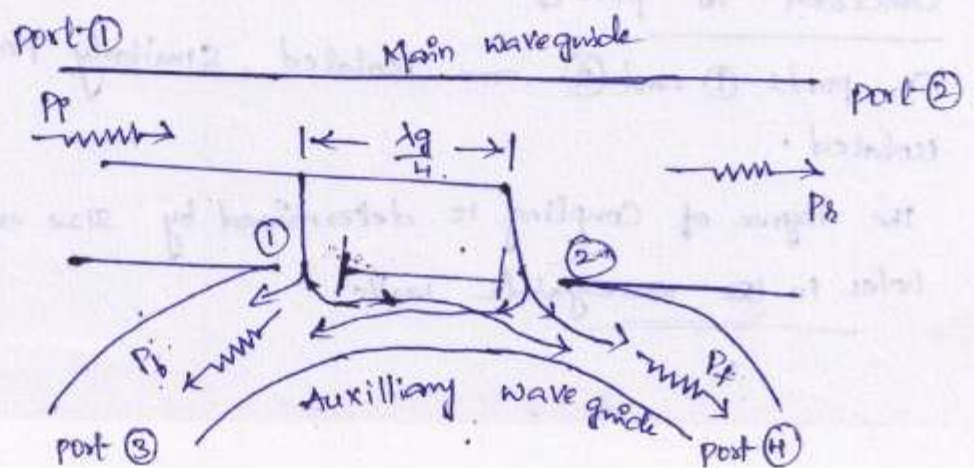
→ If holes are on the common wall to couple the energy that is called "Two hole directional coupler".

→ If many holes are on the common wall to couple the energy then that directional coupler called "Multi Hole directional coupler".

*→ A Two hole directional coupler is most commonly used.

Two-hole directional coupler:

The principle of operation of a two-hole directional coupler is shown in figure below. It consists of two guides the main and the Auxiliary with two tiny holes common between them as shown.



Since there is no coupling between port ① and port ③ and ports ② and ④.

$$\text{i.e. } S_{13} = S_{31} = 0, \quad S_{24} = 0 = S_{42}$$

$$[S] = \begin{bmatrix} 0 & S_{12} & 0 & S_{14} \\ S_{12} & 0 & S_{23} & 0 \\ 0 & S_{23} & 0 & S_{34} \\ S_{14} & 0 & S_{34} & 0 \end{bmatrix}$$

from the unitary property, $[S][S]^* = [I]$

$$\begin{bmatrix} 0 & S_{12} & 0 & S_{14} \\ S_{12} & 0 & S_{23} & 0 \\ 0 & S_{23} & 0 & S_{34} \\ S_{14} & 0 & S_{34} & 0 \end{bmatrix} \begin{bmatrix} 0 & S_{12}^* & 0 & S_{14}^* \\ S_{12}^* & 0 & S_{23}^* & 0 \\ 0 & S_{23}^* & 0 & S_{34}^* \\ S_{14}^* & 0 & S_{34}^* & 0 \end{bmatrix} = \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{bmatrix}$$

$$R_1 C_1 \Rightarrow |S_{12}|^2 + |S_{14}|^2 = 1 \longrightarrow \textcircled{1}$$

$$R_2 C_2 \Rightarrow |S_{12}|^2 + |S_{23}|^2 = 1 \longrightarrow \textcircled{2}$$

$$R_3 C_3 \Rightarrow |S_{23}|^2 + |S_{34}|^2 = 1 \longrightarrow \textcircled{3}$$

$$R_4 C_4 \Rightarrow |S_{14}|^2 + |S_{34}|^2 = 1 \longrightarrow \textcircled{4}$$

$$R_1 C_3 \Rightarrow S_{12} S_{23}^* + S_{14} S_{34}^* = 0 \longrightarrow \textcircled{5}$$

from ① & ②, $S_{14} = S_{23}$

from ① & ③, $S_{12} = S_{34}$

Let us assume that S_{12} is real and positive = P

$$S_{12} = S_{34} = P = S_{34}^* \longrightarrow \textcircled{6}$$

from eqns ⑤ & ⑥,

$$P S_{23}^* + S_{23} P = 0$$

$$P [S_{23} + S_{23}^*] = 0$$

$$\text{Since } P \neq 0, \quad S_{23} + S_{23}^* = 0$$

$$\text{if } S_{23} = jy, \quad S_{23}^* = -jy$$

i.e., S_{23} must be imaginary.

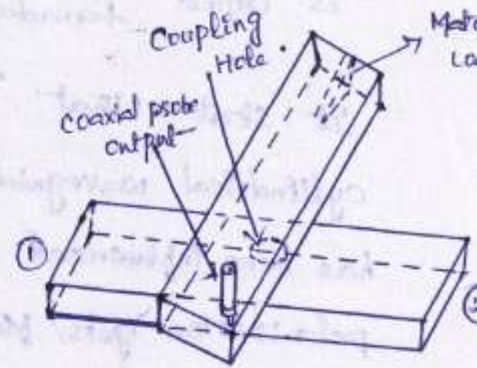
$$\text{Let } S_{23} = jy = S_{14}$$

$$\text{and therefore, } S_{12} = S_{34} = P$$

$$S_{23} = S_{14} = jy$$

The (or) Single hole Coupler:

A single hole directional coupler is shown in figure. In this directivity is improved. The power entering at port ① is coupled to the co-axial probe output and the power entering at port ② is absorbed by the matched load.



The Auxiliary guide is placed at such an angle that the magnitude of the magnetically excited wave is made equal to that of the electrically excited wave for improved directivity.

In this coupler, the waves in the auxiliary guide are generated through a single hole which includes both electric and magnetic fields. Because of the phase relationships involved in the coupling process, the signals generated by the two types of coupling cancel in forward direction and reinforce in the reverse direction.

ferrites (or ferrite devices):

ferrites are non-metallic materials with high resistivities and dielectric constants. They show magnetic properties similar to ferromagnetic metals. ferrite materials are oxide based components with general composition of the form $MnO \cdot Fe_2O_3$. They are obtained by firing powdered oxides of materials at $1100^\circ C$ or more and pressing them into different shapes. This gives additional features of ceramic insulators so that they can be used at microwave frequencies. ferrites exhibit strong magnetic property due to the magnetic dipole moment associated with the electron spin.

- Because of high resistivity, ferrites can be used up to 100 GHz.
- ferrites have a peculiar property which is non-reciprocal property so which is useful at microwave frequencies.

The two holes are at a distance of $\lambda_g/4$ where $\lambda_g \rightarrow$ Guide wavelength.

When the wave is incident - port ①, splits into two ^{waves} at two holes. They are Direct wave and Indirect wave.

Direct wave travels from the main waveguide to auxiliary waveguide through first hole directly.

Indirect wave travels from the main waveguide to auxiliary waveguide through ^{the} second hole, [when the incident signal at port ①], and reaches to the place of 1 hole in the auxiliary waveguide. by (passes) travelling the distance of $\lambda_g/2$ additionally than direct wave i.e. $[\lambda_g/4 + \lambda_g/4 = \lambda_g/2]$ signal travels two way path of $\lambda_g/4$ each. Since $\lambda_g/4$ gives 90° of phase shift.

At port ③, two waves cancelled due to the phase difference of 180° between them. So no signal is delivered from port ③.

→ Waves from port-1 reaches port-2 directly.

→ Waves from port-1 reaches port-4 either by passing through hole 1 or hole 2. Either of these signals travels same distance of $\lambda_g/4$ when we observe at port 4. So the phase difference is zero. Then those two signals added in port 4.

→ When the wave incident at port ①, waves [appears at] delivered from ports ② and ④.

Similarly when the wave incident at port ③, waves delivered from ports ① and ③. waves [direct and indirect waves] are cancelled at port ④ when power incident at ③.

So ports ① and ③ are isolated, similarly ports ② and ④ are isolated.

The degree of Coupling is determined by size and location of the holes in the waveguide walls.

The ferrite rod is tapered at both ends to reduce the attenuation and also for smooth rotation of the polarized wave.

→ When a wave enters port (1), its plane of polarization rotates by 90° because of the twist in the waveguide. It again undergoes Faraday rotation through 90° because of ferrite rod and the wave which comes out of port (2) will have a phase shift of 180° compared to the wave entering port (1).

But when the same wave (TE_{10} mode signal) enters port (2) it undergoes Faraday rotation through 90° in the same anti-clockwise direction. Because of the twist, this wave gets rotated back by 90° coming out of port (1) with 0° phase shift. Hence a wave at port (1) undergoes a phase shift of π radians (or 180°) but a wave fed from port (2) does not change its phase in a Gyrotator.

Isolator : —

An isolator is a 2-port device which

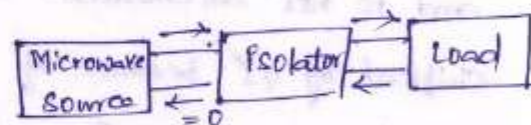
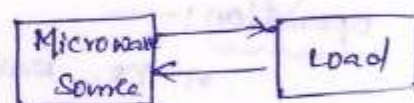
- provides very small amount of attenuation for transmission from port (1) to port (2) but provides maximum attenuation for

transmission from port (2) to port (1). This requirement is very much desirable when we want to match a source with a variable load.

In most microwave generators, the output amplitude and frequency tend to fluctuate very significantly with changes in load impedance. This is due to mismatch of generator output to the load resulting in reflected wave from load. But these reflected waves should not be allowed to reach the microwave generator, which will cause amplitude and frequency

instability of the microwave generator.

When an isolator is inserted between generator and load, the generator is coupled to the load with zero attenuation and reflections if any from the load side are completely absorbed by the isolator without affecting the



Faraday Rotation principle:

Plane of polarization of a wave rotate with distance is called "Faraday Rotation".

It states that "If a circularly polarized wave (TE_{11} in a cylindrical waveguide) is made to pass through a ferrite rod, which has been influenced by an axial magnetic field B , the axis of polarization gets tilted in clockwise direction and the amount of tilt depends upon the strength of magnetic field and geometry of the ferrite."

→ Microwave components which make use of this phenomenon are

*→ Gyrotator *→ Isolator *→ Circulator

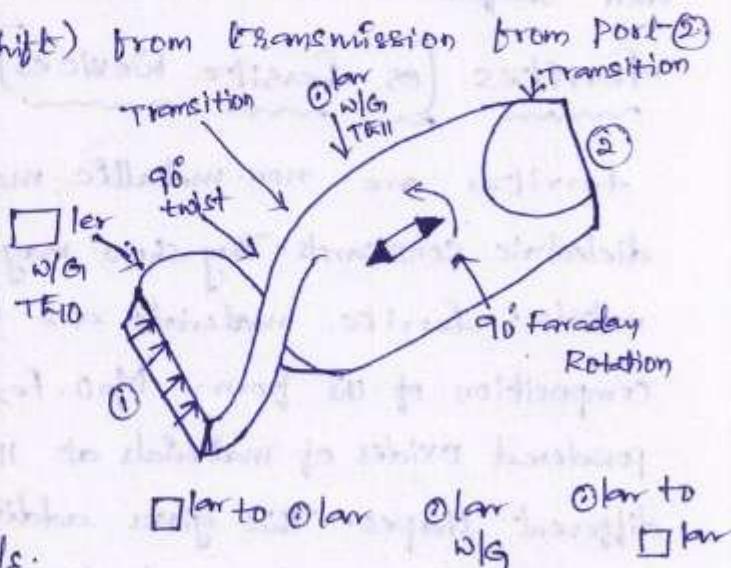
Gyrotator:

It is a two port device that has a relative phase difference of 180° for transmission from port ① to port ② and no phase shift (0° phase shift) for transmission from port ② to port ①.



operation:

The construction of a Gyrotator is shown in figure. It consists of transitions to a standard rectangular waveguide with dominant mode (TE_{10}) at both ends.



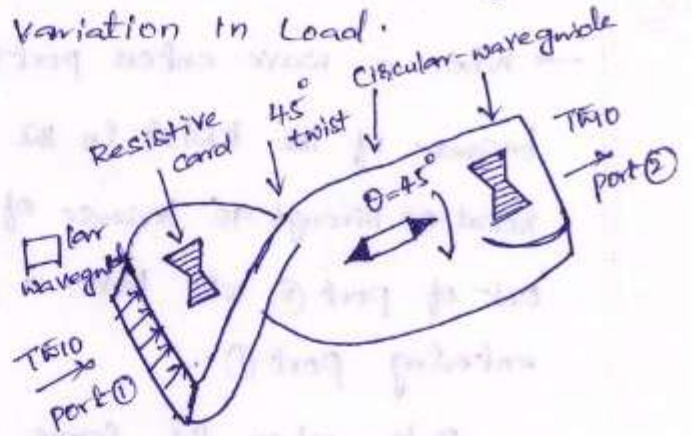
A thin circular ferrite rod tapered at both ends is located inside the circular waveguide supported by polyfoam and the waveguide is surrounded by a permanent magnet which generates DC magnetic field for proper operation of ferrite.

At the input end a 90° twisted rectangular waveguide is connected.

generator output. Hence the generator appears to be matched for all loads in the presence of isolator so that there is no change in frequency and output power due to variation in load.

Construction:—

The construction of isolator is similar to Gyrator except that an isolator makes use of 45° twisted rectangular waveguide (instead of 90° twist) and 45° Faraday rotation ferrite rod (instead of 90° in Gyrator), a resistive card is placed along the larger dimension of the rectangular waveguide, so as to absorb any wave whose plane of polarization is parallel to the plane of resistive card. The resistive card does not absorb any wave whose plane of polarization is perpendicular to its own plane.



operation:—

A TE_{10} wave passing from port ① through the resistive card and is not attenuated. After coming out of the card, the wave gets shifted by 45° because of the twist in anticlockwise direction and then by another 45° in clockwise direction because of the ferrite rod and hence comes out of port ② with the same polarization as at port ① without any attenuation.

But a TE_{10} wave fed from port ② gets a pass from the resistive card placed near port ② since the plane of polarization of the wave is perpendicular to the plane of the resistive card. Then the wave gets rotated by 45° due to Faraday rotation in clockwise direction and further gets rotated by 45° in clockwise direction due to the twist in the waveguide.

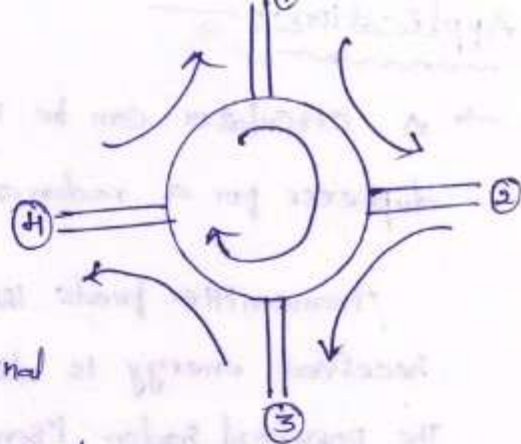
Now the plane of polarization of the wave will be parallel with that of the resistive card and hence the wave will be completely absorbed by the resistive card and the output at port ① will be zero.

Circulator:

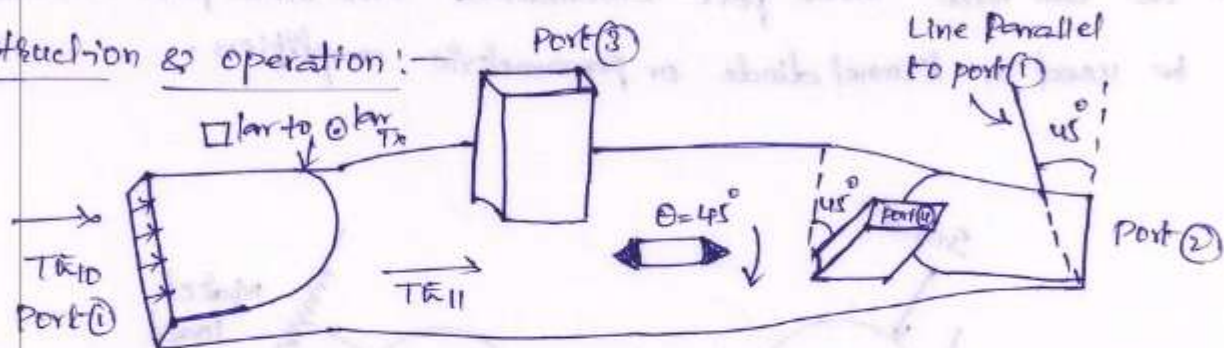
A Circulator is a four-port microwave device which has a peculiar property that each terminal is connected only to the next clockwise terminal i.e. port-① is connected to port-② only and

not to port-③ and ④ and port-② is connected only to port-③ etc.

There is no restriction on the number of ports. four port circulator is most commonly used. They are useful in parametric amplifiers, tunnel diode amplifiers and duplexers in radars.

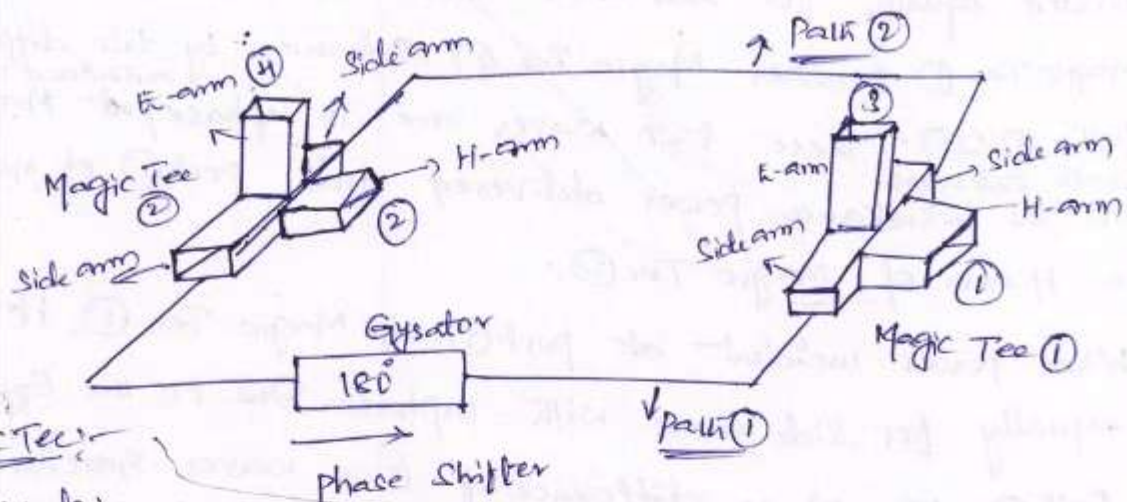


Construction & operation:-



A four port faraday rotation circulator is shown. The power entering Port ① is TE_{10} mode and is converted to TE_{11} mode because of gradual rectangular to circular transition. This power passes port ③ unaffected since the electric field is not significantly cut and is rotated through 45° due to the ferrite, passes port ④ unaffected (for the same reason as it passes port ③) and finally emerges out of port ②. Power from port ② will have plane of polarization already tilted by 45° with respect to port ①. This power passes port ④ unaffected because again the electric field is not significantly cut. This wave gets rotated by another 45° due to ferrite rod in the clockwise direction. This power whose plane of polarization is tilted through 90° finds port ③ suitably aligned and emerges out of it. Similarly port ② is coupled only to port ④ and port ④ to port ①.

Circulator construction using Magic Tees:



Magic Tee:

S-matrix of a Magic Tee

$$[S] = \frac{1}{\sqrt{2}} \begin{bmatrix} 0 & 0 & 1 & 1 \\ 0 & 0 & 1 & -1 \\ 1 & 1 & 0 & 0 \\ 1 & -1 & 0 & 0 \end{bmatrix}$$

→ When power is incident into port 3, power delivered from ports 1 and 2, in phase.

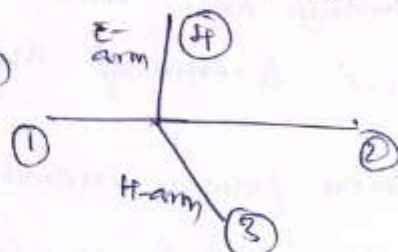
→ When power incident into port 4, power delivered from ports 1 and 2, but they are out of phase.

**

When power incident into ports 1 and 2 with same amplitude and phase, then power delivered from sum port i.e. H-port or port 3.

**

When power incident into ports 1 and 2 with same amplitude and opposite phase, then power delivered from difference port or E-arm or port 4.

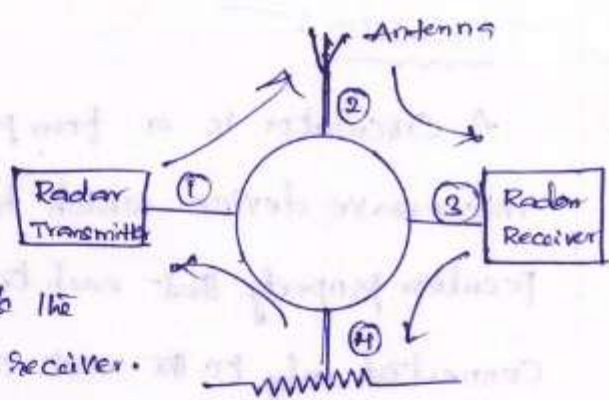


→ When power incident into port 3, power delivered through ports 1 and 2 equally with in phase.

→ When power incident into port 4, power delivered for ports 1 and 2 equally but opposite phase.

Applications:

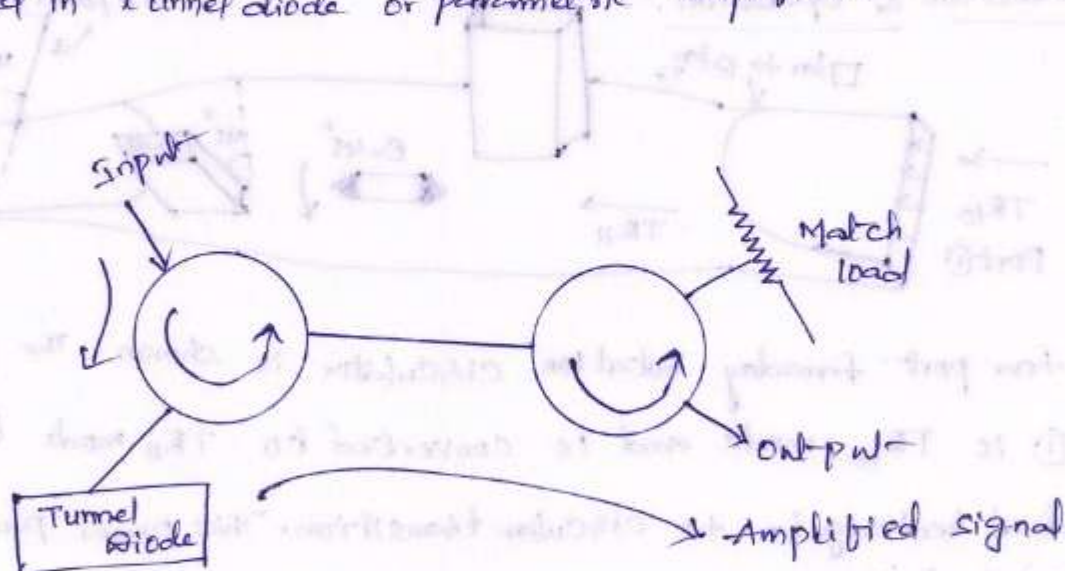
→ A circulator can be used as a duplexer for a radar antenna system.



Transmitter feeds the antenna while the received energy is directed to the receiver.

The powerful radar transmitter is isolated from the sensitive receiver and also the same antenna can be used for both transmission and reception. This is the duplexing action being formed by a circulator.

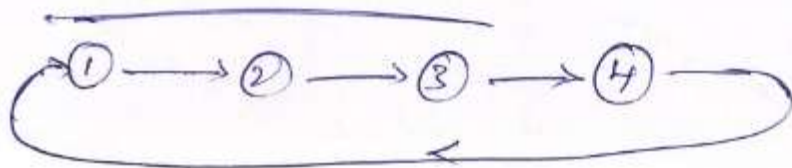
→ We can have three port circulators. Two three port circulators can be used in tunnel diode or parametric amplifiers.



→ Circulators can be used as low power devices as they can handle low powers only.

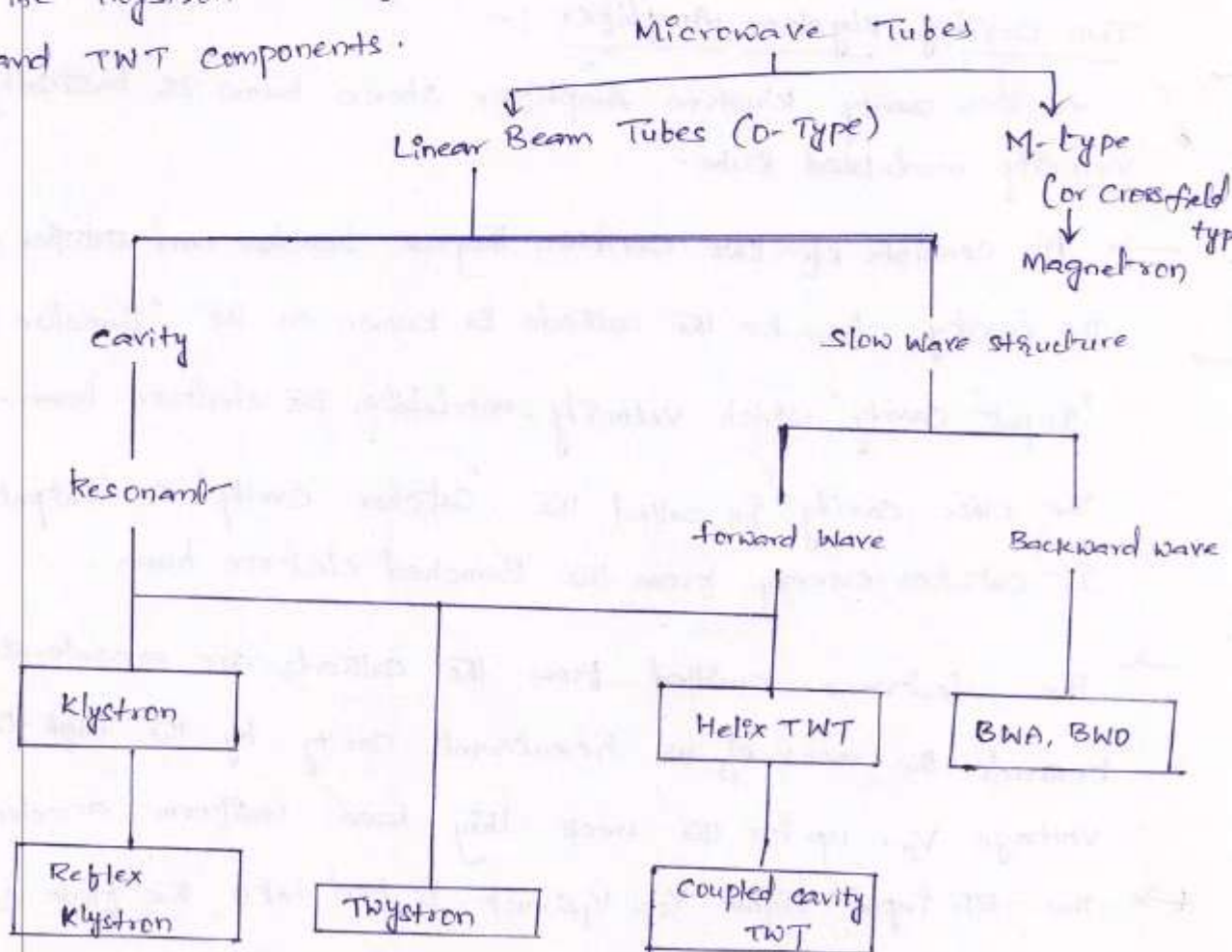
operation:-

- When power incident - into Magic Tee ① through port-①, it is divided equally for side arms. These waves which are divided in Magic Tee ① reaches Magic Tee ② side arms by two different paths ① & ②. These two waves are in inphase ^{and entered into} Magic Tee ② ^{through side arms}. So the resultant power delivered from port-② of Magic Tee ② i.e. H-arm of Magic Tee ②.
- When power incident - at port-② of Magic Tee ②, it is divided equally for side arms with inphase. Due to the gyrator through path ①, the phase difference of two waves reaches Magic Tee ① are ^{same amplitude but} out of phase by 180° each other. Then the resultant power delivered from port-③. [i.e. E-arm of Magic Tee ①].
- When power incident - at port-③ i.e. E-arm of Magic Tee ①, it is divided equally for side arms but out of phase. These two signals reaches Magic Tee ② with same amplitude with out of phase each other. When these signals enter into Magic Tee ② through side arms resultant power delivered from port-④ i.e. E-arm of Magic Tee ②.
- When power incident - at port-④ i.e. E-arm of Magic Tee ②, it is divided equally for side arms but out of phase each other. These two waves passes through two different paths and reaches Magic Tee ①. When a wave passes through path ① it provides 180° of phase shift due to the gyrator. So these two waves are ⁱⁿ same amplitude and phase when they are entered into Magic Tee ①. Due to the same amplitude and phase, resultant power delivered from port-①.



Micro Wave Tubes

- The basic principle of operation of the microwave tubes involves transfer of power from a source of DC voltage to a source of AC voltage by means of current density modulation.
- At low microwave frequencies, conventional tubes like triodes, tetrodes and pentodes are used as signal sources of low output power.
- output power can be increased by using Microwave tubes at Microwave frequencies.
- The most important microwave tubes at present are the linear-beam tubes (O-type - original type). They are two cavity klystron, Reflex klystron, Helix Travelling Wave Tube (TWT), coupled cavity TWT, forward wave Amplifier (FWA), Backward wave Amplifier and Oscillator (BWA or BWO).
- BWA or BWO have nonresonant periodic structure for electron interactions.
- The Thyratron is a hybrid amplifier that uses combinations of klystron and TWT components.



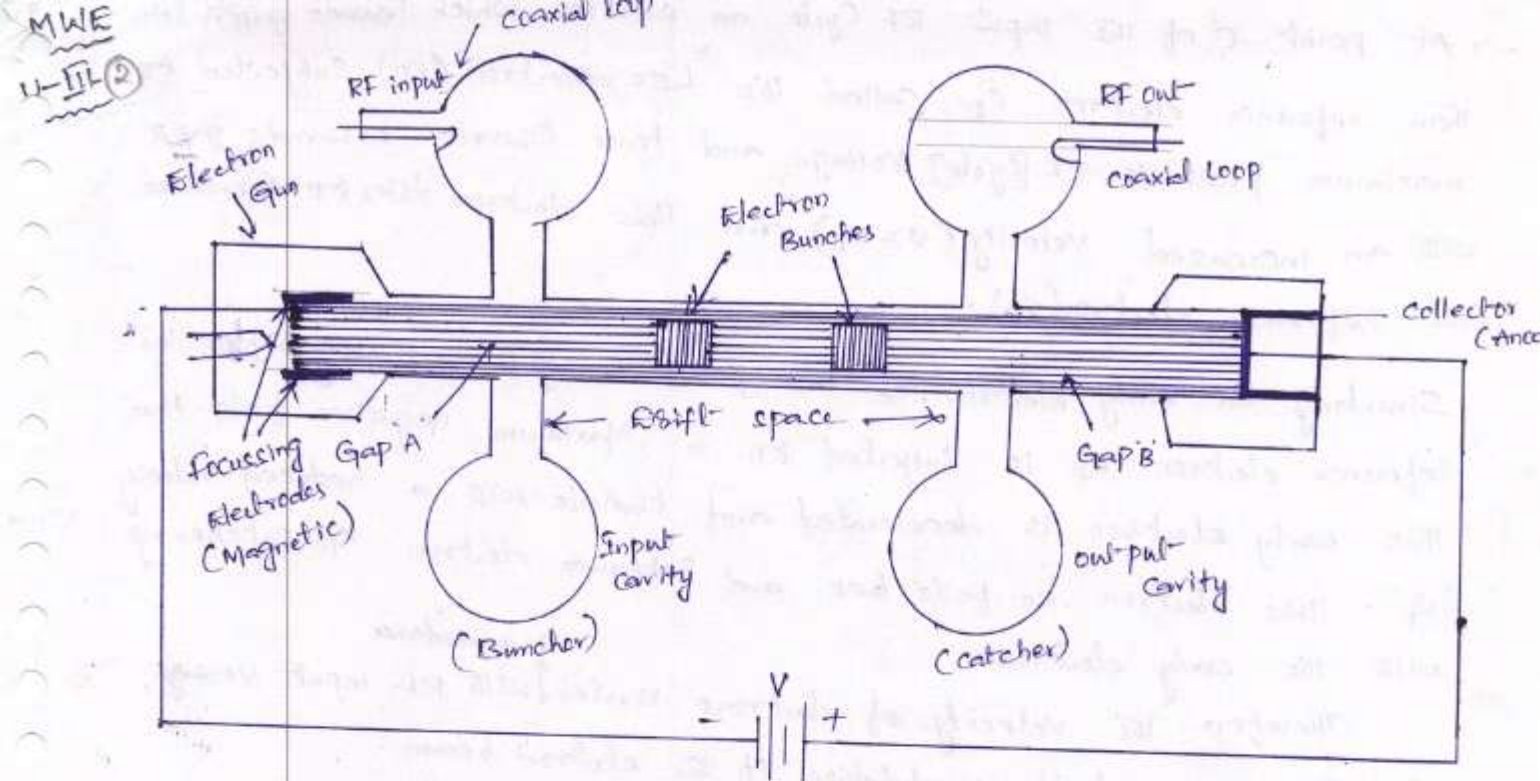
Radiation Losses:

Wherever the dimensions of the wire approaches the wavelength it will emit radiation i.e. radiation losses increase with increase in frequency.

By proper shielding of the tubes and its circuitry, we can mini-

-mize these losses.





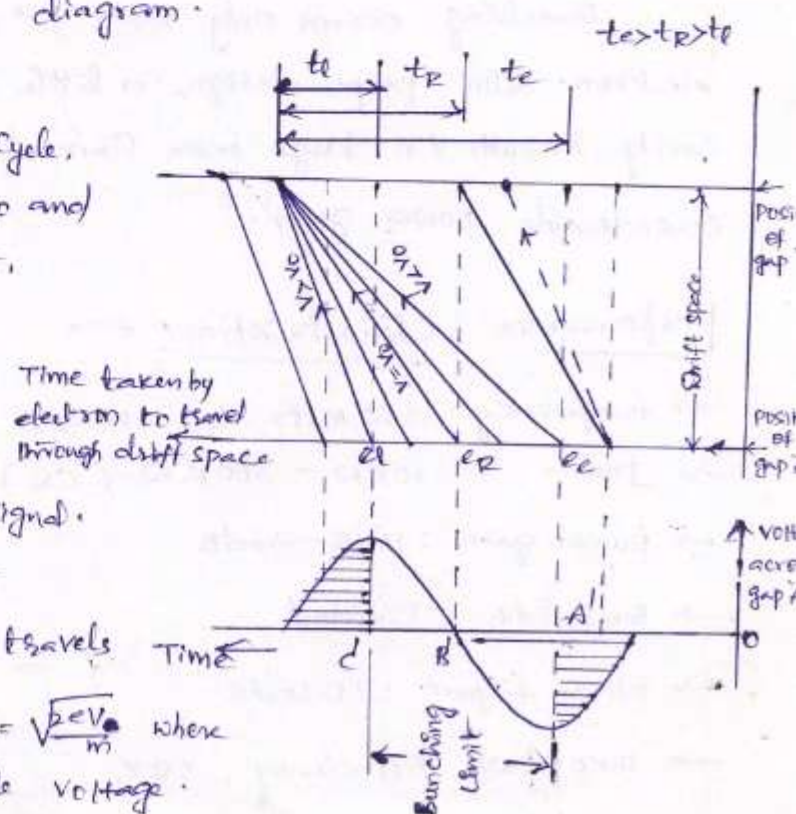
→ The Anode is kept at a positive potential with respect to cathode.
operation:-

The RF signal to be amplified is used for exciting the input bunching cavity thereby developing an alternating voltage of signal frequency across the gap A.

Let us now consider the effect of this gap voltage on the electron beam passing through gap A. The situation is best explained by means of an Apple gate diagram.

At point B' on the input RF cycle, the alternating voltage is zero and going positive. At this instant, the electric field across gap A is zero and an electron which passes through gap A at this instant is unaffected by the RF signal.

Let this electron be called the "Reference electron" (e_R), which travels with an unchanged velocity $v_0 = \sqrt{\frac{2eV_0}{m}}$ where V_0 is the anode to cathode voltage.



Transit Time Effect:-

Transit time is the time taken for the electron to travel from cathode to anode. i.e.,

$$\text{Transit time} = \tau = \frac{d}{v_0}$$

where $d \rightarrow$ distance between anode and cathode
 $v_0 \rightarrow$ velocity of electrons

Static energy of electrons = eV , $V \rightarrow$ voltage

$$\text{Kinetic energy} = \frac{1}{2} m v_0^2$$

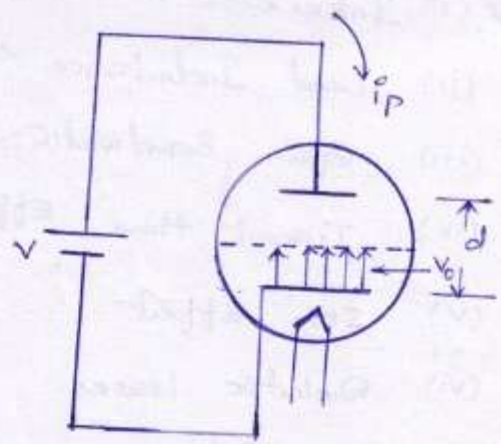
Under equilibrium, static energy = kinetic energy

i.e.,

$$eV = \frac{1}{2} m v_0^2$$

$$v_0 = \sqrt{\frac{2eV}{m}}$$

$$\tau = \frac{d}{\sqrt{\frac{2eV}{m}}}$$



At low frequencies, transit time is negligible compared to the period of the signal.

At high frequencies, transit time is comparable with the period of the signal which is very small - nano seconds.

To minimise transit time (τ), the separation between electrodes can be decreased and the plate to cathode potential ' V ' can be increased.

Gain Bandwidth Limitation:-

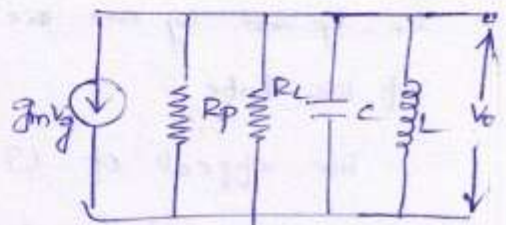
Maximum gain is achieved when the tuned circuit is at resonance.

The Maximum gain at resonance is

$$A_{\max} = g_m / G$$

Gain Bandwidth product = $A_{\max} \cdot BW$

$$= \frac{g_m}{G} \cdot \frac{G}{C} = g_m / C$$



The Gain Bandwidth product is independent of frequency.

Applications:

→ ① As power output tubes

— (a) in VHF TV Transmitters

(b) in troposphere scatter Transmitters

(c) satellite communication ground stations

(d) Radar Transmitters.

→ ② As a power oscillator if used as a Klystron oscillator.

Mathematical Analysis of a Klystron Amplifier:

Let the dc voltage between cathode and anode be V_0 and v_0 be the velocity of the electron, L be the drift space length and the RF input signal to be amplified by the Klystron be V_s .

Then,

$$v_0 = \sqrt{\frac{2eV_0}{m}} = 5.93 \times 10^5 \sqrt{V_0} \text{ m/sec} \rightarrow \textcircled{1}$$

$$V_s = V_1 \sin \omega t \rightarrow \textcircled{2}$$

where V_1 = amplitude of the signal and $V_1 \ll V_0$ is assumed.

the energy of the electron at the time of leaving buncher cavity is given by

$$\frac{1}{2} m v_1^2 = e(V_0 + V_1 \sin \omega t_1)$$

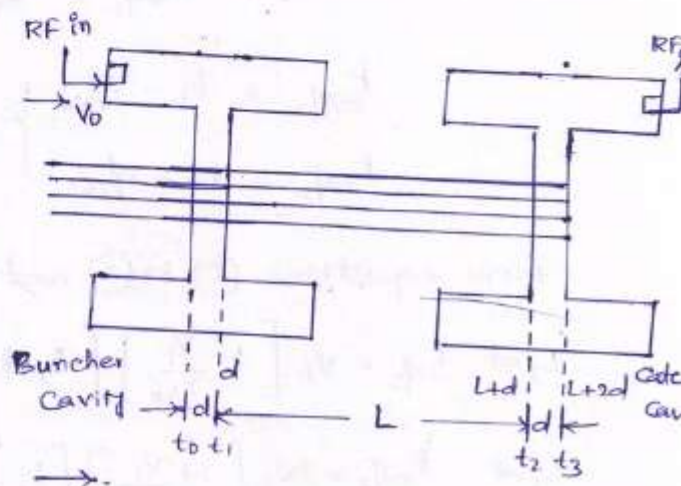
i.e.

$$v_1 = \sqrt{\frac{2e(V_0 + V_1 \sin \omega t_1)}{m}}$$

$$= \sqrt{\frac{2eV_0}{m}} \cdot \sqrt{1 + \frac{V_1 \sin \omega t_1}{V_0}}$$

Since $V_1 \ll V_0$

$$v_1 = v_0 \left[1 + \frac{V_1}{V_0} \sin \omega t_1 \right]^{1/2}$$



Expanding binomially and neglecting higher powers of $\sin \omega t_1$, we get

$$v_1 = v_0 \left[1 + \frac{V_1}{2V_0} \sin \omega t_1 \right] \rightarrow \textcircled{3}$$

This is the equation of Velocity Modulation.

Applications of TWT:

- (1) Low Noise RF amplifier in Broad Band Microwave receivers.
- (2) Repeater amplifier in wide band communication links and co-axial cables.
- (3) Due to long tube life, TWT is used as power output tube in communication satellites.
- (4) Continuous wave high power TWT's are used in troposcatter links.
- (5) Airborne and shipborne pulsed high power radars ECM ground based radars use a TWT.

Performance Characteristics of TWT:

- (1) frequency of operation: 0.5 GHz to 95 GHz.
- (2) power outputs: 5mW (10-40 GHz) (low power TWT)
250 kW (CW) at 3 GHz (High power TWT)
10 MW (pulsed) at 3 GHz.
- (3) Efficiency: 5 to 20%.
- (4) Noise figure: 4-6 dB (Low power TWT 0.5 to 16 GHz)
25 dB (High power TWT at 40 GHz).

Limitations of Conventional Tubes: — [at Microwave frequencies]

The size of electronic devices required for generation of microwave energy becomes very smaller at microwave frequencies. Because of small size, these devices increase the noise levels and result in lesser power handling capacity. So, at the microwave frequencies, the microwave tubes are used because they can provide higher output power, lesser noise, better reliability with reduced output power levels.

Due to some characteristics, conventional tubes and transistors are not used at high frequencies mentioned below.

U-III-H
If the distance has to be the same for $-\pi/2, 0, +\pi/2$ bunches, L for all three should be equal to $v_0(t_1 - t_0)$.

i.e.,
$$\frac{\pi}{2\omega} - \frac{v_1}{2v_0} (t_1 - t_0) - \frac{v_1}{2v_0} \frac{\pi}{2\omega} = 0$$

$$\begin{aligned} -(t_1 - t_0) &= \frac{\pi}{2\omega} \left[\frac{v_1}{2v_0} - 1 \right] \cdot \frac{2v_0}{v_1} \\ &= \frac{\pi}{2\omega} - \frac{\pi}{\omega} \frac{v_0}{v_1} \end{aligned}$$

As $\frac{v_0}{v_1}$ is very high, $\frac{\pi}{2\omega}$ can be neglected.

$$\therefore \boxed{t_1 - t_0 \approx \frac{\pi v_0}{\omega v_1}} \longrightarrow (13)$$

Substituting this in eqn (7), we get

$$L = v_0 \left[\frac{\pi v_0}{\omega v_1} \right] \longrightarrow (14)$$

Bunching occurs as the RF signal changes from $-\pi/2$ to $\pi/2$ i.e. π .
For a value of $\pi = 3.682$, optimum bunching occurs and

$$L_{\max} = 3.682 \left[\frac{v_0 v_0}{\omega v_1} \right] \longrightarrow (15)$$

If beam coupling coefficient of input cavity is β , given by

$$\beta = \frac{\sin(\theta_g/2)}{\theta_g/2}$$

where $\theta_g = \frac{\omega d}{v_0}$ (average gap transit angle)

Then, $L_{\max}$ is given by

$$L_{\max} = 3.682 \frac{v_0 v_0}{\omega \beta v_1} \longrightarrow (16)$$

Efficiency is given by

$$\boxed{\eta = \frac{P_{\text{out}}}{P_{\text{in}}}}$$

Output power (P_{out}):—

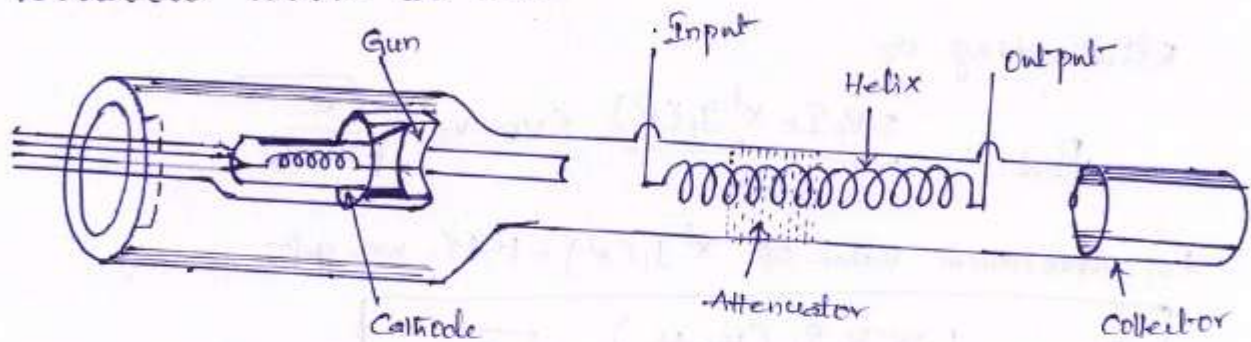
At the catcher cavity, RF voltage = $V_2 \sin \omega t_2$

Energy given by the electron to the bunch,

$$= (-e) V_2 \sin \omega t_2 = -e V_2 \sin \omega t_2$$

- The TWT makes use of a distributed interaction between an electron beam and a travelling wave. To prolong the interaction between an electron beam and the RF field it is necessary to ensure that they are both travelling in the same direction with nearly the same velocity. Thus it differs from klystron in which the electron beam travels and the RF field remains stationary.
- The electron beam travels with a velocity governed by the anode voltage. The RF field propagates with a velocity equal to velocity of light c . The interaction between the RF field and the moving electron beam will take place only when the RF field is retarded by some means. Normally slow wave structures are utilised to retard the RF field, like helix or a waveguide arrangement.

Constructional features of TWT:



It has an electron gun as used in klystrons, which is used to produce a narrow constant velocity electron beam. This electron beam is in turn passed through the centre of a long axial helix. A magnetic focussing field is provided to prevent the beam from spreading and to guide it through the center of the helix. Helix is a loosely wound thin conducting helical wire, which acts as a slow wave structure.

The signal to be amplified is applied to the end of the helix adjacent to the electron gun. The amplified signal appears at the output or other end of the helix under appropriate operating conditions.

Operation:

When the applied RF signal propagates around the turns of the helix, it produces an electric field at the centre of the helix. The RF field propagates with velocity of light. The axial electric field due to RF signal travels with velocity of light multiplied by the ratio of helix pitch to helix circumference. When the velocity of the electron beam travelling through the helix approximates the rate of advance of the axial field, then

Energy transferred = $-I_0 V_2 J_1(x) \sin \theta_0$

Maximum value of $J_1(x) = 0.58$ for $x = 1.84$.

for Maximum Energy transfer,

$$P_{\max} = -I_0 V_2 (0.58) \sin \theta_0$$

$$\sin \theta_0 = -1, \theta_0 = 2\pi - \pi/2$$

$$\therefore \text{The out-put-power, } P_{\text{out}} = P_{\max} = 0.58 I_0 V_2$$

Input power (P_{in}): The input-power is basically the dc input given by

$$P_{\text{in}} = I_0 V_0$$

The efficiency (η) is given by

$$\eta = \frac{P_{\text{out}}}{P_{\text{in}}} = \frac{0.58 I_0 V_2}{I_0 V_0} = 0.58 \frac{V_2}{V_0}$$

As V_2 is always less than V_0 , the maximum efficiency that can be attained is 0.58 or 58%. further efficiency is a function of the transit angle and is maximum for a catcher gap transit angle of $\theta_g = \theta_0 = -\pi/2$ and is zero for $\theta_g = 2\pi$.

→ for Maximum power transfer, the electron gun anode voltage required is given by $(V_1/V_0)_{\max}$. we know

$$x = \frac{V_1}{2V_0} \theta_0$$

for Maximum transfer of energy, $x = 1.84, \theta_0 = 2\pi - \pi/2$

$$\left[\frac{V_1}{V_0} \right]_{\max} = \frac{2 \times 1.84}{2\pi - \pi/2} = \frac{3.68}{[2\pi - \pi/2]}$$

Efficiency of Reflex klystron:

Maximum power is transferred to the output when the returning bunched electrons arrive at the cavity and when the field is at antinode power output.

DC power supplied by beam voltage V_0 is

$$P_{dc} = V_0 I_0 \longrightarrow (19)$$

AC power delivered is

$$P_{ac} = I_0 V_2 J_1(x') \sin \theta_0'$$

As the current flows in the negative direction, the negative sign becomes positive and $\sin \theta_0'$ is 1 and V_2 is V_1 being single and same cavity.

$$P_{ac} = I_0 V_1 J_1(x') \longrightarrow (20)$$

where $x' = \frac{V_1}{2V_0} \theta_0' \quad [\because \theta_0' = 2n\pi - \pi/2]$

$$\therefore \frac{V_1}{V_0} = \frac{2x'}{2n\pi - \pi/2}$$

From eqn (20),

$$P_{ac} = \frac{2V_0 I_0 x' J_1(x')}{(2n\pi - \pi/2)} \longrightarrow (21)$$

$$\therefore \text{Efficiency} = \frac{P_{ac}}{P_{dc}} = \frac{2V_0 I_0 x' [J_1(x')]}{V_0 I_0 (2n\pi - \pi/2)}$$

$$\boxed{\eta = \frac{2x' J_1(x')}{(2n\pi - \pi/2)}} \longrightarrow (22)$$

The factor $x' J_1(x')$ reaches a maximum value of 1.252 at $x' = 2.408$ and $J_1(x') = 0.52$. Maximum power output is obtained when $n=2$.

$\therefore$ Maximum Theoretical Efficiency is

$$\eta_{\text{Theor(max)}} = \frac{2(2.408)(0.52)}{2\pi(2) - \pi/2} = 22.78\%$$

Practical values are around 20%.

Reflex Klystron : —

The reflex klystron is a single cavity variable frequency microwave generator of low power and low efficiency.

Construction : —

It consists of an electron gun, a filament surrounded by cathode and a focussing electrode at cathode potential as shown.

The electron beam is accelerated towards the anode cavity (at positive potential). After

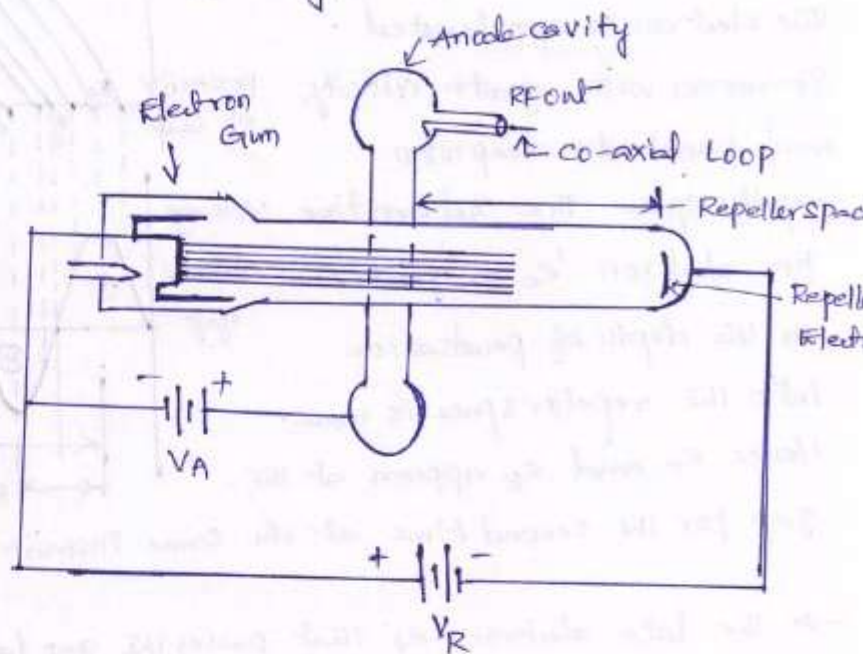
passing the gap in the cavity,

electrons travel towards a repeller electrode which is at a high negative potential V_R . The electrons never reach the repeller because of the negative field and are returned back towards the gap. Under suitable conditions the electrons give more energy to the gap than they took from the gap on their forward journey and oscillations are sustained.

operation : —

It is assumed that the oscillations are set up in the tube initially due to noise or switching transients and these oscillations are sustained by device operation. This can be explained well by Applegate diagram.

→ The RF voltage that is produced across the gap by the cavity oscillation act on the electron beam to cause velocity modulation. 'e_R' is the reference electron that passes through the gap when the gap voltage is 0 and going negative. Electron 'e_R' is unaffected by the gap voltage. This moves towards the repeller and gets reflected by the negative voltage on the repeller. It returns and passes through the gap for a second time.



$$v_1 = v_0 \left(1 + \frac{v_1}{v_0} \sin \omega t \right)^{1/2}$$

Since $v_1 \ll v_0$

$$\therefore v_1 = v_0 \left[1 + \frac{v_1}{2v_0} \sin \omega t \right]$$

$$\therefore \omega t_2 = \omega t_1 - \frac{2ms\omega}{e(V_R - v_0)} v_0 \left[1 + \frac{v_1}{2v_0} \sin \omega t \right] \rightarrow (11)$$

Let

$$-\frac{2ms\omega v_0}{e(V_R - v_0)} = \omega T_0' X = \theta_0' \rightarrow (12)$$

where $\theta_0' \rightarrow$ Round trip dc transit angle of centre of bunch electron.

$$\text{Let } \frac{v_1}{2v_0} \theta_0' = X' \rightarrow (13)$$

where $X' \rightarrow$ Bunching parameter.

$$\therefore \boxed{\omega t_2 = \omega t_1 + \theta_0' \left[1 + \frac{v_1}{2v_0} \sin \omega t \right]} \rightarrow (14)$$

Relation between Repeller Voltage and Accelerating Voltage:

considering equation (14) and when $v_1 \ll v_0$ (i.e. for the centre bunch of electrons which is unaffected by RF).

$$\omega t_2 = \omega t_1 + \theta_0'$$

Also, for maximum transfer of energy, the modes are $1\frac{3}{4}$ cycles apart

$$2\pi(n - \frac{1}{4}) \text{ where } n - \frac{1}{4} = \frac{3}{4}, 1\frac{3}{4} \text{ etc.}$$

$$\text{or } 2\pi(n - \frac{1}{4}) = 2\pi n - \pi/2$$

The optimum value of θ_0' is

$$\theta_0' = (2\pi n - \pi/2) \rightarrow (15)$$

from eqn (12),

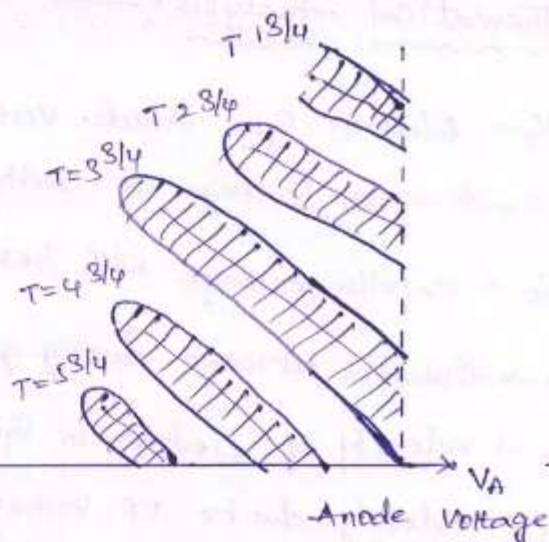
$$\theta_0' = -\frac{2ms\omega}{e(V_R - v_0)} v_0 \text{ or } v_0 = -\frac{e(V_R - v_0)}{2ms\omega} \cdot \theta_0'$$

$$v_0^2 = \frac{e^2(V_R - v_0)^2(\theta_0')^2}{4\omega^2 m^2 s^2} \rightarrow (16)$$

Since we know that $\frac{1}{2} m v_0^2 = eV_0$

$$\text{or } V_0 = \frac{m}{2e} v_0^2$$

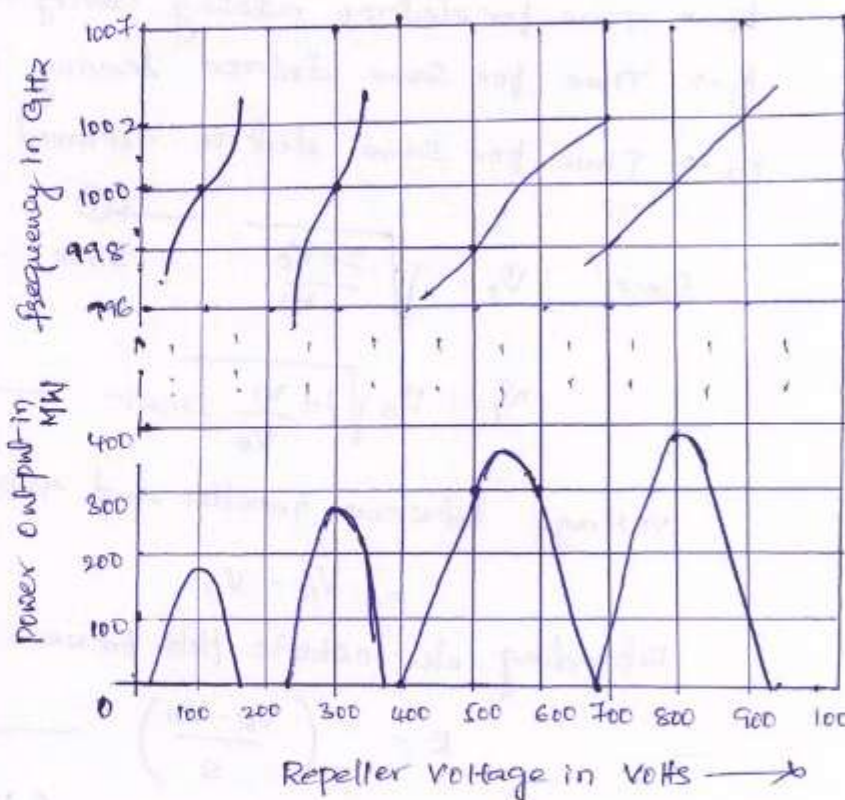
↑
VR
Repeller
Voltage



(2) power output and frequency characteristics:-

The frequency of the resonance of the cavity decides the frequency of oscillation.

→ Variations of repeller voltage slightly changes the frequency.



Applications:-

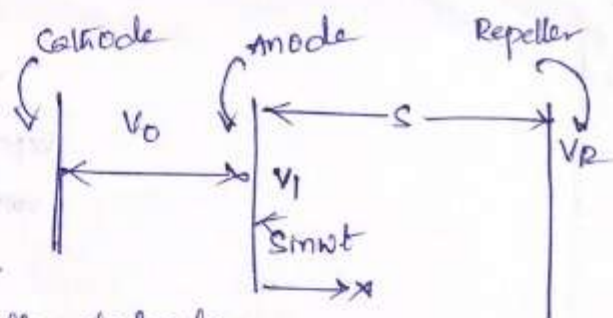
- In Radar Receivers
- Local oscillator in Microwave Receivers.
- Signal Source in Microwave generator of variable frequency.
- portable microwave links
- pump oscillator in parametric Amplifier.

Mathematical Analysis

V_0 = Electron Gun Anode Voltage

$V_1 \sin \omega t$ = RF Voltage at cavity gap

V_R = Repeller Voltage with respect to Cathode.



s → distance between cavity gap and repeller electrode.

v_0 → velocity of electron in Gun

v_1 → velocity due to RF voltage in addition to the electron accelerating voltage V_0 .

t_0 → Time for electron entering cavity gap at $x=0$

t_1 → Time for same electron leaving cavity gap at $x=d$

t_2 → Time for same electron returned by retarding field at $x=d$.

Since $v_0 = \sqrt{\frac{2eV_0}{m}}$ → (1) Since $\frac{1}{2}mv_0^2 = eV_0$

$v_1 = v_0 \sqrt{1 + \frac{V_1}{V_0} \sin \omega t}$ → (2)

Voltage between repeller and anode is

$$= V_R - V_0$$

Retarding electrostatic field between repeller and anode is given by

$E = -\left(\frac{V_R - V_0}{s}\right)$ → (3) $(E = \frac{V}{r})$ Since

force on electron $= -eE = +e\left(\frac{V_R - V_0}{s}\right)$ → (4)

Also, force on electrons = mass $\times$ acceleration $= m \frac{d^2x}{dt^2}$

$\therefore m \frac{d^2x}{dt^2} = \frac{e}{s} (V_R - V_0)$

$\Rightarrow \frac{d^2x}{dt^2} = \frac{e}{ms} (V_R - V_0)$ → (5)

By integrating the above equation,

$\frac{dx}{dt} = \frac{e}{ms} (V_R - V_0)t + c$ → (6)

At $t = t_1$, $\frac{dx}{dt} = v_1$

$$v_1 = \frac{e}{ms} (V_R - V_0) t_1 + c$$

$$\Rightarrow c = v_1 - \frac{e}{ms} (V_R - V_0) t_1$$

$$\therefore \frac{d\eta}{dt} = \frac{e}{ms} (V_R - V_0) t + v_1 - \frac{e}{ms} (V_R - V_0) t_1$$

$$\frac{d\eta}{dt} = \frac{e}{ms} (V_R - V_0) (t - t_1) + v_1 \quad \rightarrow (7)$$

By Integrating the above equation

$$\eta = \frac{e}{2ms} (V_R - V_0) (t - t_1)^2 + v_1 t + c_1 \quad \rightarrow (8)$$

At $\eta = 0$, i.e. at the point of return from repeller space, $t = t_2$

$$0 = \frac{e}{2ms} (V_R - V_0) (t_2 - t_1)^2 + v_1 t_2 + c_1$$

$$\Rightarrow c_1 = -\frac{e}{2ms} (V_R - V_0) (t_2 - t_1)^2 - v_1 t_2$$

$$\therefore \eta = \frac{e}{2ms} (V_R - V_0) (t - t_1)^2 + v_1 t - \left[\frac{e}{2ms} (V_R - V_0) (t_2 - t_1)^2 + v_1 t_2 \right]$$

$$\therefore \eta = \frac{e}{2ms} (V_R - V_0) \left[(t - t_1)^2 - (t_2 - t_1)^2 \right] + v_1 (t - t_2)$$

Again when $t = t_1$, $\eta = 0$

$$\text{i.e., } -\frac{e}{2ms} (V_R - V_0) (t_2 - t_1)^2 - v_1 (t_2 - t_1) = 0$$

$(t_2 - t_1)$ is the round trip transit time and is given by

$$(t_2 - t_1) = \frac{-2ms v_1}{e (V_R - V_0)} \quad \rightarrow (9)$$

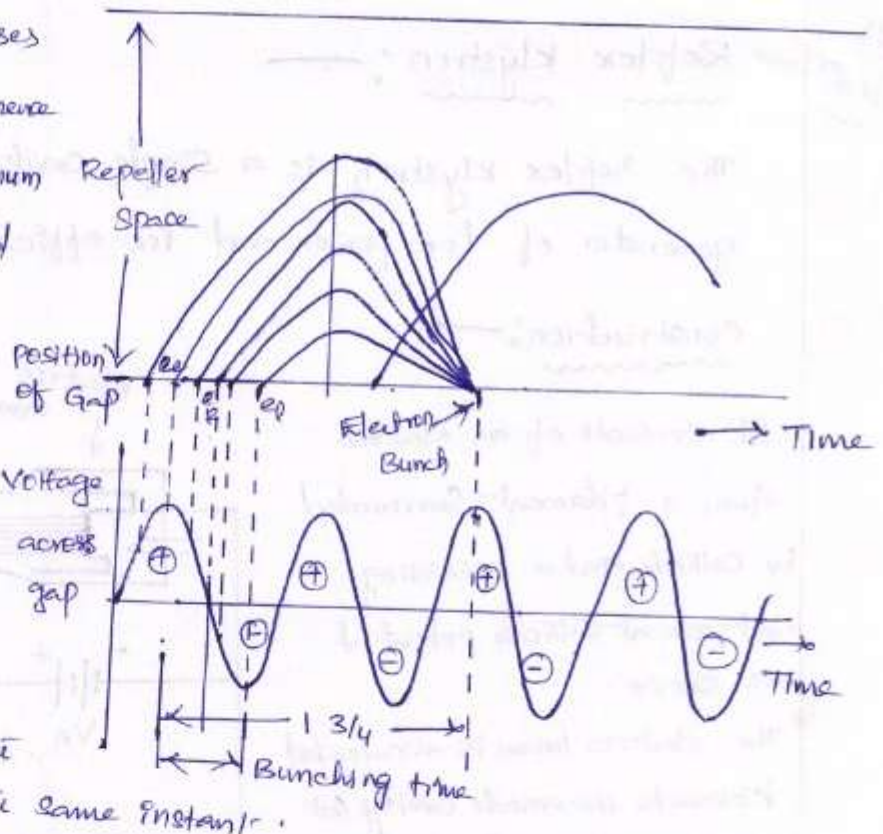
The transit angle ' ωt ' is defined as 'transit angle at time t '.

$$\omega(t_2 - t_1) = \frac{-2ms v_1 \omega}{e (V_R - V_0)}$$

$$\text{or } \omega t_2 = \omega t_1 - \frac{2ms v_1 \omega}{e (V_R - V_0)} \quad \rightarrow (10)$$

→ The early electron ' e_e ' that passes through the gap before the reference electron ' e_r ' experiences a maximum positive voltage across the gap and this electron is accelerated.

It moves with greater velocity and penetrates deep into repeller space. The return time for electron ' e_e ' is greater as the depth of penetration into the repeller space is more. Hence e_e and e_r appears at the gap for the second time at the same instant.



→ The late electron ' e_l ' that passes the gap later than reference electron ' e_r ' experiences a maximum negative voltage and moves with retarding velocity. The return time is shorter as the penetration into repeller space is less and catches up with e_r and e_e electrons forming a bunch.

Bunches occur once per cycle centred around the reference electron ' e_r ' and these bunches transfer maximum energy to the gap to get sustained oscillations.

→ For oscillations to be sustained, the time taken by the electrons to travel into the repeller space and back to the gap (called transit time) must have an optimum value.

The optimum transit time should be

$$T = n + \frac{3}{4} \quad \text{where } n \rightarrow \text{integer}$$

This depends on repeller and anode voltages.

Operating characteristics:

① Voltage characteristics:

Oscillations can be obtained only for specific combinations of anode and repeller voltages that give a favourable transit time

→ The earlier the mode the larger the output power.

$$(T = n + \frac{3}{4})$$

Substituting for V_0^2 from eqn (16)

$$V_0 = \frac{m}{2e} \cdot \frac{e^2 (V_R - V_0)^2}{4\omega^2 m^2 s^2} (\theta_0')^2$$

where $\theta_0' = 2\pi n - \pi/2$ [from eqn (15)]

$$\therefore \frac{V_0}{(V_R - V_0)^2} = \frac{m}{2e} \cdot \frac{e^2}{4\omega^2 m^2 s^2} (2\pi n - \pi/2)^2$$

Simplifying,

$$\boxed{\frac{V_0}{(V_R - V_0)^2} = \frac{1}{8} \cdot \frac{1}{\omega^2 s^2} \cdot \frac{e}{m} (2\pi n - \pi/2)^2} \rightarrow (17)$$

Expression for change in frequency due to repeller voltage:—

$$(V_R - V_0)^2 = \frac{8ms^2 V_0}{[2\pi n - \pi/2]^2 \cdot e} \omega^2$$

Differentiating V_R in the above equation with respect to V_0 then

$$2(V_R - V_0) \frac{dV_R}{d\omega} = \frac{16ms^2 V_0}{(2\pi n - \pi/2)^2 \cdot e} \omega$$

$$\text{or } \frac{dV_R}{d\omega} = \frac{8ms^2 V_0 \omega}{e [2\pi n - \pi/2]^2 (V_R - V_0)}$$

$$\therefore V_R - V_0 = \sqrt{\frac{8ms^2 V_0 \omega^2}{e [2\pi n - \pi/2]^2}}$$

$$\therefore \frac{dV_R}{d\omega} = \frac{8ms^2 V_0 \omega}{e [2\pi n - \pi/2]^2} \cdot \sqrt{\frac{e [2\pi n - \pi/2]^2}{8ms^2 V_0 \omega^2}}$$

$$\frac{dV_R}{d\omega} = \frac{s}{(2\pi n - \pi/2)} \sqrt{\frac{8mV_0}{e}}$$

Since $\omega = 2\pi f$

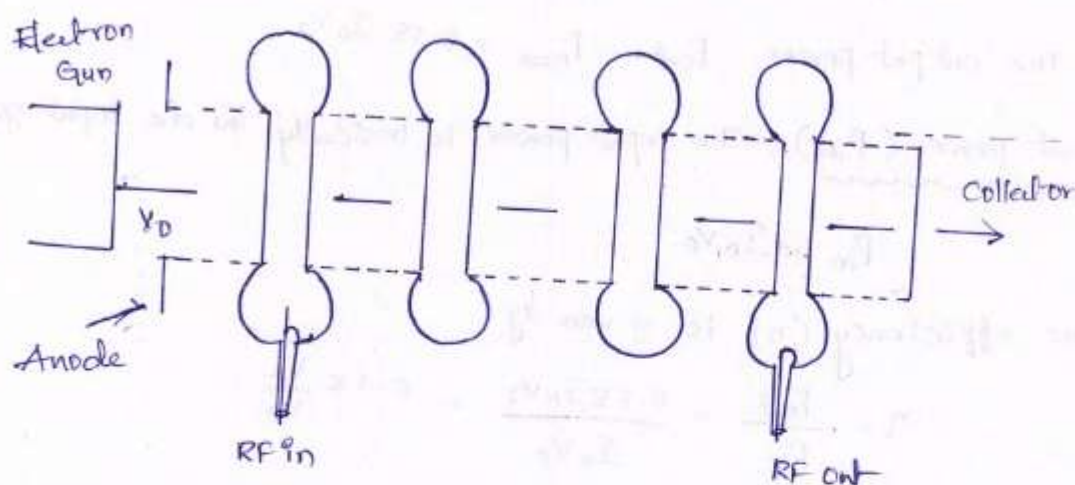
$$\boxed{\frac{dV_R}{df} = \frac{2\pi s}{(2\pi n - \pi/2)} \sqrt{\frac{8mV_0}{e}}} \rightarrow (18)$$

→ Variation in frequency is quite sensitive to repeller voltage variations.

Multi Cavity Klystron:-

→ Gain of 10 to 20 dB achieved practically with two cavity klystron.

Higher overall gain can be achieved by connecting several two cavity tubes in cascade, feeding the output of each of the tubes to the input of the succeeding one.



Here each of the intermediate cavities act as a buncher with the passing electron beam inducing an enhanced RF voltage than the previous cavity.

→ The cavities can all be tuned to the same frequency (synchronous tuning) for narrow band operation.

Band width can be improved by stagger tuning of cavities with reduction in gain. This stagger tuning is employed in VHF klystrons for TV transmitter output tubes and in satellite earth station transmitters as power amplifiers.

→ Two cavity klystron can be operated as an oscillator by feeding back a part of the catcher output into the buncher in proper phase so as to satisfy Barkhausen criterion.

Power output in terms of repeller voltage V_R :-

$$P_{out} = \frac{2V_0 I_0 x' J_1(x')}{2n\pi - \pi/2}$$

$$2n\pi - \frac{\pi}{2} = \theta_0' = \omega T_0' = \frac{2m\omega}{e(V_R - V_0)} v_0$$

The negative sign is not taken as electron bunch travels in the reverse direction, $-x$.

$$P_{out} = \frac{2V_0 I_0 x' J_1(x')}{2m\omega v_0} * (V_R - V_0)e$$

Eliminating v_0

$$P_{out} = \frac{2V_0 I_0 x' J_1(x')}{\omega s} * (V_R - V_0) \sqrt{\frac{e}{2mV_0}}$$

For Maximum value of $x' J_1(x') = 1.75$, we get

$$P_{max} = \frac{1.25 V_0 I_0 (V_R - V_0)}{\omega s} * \sqrt{\frac{e}{2mV_0}}$$

Performance characteristics of Reflex klystron:-

- (1) Frequency Range : 4 to 200 GHz
- (2) Output power : 1.0 mW to 2.5 W
- (3) Theoretical Efficiency : 22.78%.
- (4) Practical Efficiency : 10% to 20%.
- (5) Tuning range : 5 GHz at 2 Watts to 30 GHz at 10 mW.

Travelling Wave Tube (TWT):-

Klystrons are essentially narrow band devices as they utilise cavity resonators to velocity modulate the electron beam over a narrow gap whereas TWTs are broad band devices in which there are no cavity resonators. The interaction space in a TWT is extended and the electron beam exchanges energy with the RF wave over the full length of the tube.

The average energy given to the RF field in a cycle

$$P_{av} = \frac{1}{2\pi} \int_{\omega t_1=0}^{\omega t_2=2\pi} (-eV_2 \sin \omega t_2) d\omega t_1 \longrightarrow (17)$$

In the field free space between cavities, the transit time for velocity modulated electron is given by

$$\begin{aligned} T = t_2 - t_1 &= \frac{L}{v_1} = \frac{L}{v_0 \left[1 + \frac{v_1}{v_0} \sin \omega t_1 \right]^{1/2}} \\ &= \frac{L}{v_0} \left[1 - \frac{v_1}{2v_0} \sin \omega t_1 \right] \longrightarrow (18) \end{aligned}$$

Multiplying by ω ,

$$\omega T = \omega(t_2 - t_1) = \frac{\omega L}{v_0} \left[1 - \frac{v_1}{2v_0} \sin \omega t_1 \right] \longrightarrow (19)$$

In the above equation, $\frac{L}{v_0} = T_0$, the transit time without RF voltage V_1 in buncher cavity and $\frac{\omega L}{v_0} = \omega T_0 = \theta_0 = 2\pi N$ is the transit angle without RF voltage V_1 in buncher cavity and N is the number of electron transit cycles in drift space.

The bunching parameter X of a klystron is defined by the equation,

$$X = \frac{V_1}{2V_0} \theta_0 \longrightarrow (20)$$

which is a dimensionless quantity and proportional to input power. Equation (17) can be written using equation (19),

$$\begin{aligned} P_{av} &= -\frac{eV_2}{2\pi} \int_0^{2\pi} \sin(\omega t_1 + T) d\omega t_1 \\ \text{i.e., } P_{av} &= -\frac{eV_2}{2\pi} \int_0^{2\pi} \sin \left[\omega t_1 + \theta_0 \left(1 - \frac{v_1}{2v_0} \sin \omega t_1 \right) \right] d\omega t_1 \end{aligned}$$

This is a Bessel function and its solution is given by

$$P_{av} = -eV_2 J_1(X) \sin \theta_0 \longrightarrow (21)$$

where $J_1(X)$ = Bessel function of the first order for the argument X .

for N electron transit cycles

$$\text{Energy transferred} = N P_{av} = -N e V_2 J_1(X) \sin \theta_0$$

$N e = I_0$, the output current

Interaction takes place between them in such a way that on an average the electron beam delivers energy to the RF wave on the helix. Thus, the signal wave grows and amplified output is obtained at the output of the TWT.

The axial phase velocity v_p is given by

$$v_p = v_e \left(\frac{\text{Pitch}}{2\pi r} \right)$$

where $r \rightarrow$ radius of the helix and is essentially constant over a range of frequencies and this characteristic of helix slow wave structure enables TWT to have broadband operation.

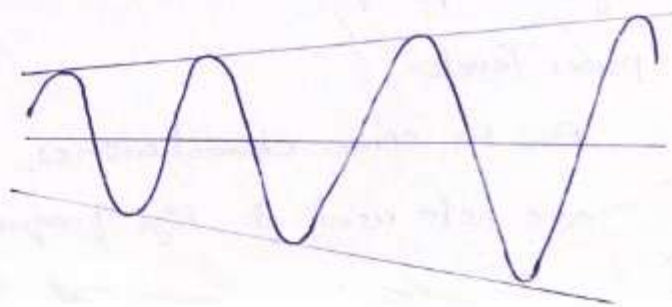
Helix provides the least change in v_p with frequency, is preferred over other slow wave structures of TWT.

$\rightarrow$ When the axial field is zero, electron velocity is unaffected. This happens at the point of node of the axial electric field.

At a point where the axial field is positive antinode, the electron coming against it is accelerated and tries to catch up with the later electron which encounters the nodal RF axial field.

At a later point where the axial RF field is negative antinode, the electrons referred before tend to overtake. The electrons get velocity modulated.

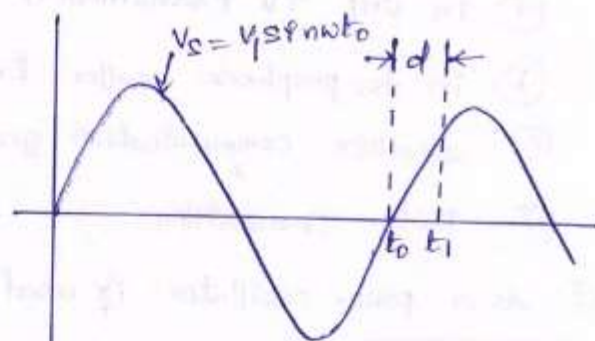
As a result of energy transfer from the electron to the RF field in phase with the RF field at the axis, a second wave is induced on the helix. This produces an axial electric field that lags behind the original electric field by $1/4$. Bunching continues to take place. The electrons in the bunch encounter retarding field and deliver energy to the wave on the helix. The output becomes larger than the input and amplification results. Thus, the energy increase in the RF is a continuous process.



$$\omega t_1 = \omega t_0 + \theta_g/2$$

where $\theta_g \rightarrow$ phase angle of the RF input voltage during which the electron is accelerated.

$$\theta_g = \omega t = \omega(t_1 - t_0) = \frac{\omega d}{v_0} \rightarrow (4)$$



Bunching process of the Electron Beam: -

Maximum velocity occurs at $+\pi/2$ so that

$$v_1(\max) = v_0 \left[1 + \frac{V_1}{2V_0} \right] \rightarrow (5)$$

and Minimum velocity at $-\pi/2$, so that

$$v_1(\min) = v_0 \left[1 - \frac{V_1}{2V_0} \right] \rightarrow (6)$$

If the distance in the drift space at which the bunching occurs from the buncher grid at time t_1 is L_1 ,

$$L_1 = v_0(t_1 - t_0) \rightarrow (7)$$

The distance L_1 at $t = -\pi/2 = v_{\min}(t_1 - t_{-\pi/2}) \rightarrow (8)$

The distance L_1 at $t = +\pi/2 = v_{\max}(t_1 - t_{+\pi/2}) \rightarrow (9)$

$$\left. \begin{aligned} t_{-\pi/2} &= t_0 - \frac{\pi}{2\omega} \\ t_{+\pi/2} &= t_0 + \frac{\pi}{2\omega} \end{aligned} \right\} \rightarrow (10)$$

from equations (5) to (7) and (10), we get

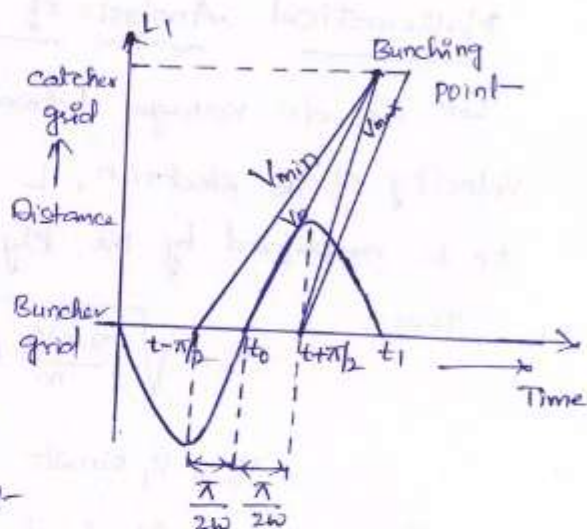
$$L_1 \text{ at } t_{-\pi/2} = v_0 \left[1 - \frac{V_1}{2V_0} \right] \left[t_1 - t_0 + \frac{\pi}{2\omega} \right] \rightarrow (11)$$

$$L_1 \text{ at } t_{+\pi/2} = v_0 \left[1 + \frac{V_1}{2V_0} \right] \left[t_1 - t_0 - \frac{\pi}{2\omega} \right] \rightarrow (12)$$

$$\therefore L_1 = v_0(t_1 - t_0) + v_0 \left[\frac{\pi}{2\omega} - \frac{V_1}{2V_0}(t_1 - t_0) - \frac{V_1}{2V_0} \frac{\pi}{2\omega} \right]$$

for Maximum value,

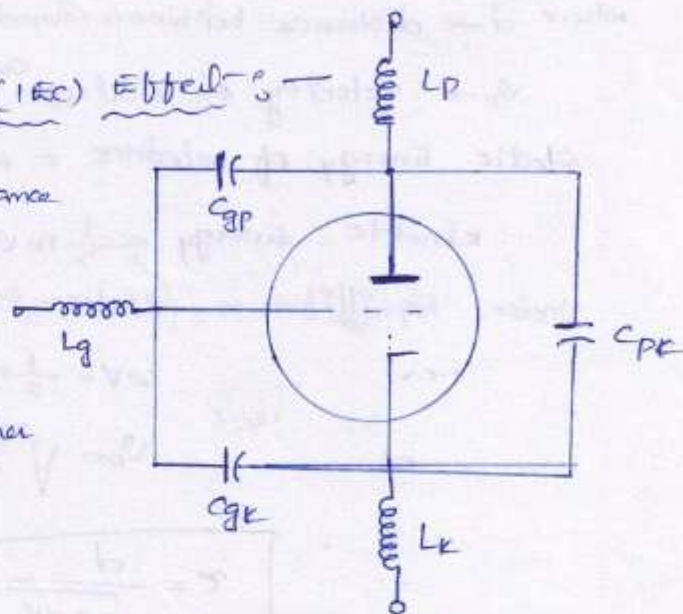
$$L_1 = v_0(t_1 - t_0) + v_0 \left[\frac{\pi}{2\omega} - \frac{V_1}{2V_0}(t_1 - t_0) - \frac{V_1}{2V_0} \frac{\pi}{2\omega} \right]$$



- U-III-12
- (i) Inter electrode capacitance
 - (ii) Lead Inductance effect
 - (iii) Gain Bandwidth Limitation
 - (iv) Transit time Effect
 - (v) Skin Effect
 - (vi) Dielectric Losses.

Inter Electrode capacitance (IEC) Effect:-

As frequency increases, the reactance $X_C = \frac{1}{2\pi f C}$, decreases and the output voltage decreases due to shunting effect. Because at higher frequencies X_C becomes almost a short. C_{gp} , C_{gk} and C_{pk} are the inter electrode capacitances, which come into effect.



The effect of IEC can be minimised by reducing the IEC's C_{gk} , C_{pk} and C_{gp} . These can be reduced by decreasing the area of the electrodes i.e. by using smaller electrodes or by increasing the distance between electrodes.

Load Inductance (LI) Effect:-

As frequency increases, the reactance $X_L = 2\pi f L$ increases and hence the voltage appearing at the active electrodes are less than the voltages at the base pins. This results in reduced gain for the tube amplifier. L_k , L_p and L_g are the lead inductances that limit the performance of the tube.

The effect of LI can be minimised by decreasing L . Since L is proportional to reactance ($L = l/\mu_0 \mu_r A$), L can be decreased by using larger sized short leads without base pins i.e. by increasing A and decreasing l . This reduces the power handling capability.

→ At point C of the input RF cycle an electron which leaves gap A later than reference electron e_R , called the "late electron" (e_L) subjected to maximum positive RF [cycle] voltage and hence travels towards gap B with an increased velocity ($v > v_0$) and this electron tries to overtake the reference electron (e_R).

Similarly an early electron e_e that passes the gap A slightly before the reference electron e_R is subjected to a maximum negative field. Hence this early electron is decelerated and travels with a reduced velocity ' v_0 '. This electron e_e falls back and reference electron e_R catches up with the early electron.

Therefore the velocity of electrons varies ^{in accordance} with RF input voltage, resulting in velocity modulation of the electron beam.

As a result of these actions, the electrons in the bunching limit (between points A' and d') gradually bunch together as they travel down the drift space, from gap A to gap B. The pulsating stream of electrons pass through gap B and excite oscillations in the output cavity (catcher). The density of electrons passing the gap B vary cyclically with time, that is the electron beam contains an ac current and is current modulated. The drift-space converts the velocity modulation into current modulation.

Bunching occurs only once per cycle centered around the reference electron. With proper design, a little RF power applied to the buncher cavity results in large beam currents at the catcher cavity with a considerable power gain.

Performance Characteristics: →

- frequency: 250 MHz to 100 GHz.
- power : 10 kW - 500 kW (CW) 30 MW (pulsed)
- power gain : 15 dB - 70 dB
- Bandwidth : Limited
- Noise figure : 15-20 dB
- Theoretical Efficiency : 58%.

In Microwave Circuits, the Gain Bandwidth Limitation can be overcome by use of

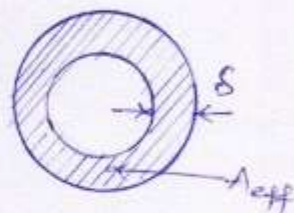
- (i) Re-entrant cavities
- (ii) Slow wave Tubes

for larger gain over a larger Bandwidth.

Effect due to RF Losses:-

(a) skin effect Losses (or Conductor or I^2R losses):-

These losses are much at higher frequencies at which the Current has the tendency to confine itself to a smaller cross section of the conductor towards its outer surface.



$$\delta = \text{skin depth} = \sqrt{2/\omega\mu\sigma}$$

$$\text{i.e., } \delta \propto \frac{1}{\sqrt{\omega}} \text{ and } \delta \propto A_{\text{eff}}$$

where A_{eff} is the effective area over which Current flows.

$$A_{\text{eff}} \propto \frac{1}{\sqrt{f}} \text{ and } R = \frac{\rho l}{A_{\text{eff}}}$$

$$R = \frac{\rho l}{\frac{1}{\sqrt{f}}} = \rho l \sqrt{f}$$

i.e., as f increases R increases.

Hence losses will increase at higher frequencies. These losses can be reduced by increasing the size of the conductors.

(b) Dielectric Losses:-

This occurs in various types of insulating materials used in the device. The loss in any of the material is generally given by

$$P = \pi f V_0^2 \epsilon_r \tan \delta$$

where ϵ_r = relative permittivity of the dielectric
 δ = loss angle of the dielectric

As f increases the power loss increases. These losses can be reduced by eliminating the tube base and to reduce the surface area of Glass.

- In a linear beam tube a magnetic field whose axis coincides with that of the electron beam is used to hold the beam together as it travels the length of the tube.
- O-type tubes abbreviated as original type tubes.
- In the Microwave interaction region the electrons are either accelerated or decelerated by the Microwave interaction region field and then bunched as they drift down the tube. The bunched electrons, in turn, induce current in the output structure.
- O-type travelling wave tubes are suitable for Amplification.

Two - cavity Klystrons:

- The two cavity Klystron is a widely used Microwave Amplifier operated by the principles of velocity and current modulation.
- All electrons injected from the cathode arrive at the first cavity with uniform velocity.

Two cavity Klystron Amplifier:

A two cavity Klystron Amplifier shown below is basically a Velocity modulated tube.

- It consists of two cavities. They are buncher and catcher cavities. The cavity close to the cathode is known as the "Buncher cavity" or "Input cavity," which velocity modulates the electron beam. The other cavity is called the "catcher cavity" or "output cavity". It catches energy from the Bunched electron beam.
- The electrons emitted from the cathode are accelerated towards the neck of the resonant cavity by the high DC Voltage V_0 . up to the neck they have uniform acceleration.
- The RF input signal $V_s = V_1 \sin \omega t$ is fed into the first cavity called "Buncher cavity".

M-Type Microwave Tubes

Microwave Tubes are two types. They are O-type or [original] and crossfield type Tubes. In original type of microwave tubes only electric fields are exists and in cross-field type tubes both electric and Magnetic fields are exists and they are perpendicular to each other. The principle in cross-field type microwave tubes is, called the M-type is the Magnetron.

Magnetrons provide Microwave oscillations of very high peak power.

In Klystrons (Two cavity klystron, Reflex klystron) the electrons carrying energy are in contact with the RF field in the resonant cavity only for a short duration, but in Magnetrons electrons move to interact with RF field for a longer duration so higher efficiency can be obtained.

Magnetrons - types:

There are three types of Magnetrons.

- (1) Negative Resistance type
- (2) Cyclotron frequency type
- (3) Travelling wave or cavity type.

Negative Resistance Type:

This type of Magnetrons make use of negative resistance between two anode segments but have low efficiency and are useful only at low frequencies ($< 500 \text{ MHz}$).

Cyclotron frequency Type: — This type of magnetrons depend upon synchronism between an alternating component of electric and periodic oscillation of electrons in a direction parallel to this field.

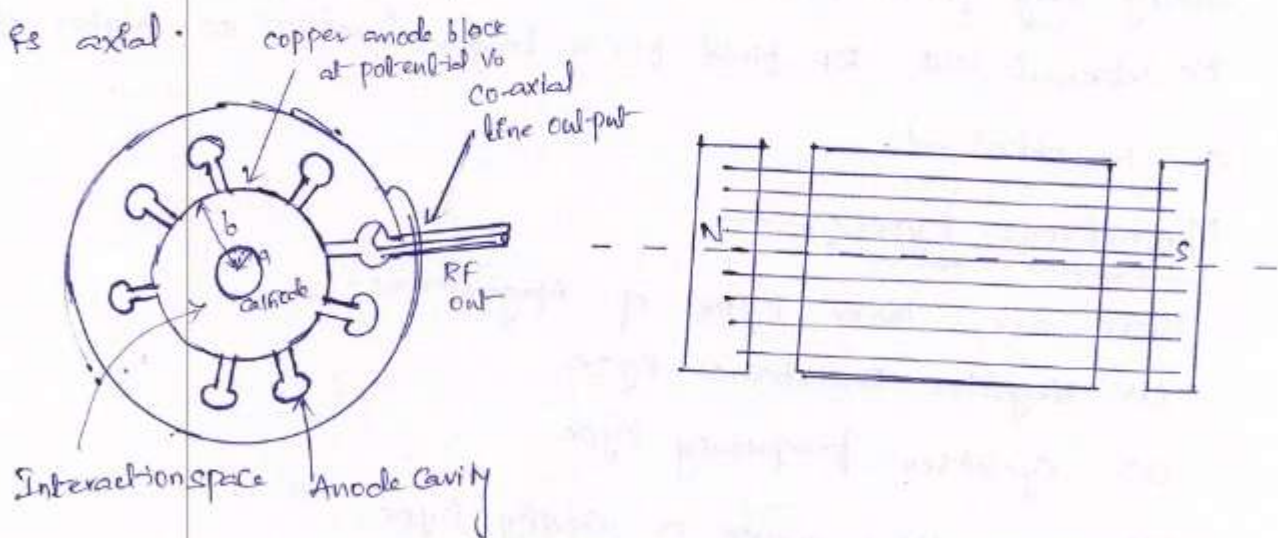
These are useful only for frequencies greater than 100MHz .

cavity types —

This type of Magnetrons depend upon the interaction of electrons with a rotating electromagnetic field of constant angular velocity. These provide oscillations of very high peak power and hence are very useful in Radar applications.

Cavity Magnetron : —

It is a diode usually of cylindrical configuration with a thick cylindrical cathode at the centre and a co-axial cylindrical block of copper as anode. In the anode block, cut a number of holes and slots which acts as resonant anode cavities. The space between anode and cathode is the interaction space and to one of the cavities is connected a coaxial line as waveguide for extracting the output. It is a cross field device as the electric field between anode and cathode is radial whereas the magnetic field produced by permanent magnets is axial.



The permanent magnet is placed such that the magnetic lines are parallel to the vertical cathode and perpendicular to the electric field between cathode and anode.

Operation : —

The cavity Magnetron shown has 8 cavities, that are tightly coupled to each other. A N-cavity tightly coupled system will have N-modes of operation each of which is uniquely characterised by a

Combination of frequency and phase of oscillation relative to the adjacent cavity. These modes must be self consistent so that the total phase shift around the ring of cavity resonators is $2n\pi$ where n is an integer. Therefore if ϕ_v represents the relative phase change of the ac electric field across adjacent cavities, then,

$$\phi_v = \frac{2\pi n}{N} \quad \text{where } n=0, \pm 1, \pm 2, \pm(\frac{N}{2}-1), \pm \frac{N}{2} \dots$$

i.e. $N/2$ mode of resonance can exist if N is an even number.

$$\text{if } n = \frac{N}{2}, \phi_v = \pi$$

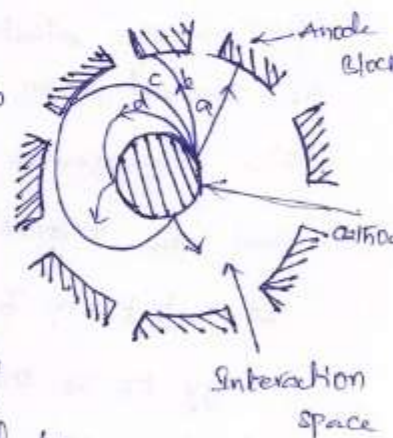
This mode of resonance is called the π -mode.

$$\text{If } n=0, \phi_v=0.$$

This is the zero mode, meaning there will be no RF electric field between anode and cathode and is of no use in magnetron operation.

Depending on the relative strengths of the electric and magnetic fields the electrons emitted from the cathode and moving towards the anode will traverse through the interaction space.

In the absence of magnetic field ($B=0$), the electron travels straight from the cathode to the anode due to the radial electric field force acting on it. (Indicated by the trajectory 'a' shown in figure).



If the magnetic field strength is increased slightly (i.e. for moderate value of B) it will exert a lateral force

bending the path of the electron as shown by path 'b' in figure. The radius of the path is given by $R = \frac{mv}{eB}$, that varies directly with electron velocity and inversely as the magnetic field strength.

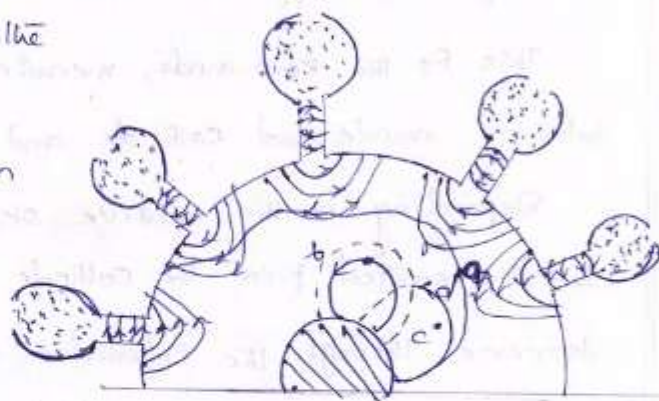
If the strength of the magnetic field is made sufficiently high so as to prevent the electrons from reaching the anode (as shown by path 'c') the anode current becomes zero. The magnetic field required to return electrons back to cathode just touching the surface of the anode is called the 'critical magnetic field' (B_c), the cut-off magnetic field.

→ If the magnetic field is made larger than the critical field ($B > B_c$), the electron experiences a greater rotational force and may return back to cathode quite faster. This is in the case of the absence of the RF field in the cavity of Magnetron.

→ If RF oscillations are assumed to be initiated due to some noise transient within the magnetron, the oscillations will be sustained by device operation. Self-consistent oscillations will be obtained if the phase difference between adjacent anode poles is $\frac{n\pi}{4}$, where 'n' is an integer. $n=4$ results in π -mode of operation. Here the anode poles are π radians apart in phase.

The dotted electron paths refer to the static fields with no RF signal.

The solid paths refer to the electron trajectory in the presence of RF oscillations in the interaction space. The electron 'a' is seen to be slowed down in presence of oscillations



thus transferring energy to the RF field are called 'favoured electron' and are responsible for bunching effect.

An electron 'b' is accelerated by the RF field and instead of transferring energy to the oscillations, takes energy from oscillations resulting in increased velocity. Hence electron spends very little time in the interaction space and is returned back to cathode. This electron does not participate in the bunching process so it is called 'unfavoured electron'. Similarly an electron 'c' which is emitted a little later to be in correct position moves faster and tries to catch up with electron 'a'. An electron emitted at 'd' will be slowed down to fall back in step with electron 'a'. This results in all favoured electrons like a, c, d to form a bunch and are confined to spokes or electron clouds.

The phase focussing effect of these favoured electrons imparts enough energy to the RF oscillations so that they are sustained.

Mathematical Analysis:-

Let the cathode and anode radius be 'a' and 'b' respectively and ϕ the angular displacement of the electron bends. Because of the cross field device electric and magnetic fields are perpendicular to each other and the path of the electrons in the presence of this cross field is naturally parabolic.

force on the electron is

$$F = Bev$$

In the direction of ϕ , the force component F_ϕ is given by,

$$F_\phi = eBv_\phi$$

where $v_\phi \rightarrow$ velocity in the direction of the radial distance r from the centre of the cathode cylinder.

Torque in ϕ direction is

$$T_\phi = r F_\phi = e r v_\phi B \quad \longrightarrow \textcircled{1}$$

Angular Momentum = angular velocity * moment of inertia

$$= \frac{d\phi}{dt} \times m r^2$$

$$\text{Time rate of angular momentum} = \frac{d}{dt} \left[\frac{d\phi}{dt} \times m r^2 \right] \longrightarrow \textcircled{2}$$

which gives the Torque in ϕ -direction.

Equating $\textcircled{1}$ and $\textcircled{2}$,

$$\frac{d}{dt} \left[\left[\frac{d\phi}{dt} \right] m r^2 \right] = e \cdot r \cdot v_\phi \cdot B$$

$$\text{i.e.} \quad 2mr \frac{d\phi}{dt} + m r^2 \frac{d^2\phi}{dt^2} = e \cdot r \cdot v_\phi \cdot B \quad \longrightarrow \textcircled{3}$$

We know that $v_\phi = \frac{dr}{dt}$, $r v_\phi = r \cdot \frac{dr}{dt}$

$$\int r v_\phi = \int r \cdot \frac{dr}{dt} = \frac{r^2}{2}$$

By integrating eqn(3) with respect to t

$$2mr \cdot \phi + m r^2 \frac{d\phi}{dt} = eB \cdot \frac{r^2}{2}$$

for a particular direction ϕ , $m p \dot{\phi}$ can be considered as a constant.

$$m p^2 \frac{d\phi}{dt} + C = e B \cdot \frac{r^2}{2} \longrightarrow (4)$$

from the Boundary conditions [i.e. at the Surface of the Cathode $r=a$ and $\frac{d\phi}{dt}=0$ being zero angular velocity at emission], we can determine the value of Constant 'C'.

$$0 + C = \frac{e B a^2}{2} \quad \text{or} \quad \boxed{C = \frac{e B a^2}{2}}$$

$\therefore$ from eqn (4), we can write

$$m p^2 \frac{d\phi}{dt} + \frac{e B a^2}{2} = e B \frac{r^2}{2}$$

$$\Rightarrow m p^2 \frac{d\phi}{dt} = \frac{e B}{2} [r^2 - a^2]$$

$$\Rightarrow \frac{d\phi}{dt} = \frac{e B}{2 m p^2} [r^2 - a^2]$$

$$\boxed{\frac{d\phi}{dt} = \frac{e B}{2 m} \left[1 - \frac{a^2}{r^2} \right]}$$

When $r=a$ [i.e. at Cathode], $\frac{d\phi}{dt}$ approaches zero.

and when $r \gg a$, $\frac{d\phi}{dt}$ approaches $(\omega)_{\max}$ (Maximum angular velocity)

$$\text{i.e.} \quad \boxed{\left[\frac{d\phi}{dt} \right]_{\max} = (\omega)_{\max} = \frac{e B}{2 m} = \frac{e B_0}{2 m}} \longrightarrow (5)$$

where $B = B_0$ is the cut-off magnetic flux density.

from the Conservation of energy,

Potential Energy of electron = Kinetic energy of electron

$$\text{i.e.} \quad e V_0 = \frac{1}{2} m v^2$$

$$\boxed{e V_0 = \frac{m}{2} [v_p^2 + v_\phi^2]} \longrightarrow (6)$$

where v_p and v_ϕ are components in r and ϕ directions in cylindrical coordinates.

$$v_p = \frac{dr}{dt} \quad \text{and} \quad v_\phi = r \cdot \frac{d\phi}{dt}$$

Re writing eqn (6),

$$eV_0 = \frac{m}{2} \left[\left(\frac{dp}{dt} \right)^2 + p^2 \left(\frac{d\phi}{dt} \right)^2 \right]$$

from eqn (5),

$$\left(\frac{d\phi}{dt} \right) = (\omega)_{\max} \left[1 - \frac{a^2}{p^2} \right]$$

$$\therefore eV_0 = \frac{m}{2} \left[\left(\frac{dp}{dt} \right)^2 + p^2 (\omega)_{\max}^2 \left[1 - \frac{a^2}{p^2} \right]^2 \right] \longrightarrow (7)$$

At anode, $p=b$, $\frac{dp}{dt}=0$. Substituting these boundary conditions in eqn (7),

$$\frac{m}{2} \left[b^2 (\omega)_{\max}^2 \left[1 - \frac{a^2}{b^2} \right]^2 \right] = eV_0 \longrightarrow (8)$$

$$\text{also, } (\omega)_{\max}^2 = \left[\frac{eB_c}{2m} \right]^2$$

where $B_c \rightarrow$ Cut-off value of the Magnetic flux density.

Substituting the value $(\omega)_{\max}^2$ in eqn (8),

$$\frac{m}{2} b^2 \left[\frac{eB_c}{2m} \right]^2 \times \left[1 - \frac{a^2}{b^2} \right]^2 = eV_0$$

$$\therefore \frac{eB_c^2 b^2}{8m} \left[1 - \frac{a^2}{b^2} \right]^2 = eV_0 \Rightarrow B_c = \frac{(8V_0 m/e)^{1/2}}{b \left[1 - \frac{a^2}{b^2} \right]}$$

Since $b \gg a$ and $\frac{a^2}{b^2}$ can be neglected.

$$\therefore B_c = \frac{1}{b} \sqrt{\frac{8mV_0}{e}}$$

$\rightarrow$ Hull's Cut-off Magnetic flux density equation

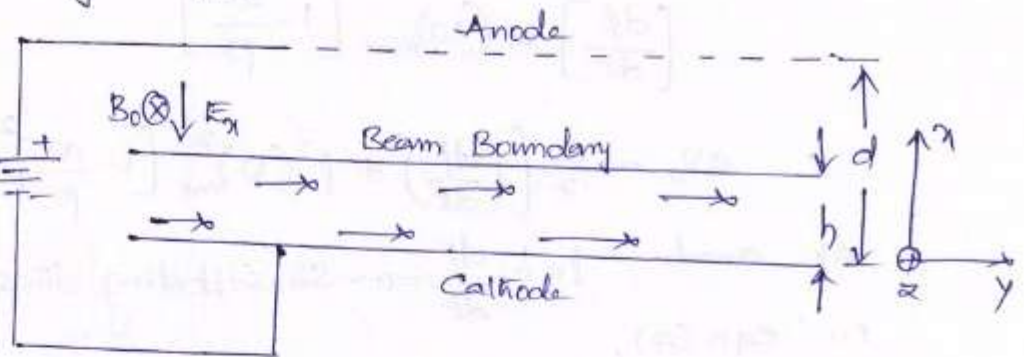
At this cut-off value of magnetic field, the electron touches the anode and the value of this magnetic field can be obtained by the anode voltage. The above equation is called the "Hull's Cut-off Voltage Equation".

$$V_0 = \frac{B_c^2 b^2}{8m} \left[1 - \frac{a^2}{b^2} \right]^2$$

Hull's Cut-off voltage equation

Hartree Condition — The Hull Cut-off Condition determines the anode voltage or magnetic field necessary to obtain non zero anode current as a function of the magnetic field or anode voltage in the absence of an electromagnetic field.

The electron beam lies within a region extending a distance V_0 from the cathode, where h is known as the hub thickness.



The spacing between the cathode and anode is d . The electron motion to be in the positive y direction with a velocity

$$v_y = \frac{En}{B_0} = \frac{1}{B_0} \frac{dV_0}{dx} \rightarrow (1)$$

where $B_0 = B_z$ is the magnetic flux density in the positive z -direction.
 V = potential.

B_0 along z
 $\rightarrow$ from Biot-Savart
 $EXH = \text{wave}$
 $EXH = v_y$
 $a_{nx} - a_z = a_y$

from the principle of energy conservation, we have

$$\frac{1}{2} m v_y^2 = eV_0 \rightarrow (2)$$

from (1) & (2),

$$\frac{1}{2} m v_y^2 = eV_0 \Rightarrow \frac{1}{2} m \left[\frac{1}{B_0} \frac{dV_0}{dx} \right]^2 = eV_0$$

$$\Rightarrow \left(\frac{dV_0}{dx} \right)^2 = \frac{2eV_0 B_0^2}{m} \Rightarrow \frac{dV_0}{dx} = \sqrt{\frac{2eV_0}{m}} B_0 \rightarrow (3)$$

$$\left(\frac{m}{2eB_0^2} \right)^{1/2} \frac{dV_0}{\sqrt{V_0}} = dx \rightarrow (4)$$

By integrating eqn(4), $\left[\frac{m}{2eB_0^2} \right]^{1/2} \frac{V_0^{-1/2+1}}{-1/2+1} = x \Rightarrow \left(\frac{m}{2eB_0^2} \right)^{1/2} * (-\sqrt{V_0}) = x$

$$\Rightarrow \frac{m}{2eB_0^2} * 4V_0 = x^2 \Rightarrow \boxed{V_0 = \frac{e B_0^2 x^2}{2m}} \rightarrow (5)$$

where the constant of integration has been eliminated for $V=0$ at $x=0$.

The potential and electric field at the hub surface are given by

$$\boxed{V(h) = \frac{e}{2m} B_0^2 h^2} \rightarrow (6)$$

and $E_x = -\frac{dV(h)}{dx} = -\frac{e}{m} B_0^2 h \rightarrow (7)$

The potential at the anode is given by

$$\begin{aligned} V_0 &= -\int_0^d E_x dx = -\int_0^h E_x dx - \int_h^d E_x dx \\ &= V(h) + \frac{e}{m} B_0^2 h(d-h) \\ &= \frac{e}{2m} B_0^2 h^2 + \frac{e}{m} B_0^2 h(d-h) \\ &= \frac{e}{m} B_0^2 h \left[\frac{h}{2} + d - h \right] = \frac{e}{m} B_0^2 h \left[\frac{h + 2d - 2h}{2} \right] \\ &= \frac{e}{m} B_0^2 h \left[\frac{2d - h}{2} \right] = \frac{e}{m} B_0^2 h \left[d - \frac{h}{2} \right] \rightarrow (8) \end{aligned}$$

The electron velocity at the hub surface is obtained from eqn (1) and eqn (7) as

$$v_y(h) = \frac{e}{m} B_0 h \rightarrow (9)$$

for synchronism, the electron velocity is equal to the phase velocity of the slow wave structure. That is,

$$\frac{\omega}{\beta} = \frac{e}{m} B_0 h \rightarrow (10)$$

for the π -mode operation, the anode potential is given by

$$V_{oh} = \frac{\omega B_0 d}{\beta} = \frac{m \omega^2}{2e \beta^2} \rightarrow (11)$$

This is the "Hartree anode voltage equation" that is a function of the magnetic flux density and the spacing between cathode and anode.

Magnetron in π -mode

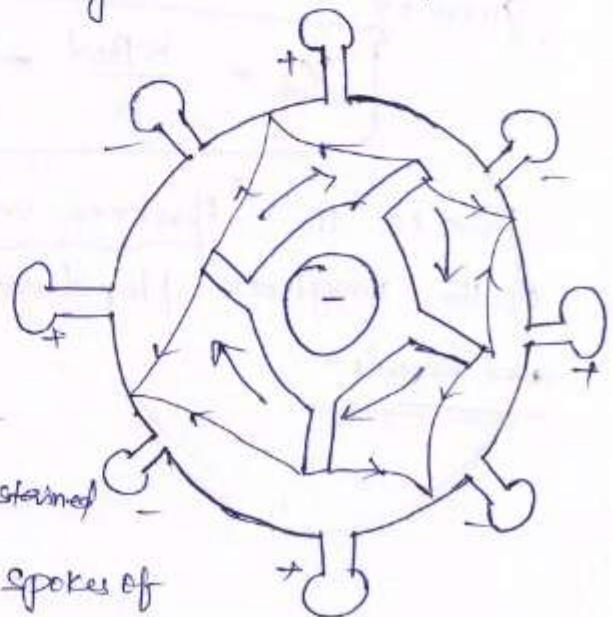
When the phase shift between two adjacent cavities is π , it is called as ' π -mode of operation'. Oscillations can build-up only in π -mode. The lines of force of the field traversed by electrons in magnetron under π -mode shown in fig. 1.



In the absence of RF E-field electron follows the path 'a' as shown. Such electron is called as 'unfavoured electrons'. These electrons which follow the path 'b' are called 'favoured electrons'. These electrons will boost the RF field in the gaps of cavities since all the cavities are tuned to a resonant frequency. Under this condition magnetron sustains oscillations and delivers a high output power.

→ When an electron cloud sweeps past a resonator, it's shock excites resonator into self oscillation. These oscillations are combined to set-up an electric wave in the interaction space between cathode and anode. Because the wave adds algebraically with the static DC potential, it causes both acceleration and deceleration of the electron cloud. Due to this action, both velocity and density modulates the cloud, causing it to bunch up into a spoked wheel formation.

→ In an oscillating magnetron, the spokes of the electron cloud wheel rotate in close synchronism with the wave phase velocity. The electron drift velocity (v_d) must match the phase velocity for sustained



oscillations. Under this condition, spokes of wheel continue to revolve, continually exciting cavity resonators and creating a sustained wave at microwave frequency. Because of this process, electrons lose energy to the wave by interaction.

MWE
13-IV-6
The electrons tend to lose velocity and fall into the cathode, where they are collected. Thus, in the oscillation mode there is no anode current.

Due to the more interaction space, large output power will be delivered from the cavity output.

Oscillations in the π -mode can occur at voltage between a critical level called "Hartree potential" and the "Full potential".

The Hartree potential is

$$V_{Hr} = \frac{2\pi f B}{N} (b^2 - a^2)$$

$N \rightarrow$ no. of cavity resonators

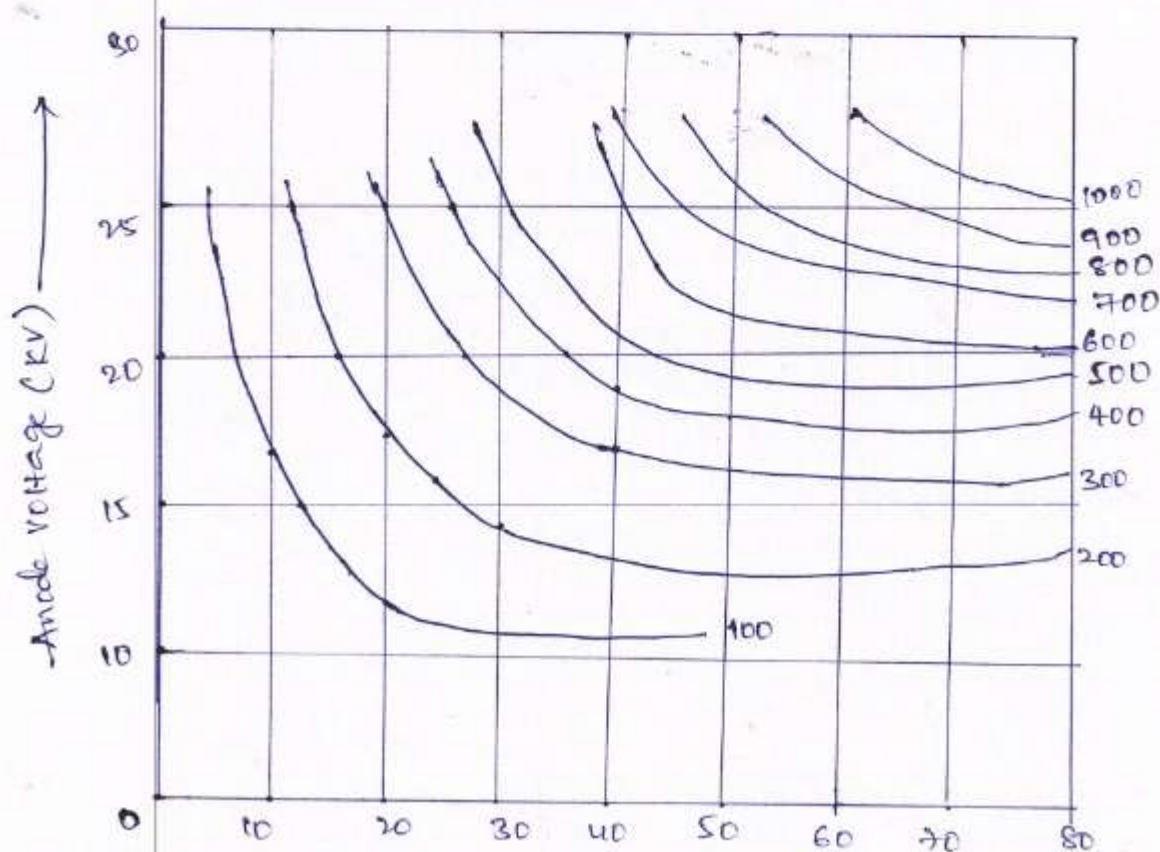
$B \rightarrow$ flux density

$f \rightarrow$ oscillating frequency.

Output characteristics:

Output characteristics are best studied by means of Rieke diagrams.

The operating conditions of Magnetron for a given load can be obtained from these diagrams, which is basically a plot of anode voltage vs current with power output, efficiency, flux density, frequency deviation as parameter.



Rieke Diagram performance chart of magnetron

Output Characteristics:

Power output: In excess of 250 kW (pulsed modulation)

10 mW (VHF Band) 2 mW (X Band) 8 kW (at 95 GHz)

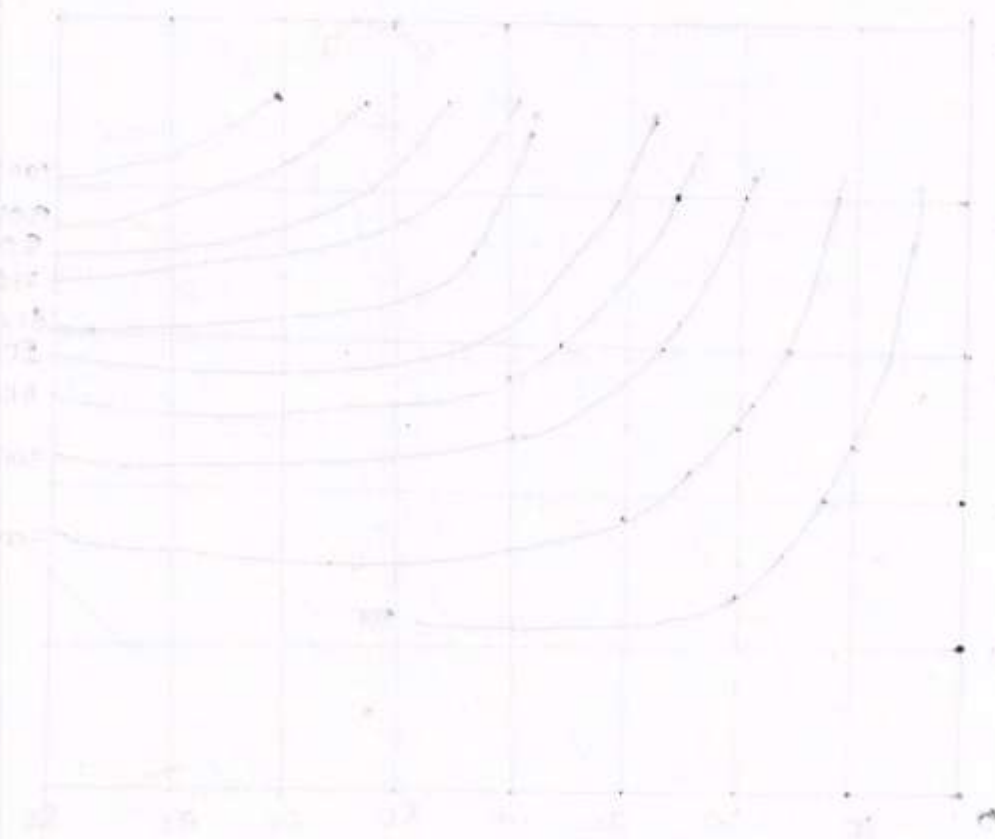
Frequency: 500 MHz to 129 GHz

Duty cycle: 0.1%.

Efficiency: 40% to 70%.

Applications of Magnetron:

- pulsed radar is the most important application with large pulse powers.
- voltage tunable magnetrons are used in sweep oscillators in telemetry and in missile applications.
- fixed frequency, continuous wave magnetrons are used for industrial heating and microwave ovens.



MW
V-Ir (3)

Microwave Solid State Devices

for development of better low noise, high frequency, greater Bandwidth lesser switching time and other improvements in the performance characteristics for achieving the amplification, oscillation switching, limiting or frequency multiplication can be possible by Microwave solid state devices.

The application of two-terminal semiconductor devices at microwave frequencies is increasing day to day. The continuous wave, average and peak power outputs of these devices at higher microwave frequencies are much larger than those obtainable with the best power transistor. The common characteristic of all active two-terminal solid state devices is their negative resistance. The real part of their impedance is negative over a range of frequencies.

In a positive resistance, the current through the resistance and the voltage across it are in phase. The voltage drop across a positive resistance is positive and a power $[I^2R]$ is dissipated in the resistance. In a negative resistance, the current and voltage are out of phase by 180° . The voltage drop across a negative resistance is negative and a power $[of -I^2R]$ is generated by the power supply associated with the negative resistance. In other words, positive resistances absorb power [passive devices], whereas negative resistances generate power (active devices). (whereas negative resistances generate power (active devices))

Classifications:

Semiconductor microwave devices include three terminal devices such as Bipolar Junction Transistor and field effect Transistors. two terminal devices such as Transferred electron devices (Gunn diodes, LSA diodes), Avalanche transit time devices (IMPATT, TRAPATT, BARITT Parametric devices), tunnel diodes, varactors, Quantum electronic devices such as MASERS, semiconductor lasers and Infrared devices.

Applications :-

- Communication systems
- Radars
- Navigation
- Medical and Biological equipment
- In industrial electronic products etc.

Transferred Electron Devices (TED's) :-

TED's are bulk devices having no junction, or gates as compared to microwave transistors which operate with either junction or gates.

TED's are fabricated from compound semiconductors such as GaAs, InP (Indium phosphate) or CdTe (Cadmium Telluride) as against Ge and Si of transistors.

TED's operate with hot electrons whose energy is very much greater than the thermal energy. Transistors operate with warm electrons whose energy is not much greater than their thermal energy. Gunn diode is an example of this kind.

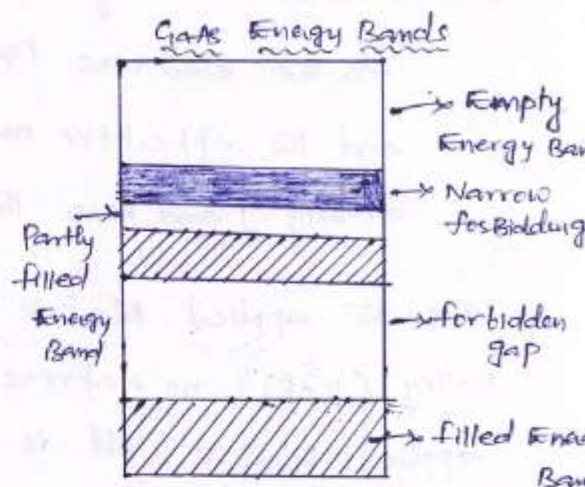
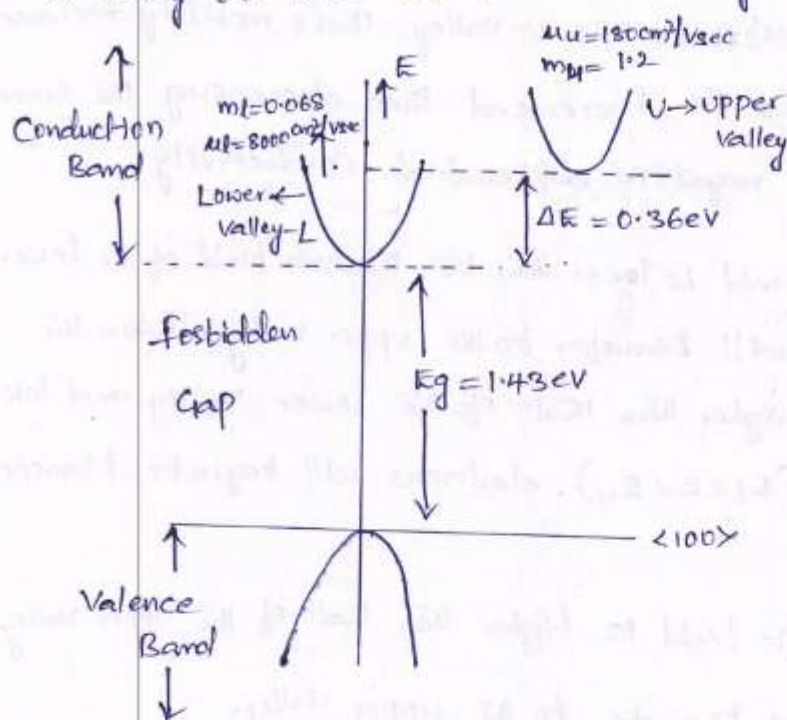
Gunn Diode :-

Gunn Diodes are named after J.B. Gunn (1963), who discovered periodic fluctuations of current passing through the n-type GaAs specimen when the applied voltage exceed a certain critical value. certain critical value.

Basic mechanism involved in the operation of bulk-n type GaAs devices is the transfer of electrons from lower conduction valley to upper subsidiary valley i.e. L-valley to U-valley.

The curvature of the two valleys in the conduction band also called the subbands are different so that an electron in

L-Valley has a smaller effective mass than one in the U-Valley. The different effective masses mean different mobilities for the L-valley and the U-Valley respectively. The ratio of density of states in the U-Valley to that in the L-Valley is about 60. ~~MARK MARK~~



Data for Two Valleys in GaAs

Valley	Effective Mass m_e	Mobility μ	Separation ΔE
Lower	$m_{eL} = 0.068$	$\mu_L = 8000 \text{ cm}^2/\text{Vsec}$	$\Delta E = 0.36 \text{ eV}$
Upper	$m_{eU} = 1.2$	$\mu_U = 180 \text{ cm}^2/\text{Vsec}$	$\Delta E = 0.36 \text{ eV}$

→ Electron densities in the lower and upper valleys remain the same under the equilibrium condition.

At low (or zero) electric fields, conduction electrons are distributed in a manner determined by energy separation, lattice temperature and the density of states. With the typical values stated that, most of the electrons at low electric fields and at low lattice temperature will occupy states in the L-valley and carry current $J = \sigma E$.

with $\sigma = n \mu_L = e n_0 \mu_L \rightarrow n_L$ - carrier concentration in L-valley and is assumed to be equal to the total carrier concentration n_0 and μ_L is the mobility in the L-valley.

As the applied field is increased, the electrons gain energy from it and move upward in the U-valley. The probability of this inter valley transfer of electrons is good as there are many available states in the U-valley.

As the electrons transfer to this U-valley, their mobility decreases and the effective mass is increased thus decreasing the current density J and hence the negative differential conductivity.

When the applied electric field is lower than the electric field of the lower valley ($E < E_L$), no electrons will transfer to the upper valley. When the applied electric field is higher than that of the lower valley and lower than that of upper valley ($E_L < E < E_U$), electrons will begin to transfer to the upper valley.

When the applied electric field is higher than that of the upper valley ($E_U < E$), all electrons will transfer to the upper valley.

If electrons density in the lower and upper valleys are n_l and n_u respectively, the conductivity of the n-type GaAs is

$$\sigma = e(\mu_l n_l + \mu_u n_u)$$

where $e \rightarrow$ electron charge
 $\mu \rightarrow$ electron mobility
 $n \rightarrow$ Electron density

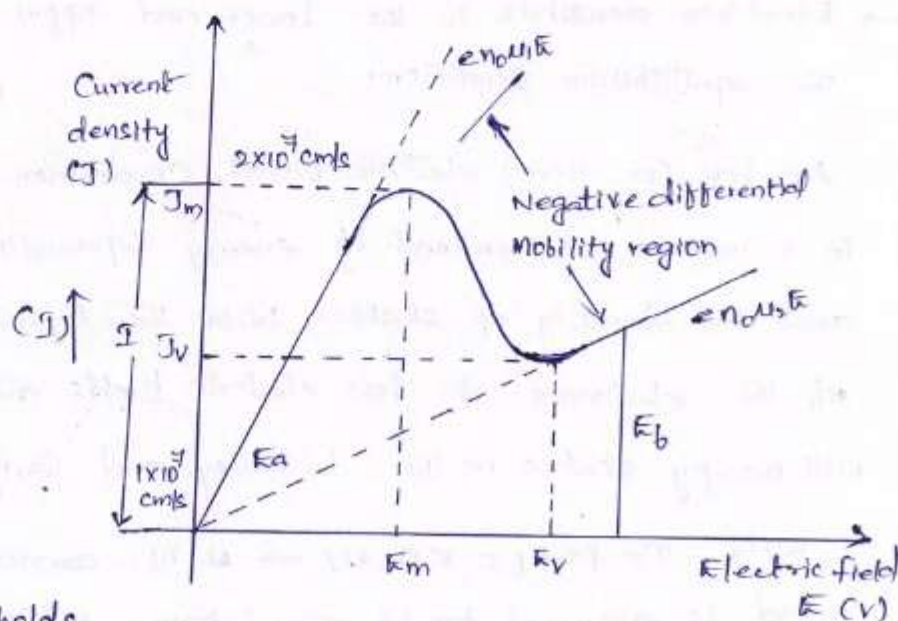
J-E characteristics:

Where
 $J_m \rightarrow$ Maximum Current density
 $J_v \rightarrow$ Valley Current density
 $E_m \rightarrow$ Maximum electric field required before the onset of negative conductance region

$E_a \rightarrow$ Maximum electric field for which $J = \sigma E$ is valid

$E_b \rightarrow$ Electric field for which $J = e n_u \mu_u E$ holds

$E_v \rightarrow$ Electric field corresponding to J_v .



with
1-IV-9

The region of the characteristic between V_{TN} and V_{TP} where current density decreases with increasing electric field is one of negative differential resistivity (NDR).

$$\text{Since } J = vE = en\mu E$$

Differentiating above equation with respect to E , we get

$$\frac{dJ}{dE} = v \left[1 + \frac{E}{v} \frac{dv}{dE} \right]$$

$$\frac{dJ}{dE} = v + E \frac{dv}{dE}$$

The condition for negative conductance (or resistance) is

$$\frac{dJ}{dE} < 0 \text{ or } 1 + \frac{E}{v} \frac{dv}{dE} < 0$$

If $\frac{dE}{dv}$ satisfies the above condition, i.e. if differential conductivity is negative, then the slope of the $V-I$ characteristic will be negative and the device will present a net negative conductance to the external circuit.

Ridley - Watkins - Hilsun (RWH) Theory : $\rightarrow$

In 1954, Shockley suggested that two terminal negative resistance devices using semiconductors may have advantages over transistor at high frequencies.

In 1961, Ridley and Watkins described a new method for obtaining negative differential mobility in semiconductors. The principle involved is to heat carriers in a light mass, high mobility subband with an electric field so that the carriers can transfer to a heavy mass, low mobility higher energy subband when they have a high enough temperature. Their theory for achieving negative differential mobility in bulk semiconductors by transferring electrons from high mobility energy bands to low mobility energy bands was taken a step further by Hilsun in 1962.

$\rightarrow$ on the basis of the Ridley - Watkins - Hilsun theory the band structure of a semiconductor must be such that

The conditions are

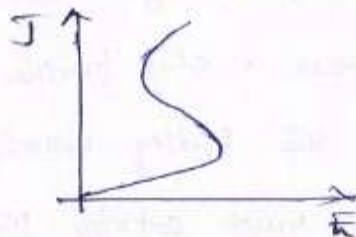
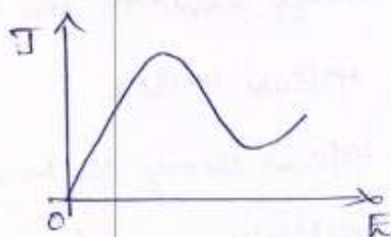
- (1) The separation energy between the bottom of the lower valley and the bottom of the upper valley must be several times greater than the thermal energy at room temperature i.e. $\Delta E > kT$ or $\Delta E > 0.026 \text{ eV}$.
- (2) The separation energy between the valleys must be smaller than the gap energy between the conduction and valance bands. This means $\Delta E < E_g$.
Otherwise the semiconductor will breakdown and become highly conductive before the electrons begin to transfer to the upper valleys because electron hole pair formation is created.
- (3) Electrons in the lower valley must have high mobility, small effective mass and low density of state whereas in the upper valley must have low mobility, large effective mass and a high density of state. Electron velocities (dE/dk) must be much larger in the lower valleys than in the upper valley.

Si and Ge do not meet all these criteria. Some compound semiconductors such as GaAs, InP and CdTe satisfies the criteria.

Differential Negative Resistances: —

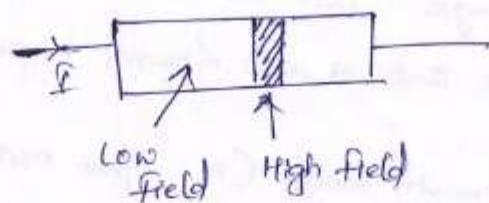
The fundamental concept of the Ridley Watkins Hilsen (RWH) theory is the differential negative resistance developed in a bulk solid state III-V compound when either a voltage (or E-field) or a current is applied to the terminals of the sample.

There are two modes of negative resistance devices: They are
(1) Voltage Controlled mode (2) Current-Controlled Mode

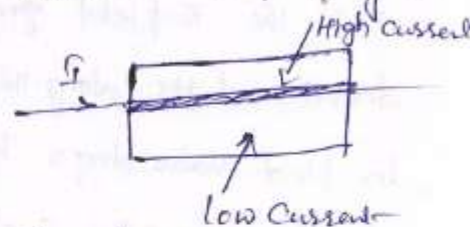


→ In the voltage controlled mode the current density can be multivalued, whereas in the current controlled mode the voltage can be multivalued.

MWE
1-IV-10
In the voltage controlled negative resistance mode high field domains are formed, separating two low field regions. The interfaces separating low and high field domains lie along equipotentials, thus they are in planes perpendicular to the current direction.

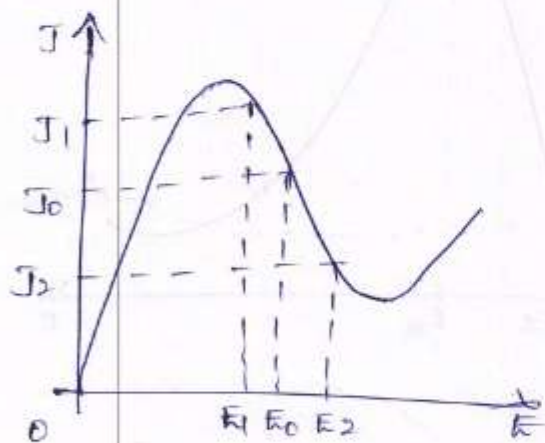


→ In the current controlled negative resistance mode, splitting the sample results in high current filaments running along the field direction.

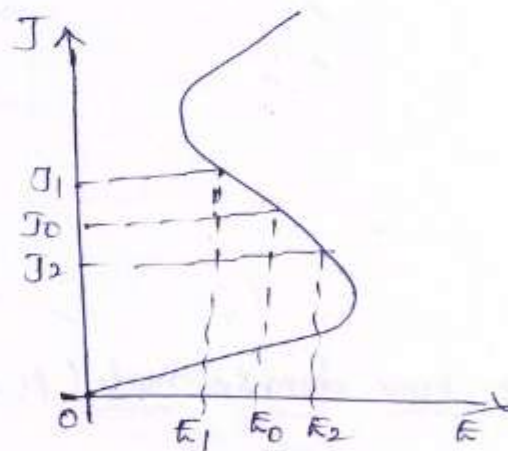


Mathematically, the negative resistance of the sample at a particular region is

$$\frac{dI}{dV} = \frac{dJ}{dE} = \text{Negative resistance.}$$



Voltage Controlled Mode



Current Controlled Mode.

Gunn oscillation Modes:

$$10^{12}/\text{cm}^2 \leq n_{0L} \leq 10^{14}/\text{cm}^2$$

Most Gunn effect diodes have the product of doping and length n_{0L} greater than $10^{12}/\text{cm}^2$. However, the mode that Gunn himself observed had a product n_{0L} that is much less.

When the product of n_{0L} is greater than $10^{12}/\text{cm}^2$ in GaAs, the space charge disturbances in the specimen increases exponentially in space and time. Thus a high field domain is formed and moves from the

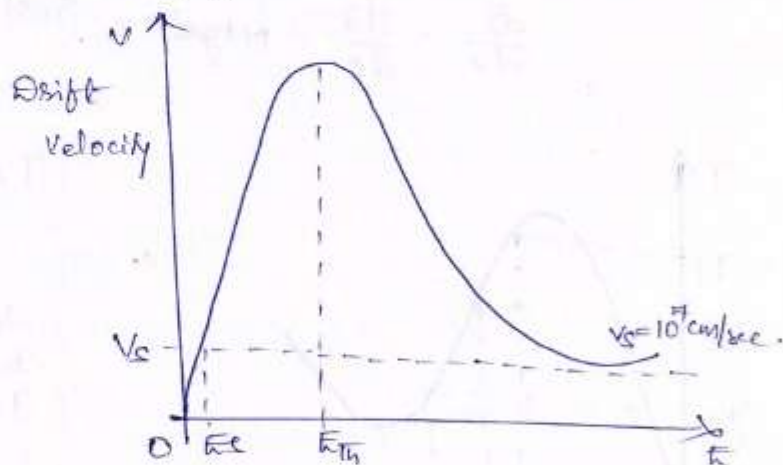
cathode to anode.

The frequency of oscillation is given by

$$f = \frac{V_{dom}}{L_{eff}} \quad \begin{matrix} \text{(domain velocity)} \\ \text{(Effective length)} \end{matrix}$$

$L_{eff} \rightarrow$ Effective Length that the domain travels from the time it is formed until the time that a new domain begins to form.

The normal Gunn domain mode (or Gunn oscillation mode) is operated with the E-field greater than the threshold field ($E > E_{Th}$). The high field domain drifts along the specimen until it reaches the anode or until the low field value drops below the sustaining field E_s required to maintain V_s . The sustaining drift velocity for GaAs is $V_s = 10^7$ cm/sec. Since the electric drift velocity v varies with the electric field. There are three possible domain modes for the Gunn oscillation modes.



Transit time domain mode ($fL = 10^7$ cm/sec)

When the electron drift velocity V_d is equal to the sustaining velocity V_s , the high field domain is stable.

$$\text{i.e. } V_d = V_s = fL = 10^7 \text{ cm/sec}$$

Then the oscillation period is equal to the transit time i.e. $T_o = T_t$.

The efficiency is below 10% because the current is collected only when the domain arrives at the anode.

Delayed Domain Mode [$10^6 \text{ cm/sec} < fL < 10^7 \text{ cm/sec}$]:-

When the transit-time is chosen such that the domain is collected while $E < E_{Th}$. A new domain cannot form until the field rises above

MWE
U-IV-11
threshold again. In this oscillation period is greater than transit time.

This delayed mode is also called 'Inhibited Mode'. The efficiency of this mode is about 20%.

Quenched domain Mode: — ($fL > 2 \times 10^7$ cm/sec)

If the bias field drops below the sustaining field E_s during the negative half cycle, the domain collapses before it reaches the anode. When the bias field swings back above threshold, a new domain is nucleated and the process repeats. Therefore the oscillation occurs at the frequency of the resonant frequency. The resonant frequency of circuit is several times the transit time frequency since one dipole does not have enough time to readjust and absorb the voltage of the other dipoles. Theoretically, the efficiency can reach 13%.

Limited space charge Accumulation (LSA) Mode: —
($fL < 2 \times 10^7$ cm/sec)

When the frequency is very high, the domains do not have sufficient time to form while the field is above the threshold. As a result, most of the domains are maintained in the negative conductance state during large fraction of the voltage cycle.

Any accumulation of electrons near the cathode has time to collapse while the signal is below threshold thus LSA mode is the simplest mode of operation and it consists of a uniformly doped semiconductor without any internal space charges. The internal electric field would be uniform and proportional to the applied voltage. The current in the device is then proportional to the drift velocity at this field level. The efficiency can reach 20%.

The oscillation period T_o should be not more than several times large than the magnitude of the dielectric relaxation time in the negative conductance region τ_d .

For the n-type GaAs, the product of doping and length (n_0L) is about $10^{12}/\text{cm}^2$.

output power of an LSA oscillator

Applications of Gunn Diode:-

- In Radar transmitters
- pulsed Gunn diode oscillators used in transponder for air traffic control (ATC) and in industry telemetry systems.
- Broadband Linear Amplifier.
- fast combinational and sequential logic circuits.
- Low and medium power oscillator in Microwave receivers.
- As pump sources in parametric amplifier.

* Disadvantage of Gunn diode is that it is very temperature dependent $0.5-3 \text{ MHz}/^\circ\text{C}$ change.

Avalanche Transit Time Devices:-

It is possible to make a microwave diode exhibit negative resistance by having a delay between voltage and current in an avalanche together with transit time through the material. Such devices are called "Avalanche transit-time devices". They use carrier impact ionization and drift in the high field region of a semiconductor junction to produce negative resistance at microwave frequencies.

There are three types of Avalanche transit time devices.

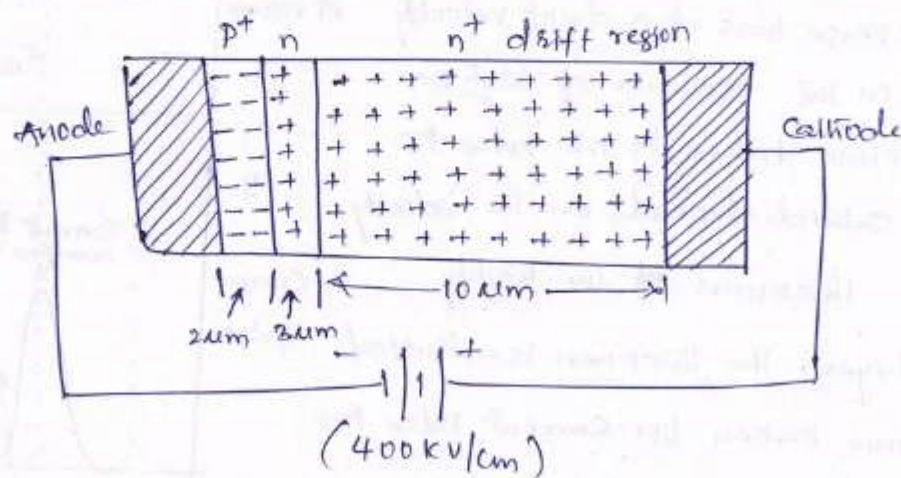
- (1) IMPATT: Impact Ionization Avalanche Transit time device
- (2) TRAPATT: Trapped Plasma Avalanche Triggered Transit device
- (3) BARITT: Barrier Injected Transit-Time device

IMPATT Diode:- [Impact Ionization Avalanche Transit Time Device]

Any device which exhibit negative resistance for DC will also exhibit it for AC i.e., If an AC voltage is applied current will rise when voltage falls at an AC rate. Hence negative resistance can also be defined as that property of a device which causes the current through it to be 180° out of phase with the voltage across it. This is the kind of negative resistance exhibited by IMPATT Diode i.e., if we show voltage and current have a phase difference of 180° , then negative resistance in IMPATT Diode is proved.

A combination of delay involved in generating avalanche current multiplication together with delay due to transit time through different space provides the necessary 180° phase difference between applied voltage and the resulting current in an IMPATT Diode.

IMPATT Diode, the junction being between the p^+ and n layers.



An extremely high voltage gradient (400 kV/cm) is applied to the IMPATT Diode eventually resulting in a very high current. A normal diode would very quickly breakdown under these conditions but IMPATT is constructed such that it will withstand these conditions repeatedly. Such a high potential gradient back biasing the diode causes a flow of minority carriers across the junction. Let us consider application of a RF AC voltage superimposed on top of the high DC voltage.

and holes by knocking them out of the crystal structure by so called "Impact Ionization". These additional carriers continue the process at the junction and it now increases into an avalanche.

If the original DC field was just at the threshold of allowing this situation to develop, this voltage will be exceeded during the whole of the RF positive cycle and the avalanche current multiplication will be taking place during this entire time. Since it is a multiplication process avalanche is not instantaneous. This process in fact takes a time such that the current pulse maximum at the junction occurs at the instant when RF voltage across the diode is zero and going negative. A 90° phase shift or phase difference between voltage and current has been achieved.

The current pulse as shown is situated at the junction. It does not stay there but moves towards the cathode due to applied reverse bias at a drift velocity dependent on the presence of high DC field. The time taken by the pulse to reach the cathode depends on the velocity and on the thickness of the highly doped 'n' layer. The thickness is adjusted such that time taken for current pulse to move from $v=0$ position to $v=\text{negative maximum}$ of RF cycle is exactly 90° . Hence voltage and current are 180° out of phase and "dynamic RF negative resistance has been proved to exist". Hence IMPATT Diode is useful both as an oscillator and as an amplifier. The resonant frequency is given by

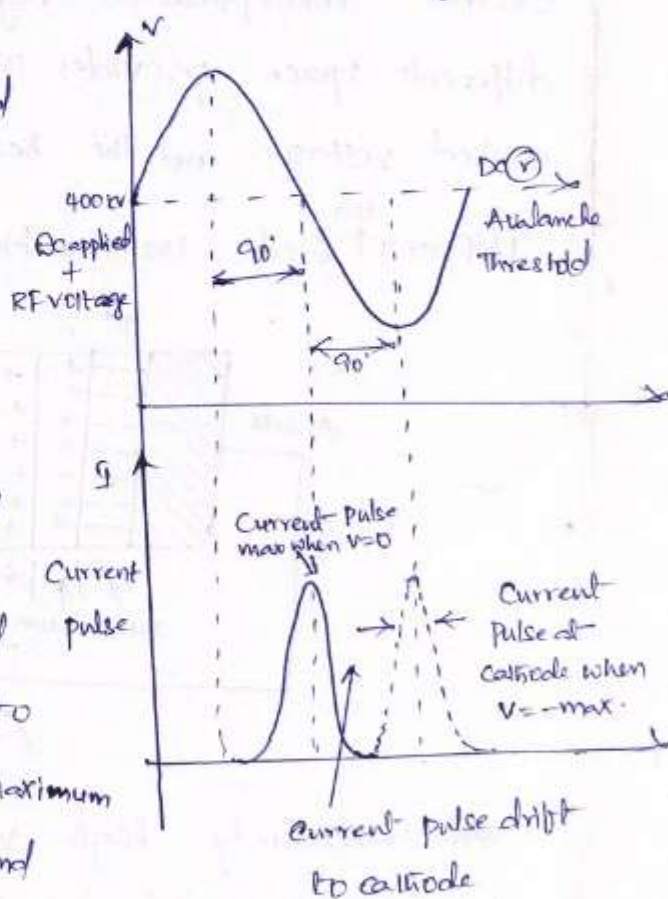


fig: V and I vs t characteristics

$$f = \frac{vd}{2L}$$

$v \rightarrow$ carrier drift velocity

$L \rightarrow$ length of the drift space charge region.

The efficiency η of IMPATT diode is given by

$$\eta = \left[\frac{P_{ac}}{P_{dc}} \right] = \frac{V_a}{V_d} \left[\frac{I_a}{I_d} \right]$$

where $P_{ac} \rightarrow$ AC power

$P_{dc} \rightarrow$ DC power

V_a and $I_a =$ AC voltage and current

V_d or $I_d =$ DC voltage and current

Performance characteristics:

Theoretical Efficiency $\eta = 30\%$

frequency : 1 to 300 GHz

Maximum output power for single diode: 5W in X band to 0.5W at 300 GHz.

pulsed powers = 4 kW.

Disadvantages of IMPATT Diode:

→ it is very noisy because avalanche is a noisy process.

→ Tuning range is not as good as Gunn Diode.

Applications of IMPATT Diode:

IMPATT is used as Microwave oscillator such as (i) Microwave Generator (ii) Modulated output oscillators (iii) Receiver local oscillators (iv) par amp pumps.

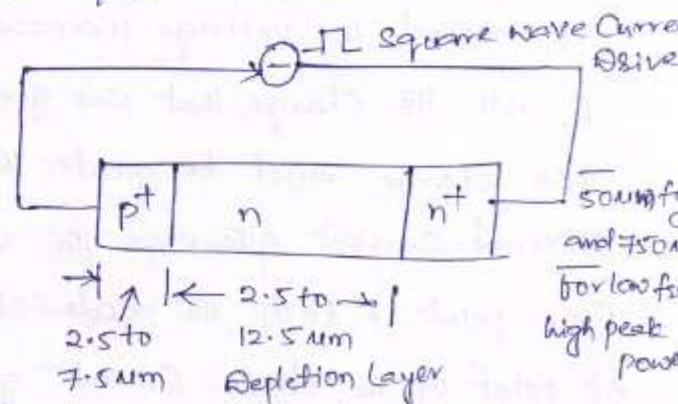
TRAPATT Diode [Trapped Plasma Avalanche Triggered Transit Device]

It is a high efficiency microwave generator capable of operating from several hundred MHz to several GHz.

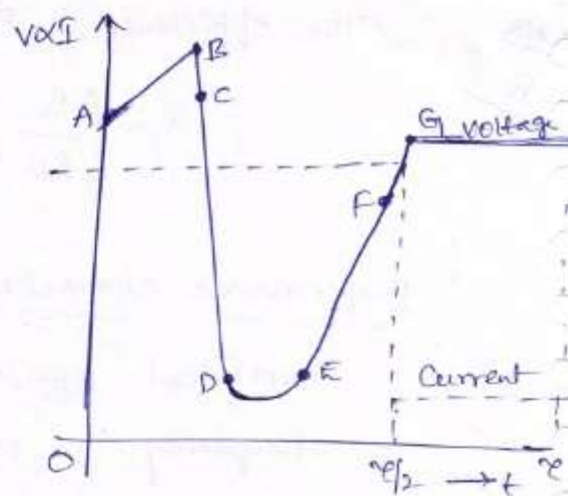
It is typically $p^+ - n - n^+$ Si or GaAs structure.

Operation:

A high field avalanche zone propagates through the diode and fills the depletion region with a dense plasma of electrons and holes that become trapped in the low field region behind the zone.



In figure, AB shows charging, BC shows plasma formation, DE shows plasma extinction, EF shows residual extinction and FG shows charging.



→ At point A, the electric field is uniform throughout the sample and its magnitude is large but less than the value required for avalanche breakdown. At A, charge carriers due to thermal generation results in charging of the diode like a linear capacitance AB driving the magnitude of the electric field above the breakdown voltage. When a sufficient number of carriers are generated, the particle current exceeds the external current and the electric field is depressed throughout the depletion region causing the voltage to decrease (BC). During this time interval the electric field is sufficiently large for the avalanche to continue and a dense plasma of electrons and holes are created. As some of the electrons and holes drift out of the ends of the depletion layer, the field is further depressed and traps the remaining plasma. The voltage decreases to point D. A long time is required to remove the plasma because the total plasma charge is large compared to the charge per unit time in the external current. At point E, the plasma is removed but a residual charge of electrons remains in one end of the depletion layer, the field is further depressed and a residual charge of holes on the other end. As the residual charge is removed, the voltage increases from point E to point F. At point F, all the charge that was generated internally has been removed. This charge must be greater than or equal to that supplied by the external current otherwise the voltage will exceed that at point A. From point F to G the diode charges up again like a fixed capacitor. At point G the diode current goes to zero for half a period and the voltage remains constant at V_A until the current comes back on and the cycle repeats.

MWE
1-IV-14

The Avalanche zone Velocity ' v_z ' is given by

$$v_z = \frac{dx}{dt} = \frac{J}{qNA}$$

where J = Current density

q = Electron charge, $1.6 \times 10^{-19} \text{ C}$

NA = Doping concentration.

The time delay of carriers in transit is utilized to obtain a current phase shift suitable for oscillation.

specifications:-

- (1) CW power : 1.3W between 8GHz to 0.5GHz
- (2) pulse power : 1.2kW at 1.1GHz
- (3) operating Voltage : 60-150V
- (4) Efficiency : 15 to 40%
- (5) Noise figure : > 30dB
- (6) frequency : 3 to 50 GHz.

→ Disadvantage of TRAPATT is it is very noisy.

Applications:

- Low power doppler Radars
- Local oscillators for Radars
- Microwave beacon Landing System
- Radio altimeter
- phased array Radars etc.

BARITT Diode [Barrier Injected Transit Time Device]

BARITT Diodes having long drift regions similar to those of IMPATT diodes. The carrier ~~transversing~~ traversing the drift region of BARITT diodes however are generated by minority carrier injection from forward biased junctions instead of being extracted from the plasma of an avalanche region.

→ BARITT Diodes are much less noisy than IMPATTs.

IMPATT devices employ impact ionization technique which is too noisy. In order to achieve low noise figures, impact ionization is

provided in BARITT devices. The minority injection is provided by punch through of the intermediate region (depletion region). This process is different from the sense that it has lower noise than impact ionization responsible for current injection in an IMPATT diode.

The negative resistance in BARITT device is obtained on account of the drift of the injected holes to the collector end of the diode, made of p type material.

→ An essential requirement

for the BARITT device is that the intermediate drift region be entirely depleted to cause punch through to the emitter-base junction without causing available breakdown of the base collector junction.

The Breakdown voltage is given by

$$V_{bd} = \frac{2 \cdot qNL^2}{2\epsilon_s} = \frac{qNL^2}{\epsilon_s}$$

The Breakdown electric field (E_{bd}) is given by

$$E_{bd} = \frac{V_{bd}}{L} = \frac{qNL}{\epsilon_s}$$

Performance characteristics:

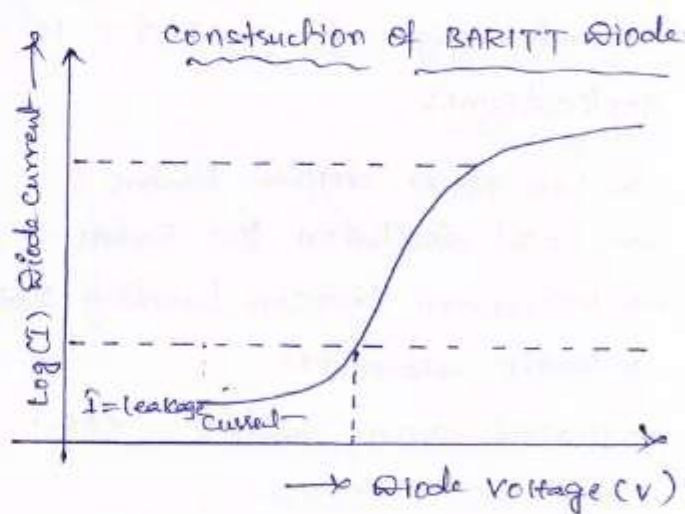
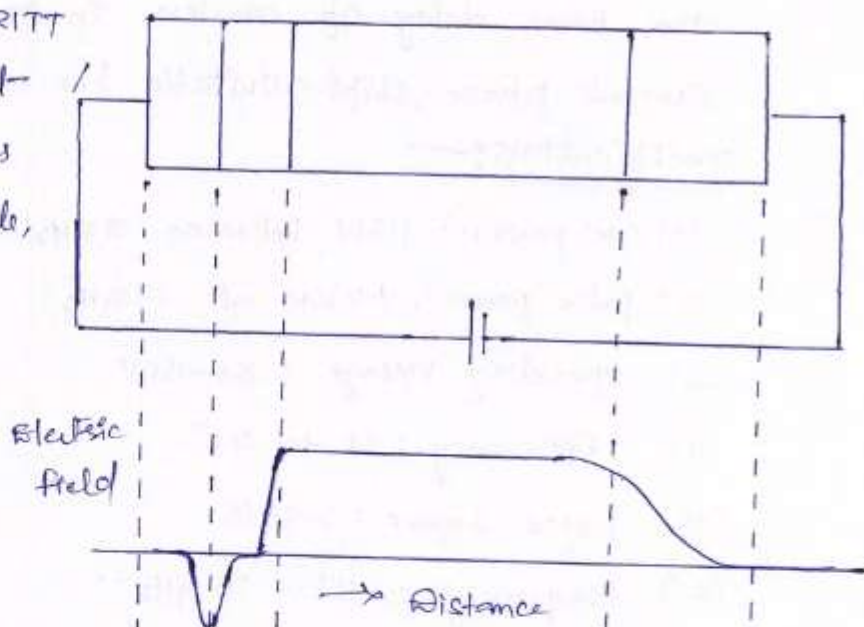
→ frequency: 4-10 GHz

→ Power: 50 mW at 4.9 GHz

→ Efficiency: 1.8%.

Applications:

→ Because of the low efficiency, BARITT is primarily used for amplifier.



V-I - characteristics of BARITT

→ Noise figure: 9 dB at 6.35 GHz
with 15 dB Gain

Microwave Measurements

Scattering Matrix:

The matrix which relates reflected normalized waves and incident normalized waves are called "Scattering Matrix".

Coefficients in the scattering matrix are called "scattering parameters" or "S-parameters" or "scattering coefficients".

$$\rightarrow [b]_{n \times 1} = [S]_{n \times n} [a]_{n \times 1}$$

$$\begin{bmatrix} b_1 \\ b_2 \\ \vdots \\ b_n \end{bmatrix} = \begin{bmatrix} S_{11} & S_{12} & \dots & S_{1n} \\ S_{21} & S_{22} & \dots & S_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ S_{n1} & S_{n2} & \dots & S_{nn} \end{bmatrix} \begin{bmatrix} a_1 \\ a_2 \\ \vdots \\ a_n \end{bmatrix}$$

$$[S] = \begin{bmatrix} S_{11} & S_{12} & \dots & S_{1n} \\ S_{21} & S_{22} & \dots & S_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ S_{n1} & S_{n2} & \dots & S_{nn} \end{bmatrix}_{n \times n} \rightarrow \text{scattering Matrix}$$

$$[b] \rightarrow \begin{bmatrix} b_1 \\ b_2 \\ \vdots \\ b_n \end{bmatrix} \rightarrow \text{Reflected Matrix}, [a] = \begin{bmatrix} a_1 \\ a_2 \\ \vdots \\ a_n \end{bmatrix} \rightarrow \text{Incident Matrix}$$

S_{nn} $\rightarrow$ power reflected from the (junction) port

$S_{mn} \rightarrow$ scattering coefficients resulting due to input at n^{th} port and output taken out of m^{th} port.

Isolator:

Isolator is a device which offers low resistance for forward direction i.e., from port ① to port ② and it offers ^{high} low resistance for reverse direction i.e., from port ② to port ①.

When the power is incident into port ① power not reflected into port ① itself so $S_{11} = 0$.

When the power is incident into port ②, power not reflected into port ② itself so $S_{22} = 0$.

→ When power is incident into port ①, power delivered from port ② so S_{21} is exists.

→ When the power is incident into port ②, power not delivered to port ①.
so S_{12} is zero.

$$[S] = \begin{bmatrix} S_{11} & S_{12} \\ S_{21} & S_{22} \end{bmatrix} = \begin{bmatrix} 0 & 0 \\ S_{21} & 0 \end{bmatrix} = \begin{bmatrix} 0 & 0 \\ 1 & 0 \end{bmatrix}$$

Circulator:—

for a N-port Circulator, Circulator has N-ports.

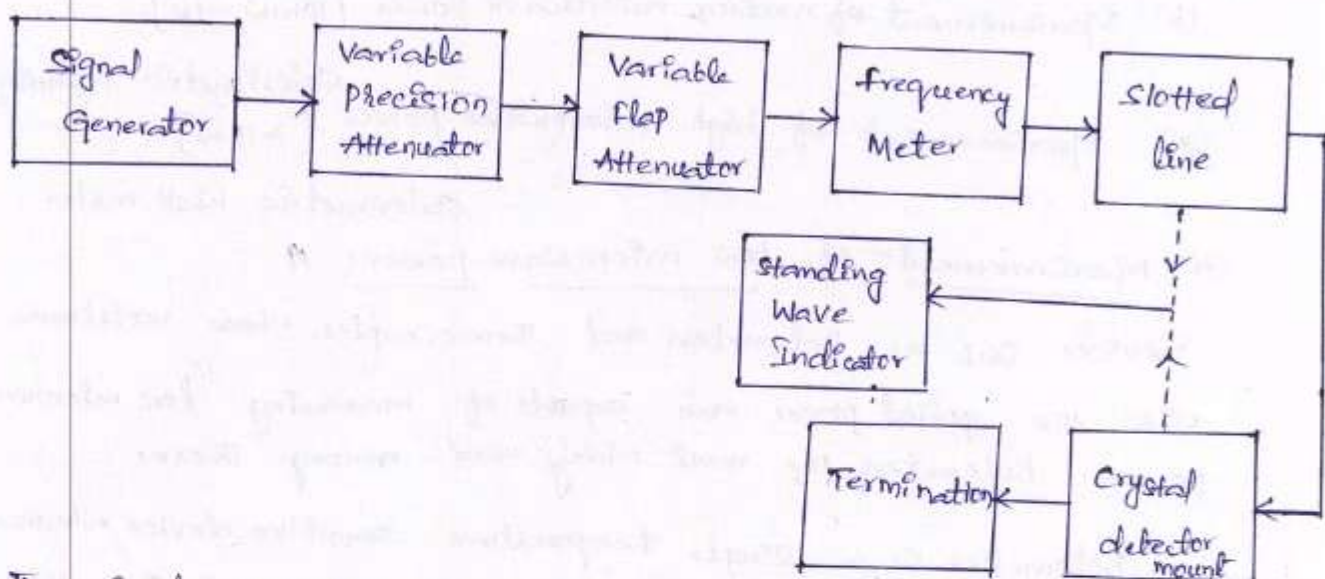
for a 3-port Circulator, Circulator consists of 3-ports.

The property of Circulator is power delivered from the next port of the power incident i.e. if power is incident into port ①, power delivered from port ②. If power is incident into port ②, power delivered from port ③. If power is incident into port ③, power delivered from port ①.

$$[S] = \begin{bmatrix} 0 & 0 & 1 \\ 1 & 0 & 0 \\ 0 & 1 & 0 \end{bmatrix} = \begin{bmatrix} S_{11} & S_{12} & S_{13} \\ S_{21} & S_{22} & S_{23} \\ S_{31} & S_{32} & S_{33} \end{bmatrix}$$

Microwave Bench:-

The general setup for measurement of any parameter in microwave is normally done by a Microwave Bench. Such a setup is shown below.



The signal generator is a Microwave source whose output is of the order of milliwatts. It could be a Gunn diode oscillator, a backward wave oscillator or a reflex klystron tube. It can provide either a continuous wave or square wave modulated at an audio rate which is normally 1 KHz. The precision attenuator can provide 0 to 50 dB attenuation above its insertion loss. The variable flap attenuator is also used, whose calibration can be checked against readings of the precision attenuator. A frequency meter is used for direct reading of frequency that consists of a single cylindrical cavity which can be adjusted resonant and is slot coupled to the waveguide. In a slotted line, a coupling probe moving along the waveguide can be used to detect the standing wave pattern present inside the waveguide. Which is used for measuring standing wave ratio. The crystal detector, inserted in the E-probe of the slotted line is contained in the crystal detector mount at the end of the waveguide, which is used to detect the modulated signal. The SWR (standing wave ratio) indicator is basically a sensitive tuned voltmeter, that provides direct readings of the SWR or its equivalent value.

Power Measurement:

The microwave power inside a waveguide is invariant with position of measurement and the power measured is the average power. The technique used to measure power depends on whether the power to be measured is low or high. Thus we divide the measurement of power into three categories.

(a) Measurement of low power $[0.01\text{ mW} - 10\text{ mW}]$ - Bolometer Technique

(b) Measurement of medium microwave power $[10\text{ mW} - 1\text{ W}]$ -

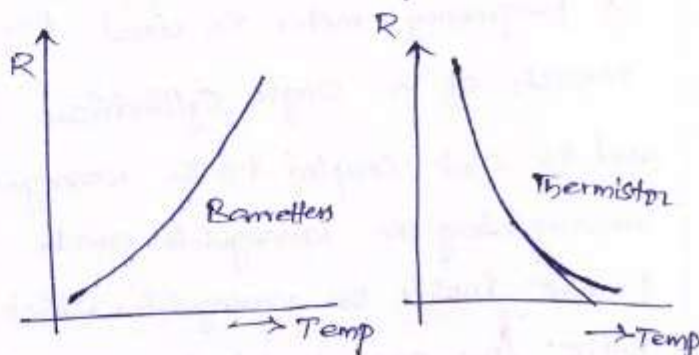
(c) Measurement of high microwave power $[>10\text{ W}]$ -
Calorimetric Technique

- Calorimetric Watt meter

(a) Measurement of low microwave power:

Devices such as Bolometers and thermocouples whose resistance changes with the applied power are capable of measuring low microwave powers. Bolometers are most widely used among these.

Bolometer is a simple temperature sensitive device whose resistance varies with temperature. These are of two types viz Barretters and thermistors. Barretters have positive temperature coefficient and their resistance increases with an increase in temperature. It basically consists of a short length of fine platinum wire mounted in a cathode, like an ordinary fuse. Thermistors have negative temperature coefficient of resistance and their resistance decreases with increase in temperature. Thermistors are basically semiconductor materials.

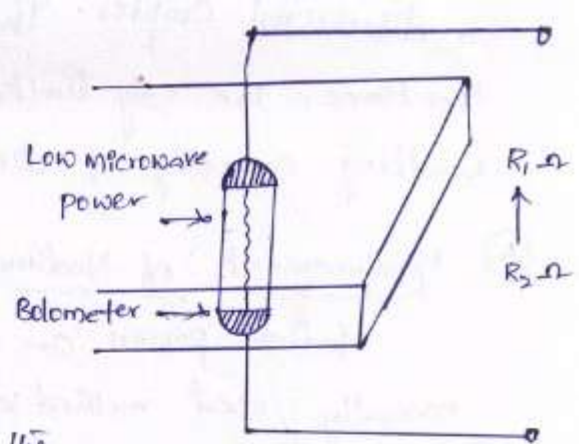


A Bolometer such as crystal diode is a square law device and it produces a current that is proportional to the applied power

i.e. square of the applied voltage, rather than the applied voltage.

Q-7-3

Bolometer is mounted inside the waveguide as shown, where the bolometer itself is used as a load, with the operation resistance as $R_1 \Omega$. Apply the low microwave power which is to be measured. Some power is absorbed in the Bolometer load and dissipated as heat and the resistance changes to R_2 . This change in resistance ($R_1 - R_2$) is proportional to the microwave power which can be measured using a bridge. Because of non-linear characteristics of Bolometer, accuracy is less.

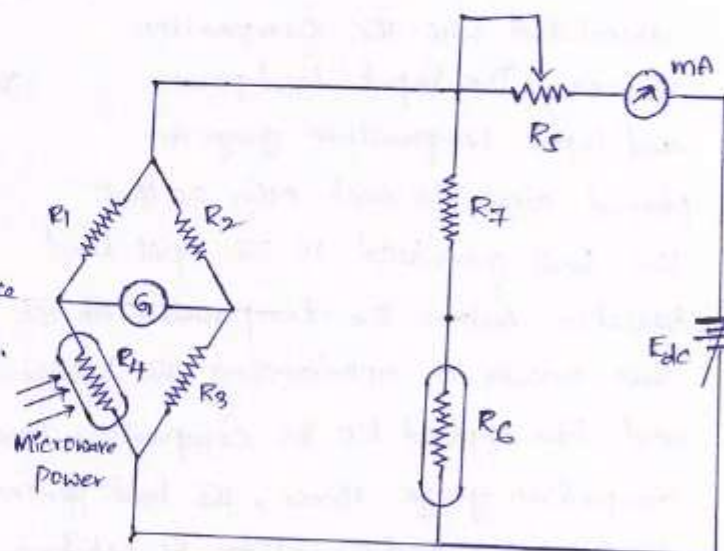


Balanced Bolometer Bridge Technique:-

In this technique, the Bolometer itself is made to be one of the arms of the bridge.

Initially, the bridge is balanced by adjusting R_5 , which varies the dc power applied to the bridge and the Bolometer element is brought to a predetermined operating resistance before microwave power is measured.

Let the voltage of the Battery be E_1 at balance. The microwave



Power is now applied and this power gets dissipated in the Bolometer then the Bolometer heats up and it changes its resistance. Therefore the bridge becomes unbalanced. The applied dc power is changed to E_2 to get back the balance and this change in dc Battery voltage ($E_1 - E_2$) will be proportional to the microwave power. Alternately, the detector G can be directly calibrated in terms of Microwave power so that when the bridge is unbalanced, the detector reads the microwave power directly.

→ Barretters and thermistors, both are limited in their power handling ability to about 10mW, so that powers greater than 10mW cannot be measured with them directly. However power measurement range can be increased by using

a directional coupler. This method extends the range of power by 1000 times. The only limitation of this method being the limited power handling capacity of directional coupler itself.

(b) Measurement of Medium power: — [10mW to 10W]

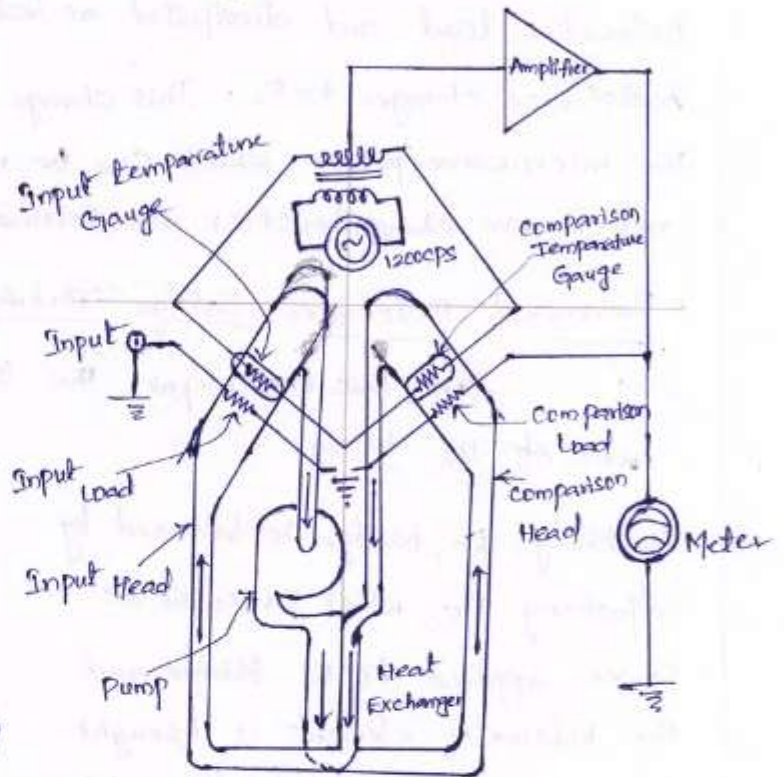
Medium powers can be measured by calorimetric techniques. The normally used method is the self balancing bridge technique.

It consists of identical temperature sensitive resistors or gauges in two arms, an indicating meter and two load resistors.

The input load resistor senses the unknown input microwave power and the comparison head is associated with the comparison power. The input load power and input temperature gauge are placed close to each other so that the heat generated in the input load resistor raises the temperature of the gauge.

This results in unbalancing the bridge. The signal due to the imbalance is amplified and then applied to the comparison load resistor which is placed closer to the comparison gauge. Hence, the heat generated in comparison load resistor is transferred to its gauge and the bridge is rebalanced. The meter measures the amount of power that is supplied to the comparison load in order to rebalance the bridge. It can be calibrated directly in terms of input microwave power. It is necessary that the characteristics of the two gauges be matched and also the heat transfer characteristics from each load be same for equal power dissipation in the two loads.

→ for quick balancing of the bridge and for efficient heat transfer from loads to the gauges, the components are immersed in an oil stream.



① High Power Measurement: [10W to 50KW]

Any power between 10W to 50KW is considered high power. These are normally measured by calorimetric watt meters. These meters can be either dry type or flow type.

A dry type calorimeter normally consists of a coaxial cable which is filled by a dielectric with a high hysteresis loss.

The flow type uses circulating water, oil or any liquid which is a good absorber of microwaves. The fluid after passing through the load experiences a temperature rise due to microwave energy. The difference between the temperature (T_1) of a known quantity of a liquid before entering the load and the temperature (T_2) of a known quantity after it emerges is a measure of the power which has been absorbed.

Attenuation Measurement: —

Attenuation is the ratio of input power to the output power and is expressed in decibels.

$$\text{Attenuation (in dBs)} = 10 \log \left[\frac{P_{in}}{P_{out}} \right]$$

where P_{in} = Input power and P_{out} = output power

The amount of attenuation can be measured by two methods.

- (a) power ratio Method (b) RF Substitution Method

④ Power Ratio Method: —

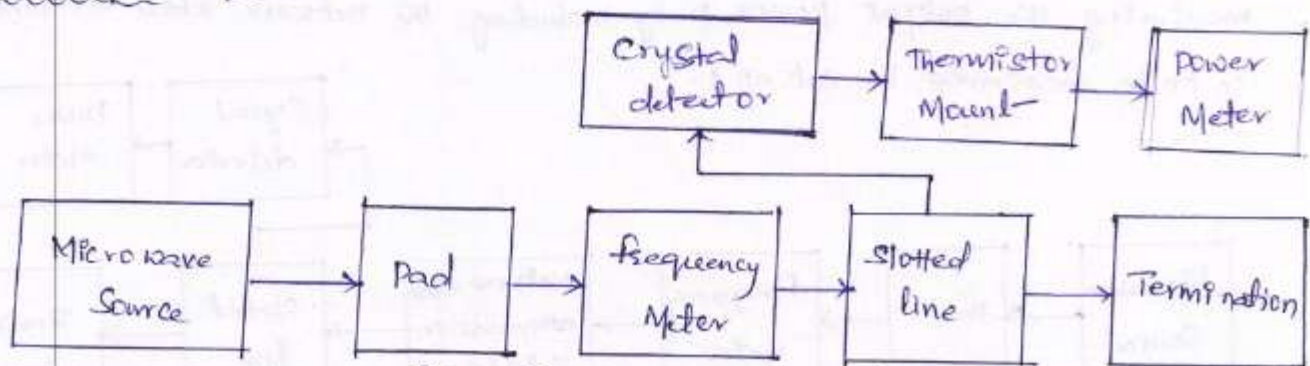


figure ①

This method involves measuring the input power and output power with and without the device whose attenuation is to be measured as shown in setup 1 and setup 2. The powers are measured in each set-up as P_1 and P_2 .

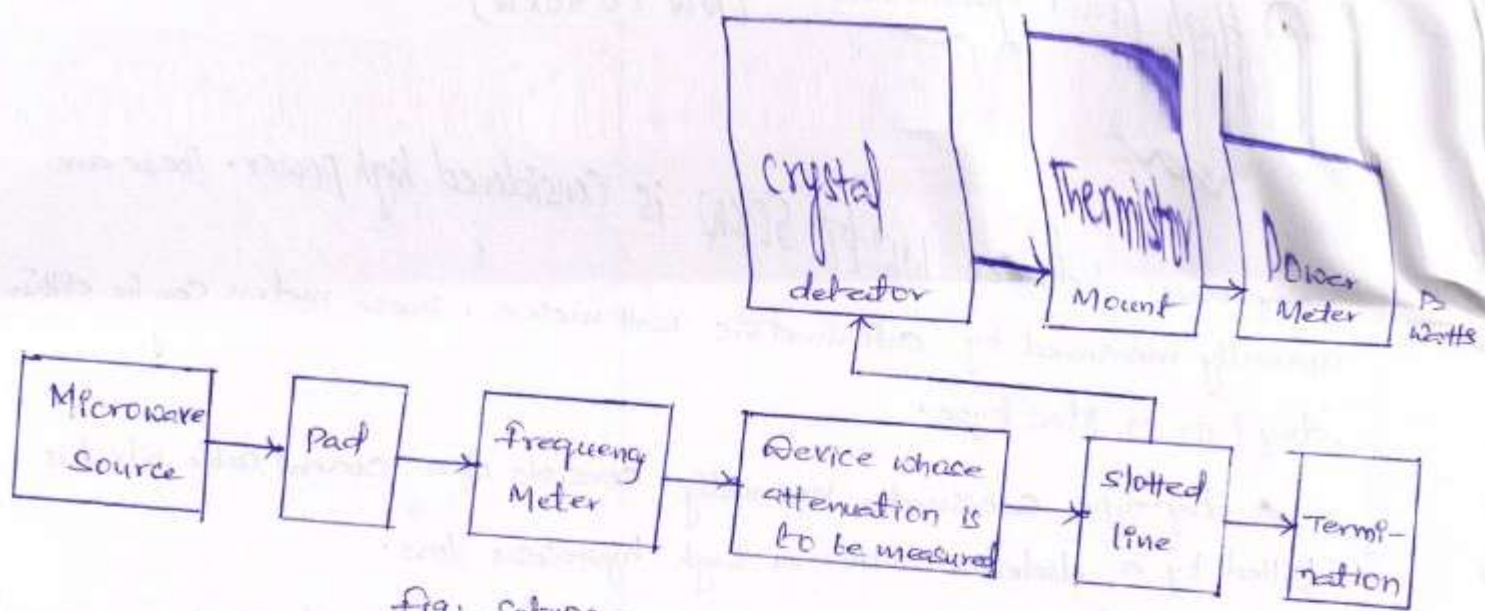
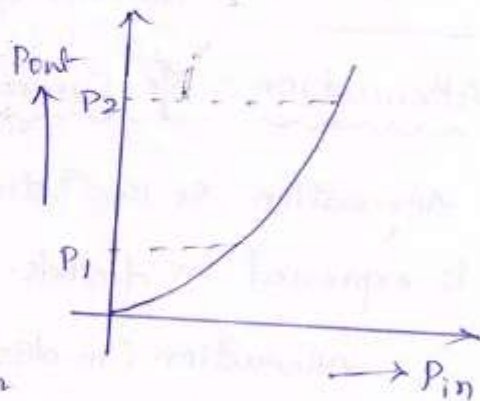


fig: Setup 2

The ratio of powers P_1/P_2 expressed in decibels gives the attenuation.
 → The drawback of this method is that the attenuation measured corresponds to two power positions on the power meter with a square law crystal detector characteristics as shown.

Due to non-linear characteristics of the two powers measured, the attenuation calculated will not be accurate. Particularly if the attenuation of the network is large and if the input power is low.



RF Substitution Method:

This method overcomes the drawback of power ratio method since here we measure attenuation at a single power position. The method consists of measuring the output power P by including the network whose attenuation is to be measured in setup 1.

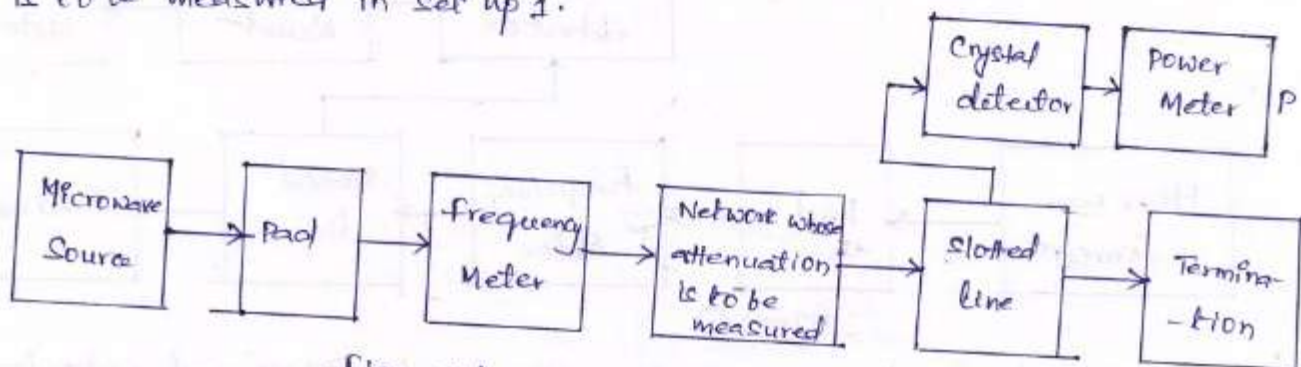


fig: Setup 1

MWE
U-V-5
Under this condition the attenuation read on the precision attenuator would give attenuation of the network directly.

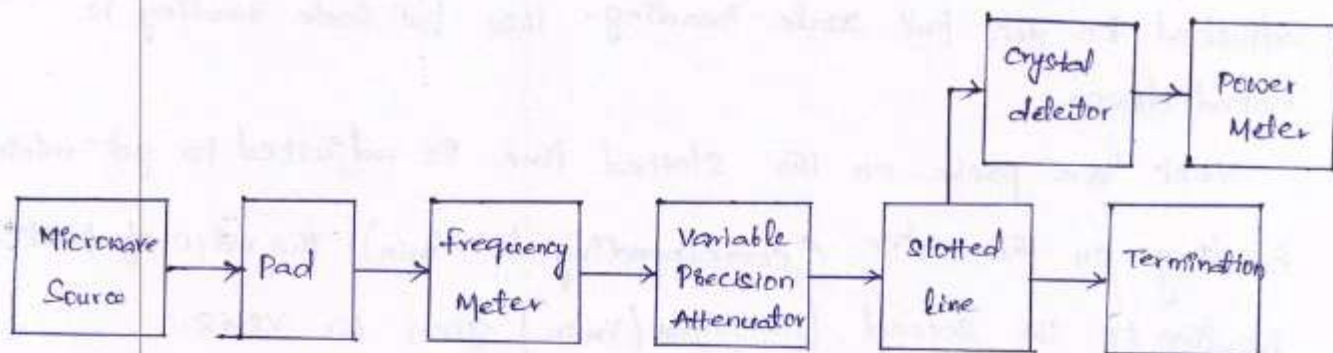
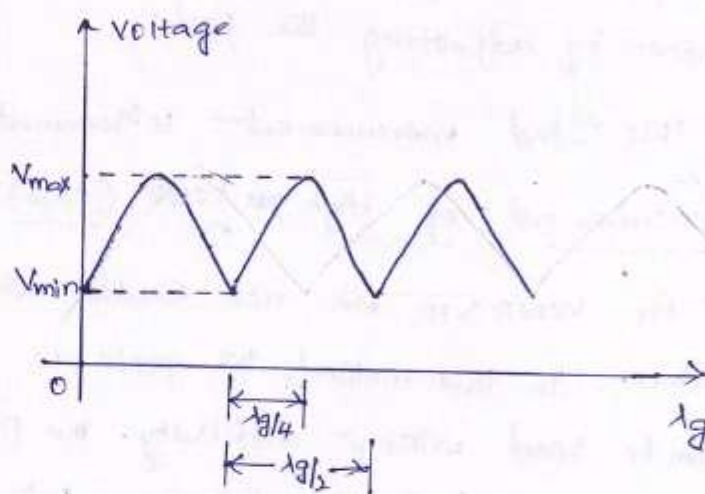


Fig: Setup 2

Measurement of Voltage Standing Wave Ratio [VSWR]:-

Any mismatched load leads to reflected waves resulting in standing waves along the length of the line.

The ratio of maximum to minimum voltage gives the VSWR.

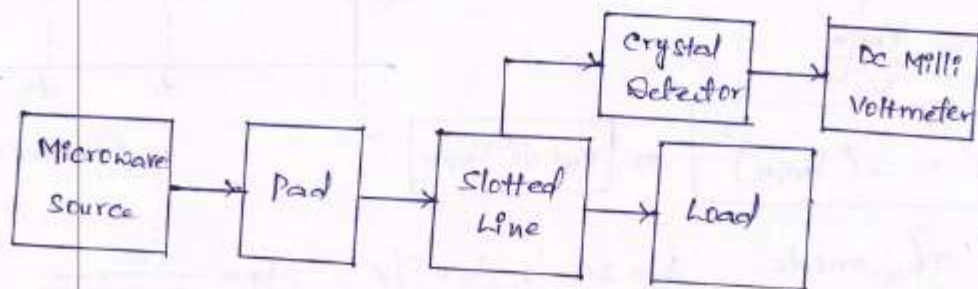


$$\text{i.e. } S = \frac{V_{max}}{V_{min}} = \frac{1+\Gamma}{1-\Gamma}$$

where $\Gamma \rightarrow$ Reflection Coefficient = $\frac{P_{reflected}}{P_{incident}}$

S varies from 1 to ∞ & Γ varies from 0 to 1.

⑨ Measurement of Low VSWR ($S < 10$):-



VSWR values not exceeding 10 are very easily measured with the VSWR meter calibrated directly.

The measurement basically consists of simply adjusting the attenuator to give an adequate reading on the meter, which is a DC millivolt meter.

The probe on the slotted waveguide is moved to get maximum reading on the meter (corresponding to V_{max}). The attenuation is now adjusted to get full scale reading. This full scale reading is noted down.

Next the probe on the slotted line is adjusted to get minimum reading on the meter (corresponding to V_{min}). The ratio of first reading to the second [i.e. V_{max}/V_{min}] gives the VSWR.

The meter itself can be calibrated in terms of VSWR. In this case, the probe carriage is moved to give maximum deflection on the VSWR meter by adjusting the pad.

→ This kind of measurement is inaccurate when $VSWR > 10$.

⑥ Measurement of High $VSWR$ ($S > 10$):-

For $VSWR > 10$, we use Double Minimum Method to measure VSWR values. In this method, the probe is inserted to a depth where the minimum can be read without difficulty. The probe is then moved to a point where the power is twice the minimum. Let this position be denoted by d_1 . The probe is then moved to twice the power point on the other side of the minimum (say d_2) as shown.

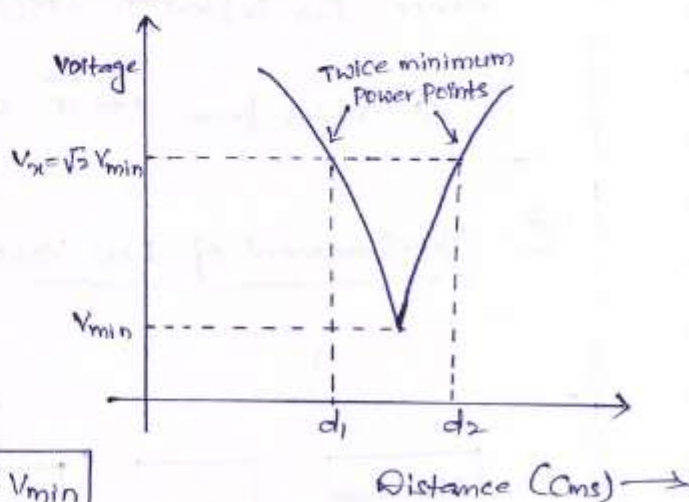
We get,

$$P_{min} \propto V_{min}^2$$

$$2 P_{min} \propto V_x^2$$

$$\therefore \frac{1}{2} = \frac{V_{min}^2}{V_x^2}$$

$$(or) \boxed{V_x^2 = 2(V_{min})^2} \quad or \quad \boxed{V_x = \sqrt{2} V_{min}}$$



further for TE₁₀ mode, $k_c = 2a$, $\lambda_0 = c/f$; $\lambda_g = \frac{\lambda_0}{\sqrt{1 - (\lambda_0/\lambda_c)^2}}$

then VSWR can be calculated using the formula. i.e.

$$\boxed{VSWR = \frac{\lambda_g}{\pi(d_2 - d_1)}}$$

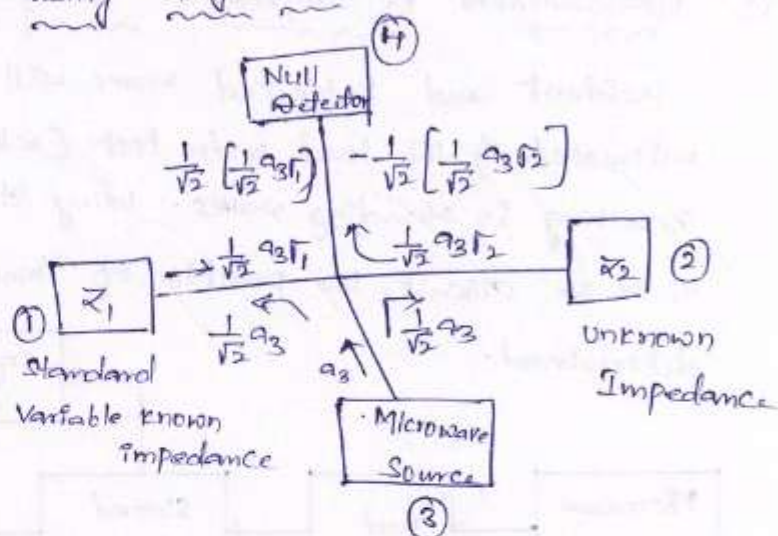
Measurement of Impedance:-

Impedance at Microwave frequencies can be measured by the following three methods.

- using Magic Tee
- using slotted line
- using Reflectometer

(a) Impedance measurement using Magic Tee:-

Microwave Source is connected to arm (3) and null detector is connected to arm (4). The unknown impedance is connected in arm (2) and a standard variable known impedance connected to arm (1).



From the properties of Magic Tee, the power from microwave source (a_3) gets divided equally between arms

$\Gamma_1 \rightarrow$ Reflection Coefficient of Z_1

$\Gamma_2 \rightarrow$ Reflection Coefficient of Z_2 .

(1) and (2) $\left[\frac{a_3}{\sqrt{2}} \right]$ (to the unknown impedance and standard variable impedances).

These impedances are not equal to characteristic impedance Z_0 and hence there will be reflections from arms (1) and (2).

If Γ_1 and Γ_2 are the reflection coefficients, then powers $\frac{\Gamma_1 a_3}{\sqrt{2}}$ and $\frac{\Gamma_2 a_3}{\sqrt{2}}$ enter the Magic Tee junction from arms (1) and (2) as shown. The resultant wave into arm (4) i.e. the null detector.

The net wave reaching the null detector is

$$= \frac{1}{\sqrt{2}} \left[\frac{1}{\sqrt{2}} a_3 \Gamma_1 \right] - \frac{1}{\sqrt{2}} \left[\frac{1}{\sqrt{2}} a_3 \Gamma_2 \right] = \frac{1}{2} a_3 [\Gamma_1 - \Gamma_2]$$

for perfect balancing of the bridge (null detector), the net wave reaching null detector is zero.

$$\text{i.e. } \frac{1}{2} a_3 (\Gamma_1 - \Gamma_2) = 0 \Rightarrow \Gamma_1 - \Gamma_2 = 0 \Rightarrow \Gamma_1 = \Gamma_2$$

$$\Rightarrow \cancel{Z_1 Z_2} + \cancel{Z_1 Z_0} - \cancel{Z_0 Z_2} - \cancel{Z_0^2} = \cancel{Z_1 Z_2} - \cancel{Z_1 Z_0} + \cancel{Z_0 Z_2} - \cancel{Z_0^2} \Rightarrow \cancel{Z_1 Z_0} = \cancel{Z_0 Z_2}$$

$$\frac{Z_1 - Z_0}{Z_1 + Z_0} = \frac{Z_2 - Z_0}{Z_2 + Z_0}$$

$$\therefore \boxed{Z_1 = Z_2}$$

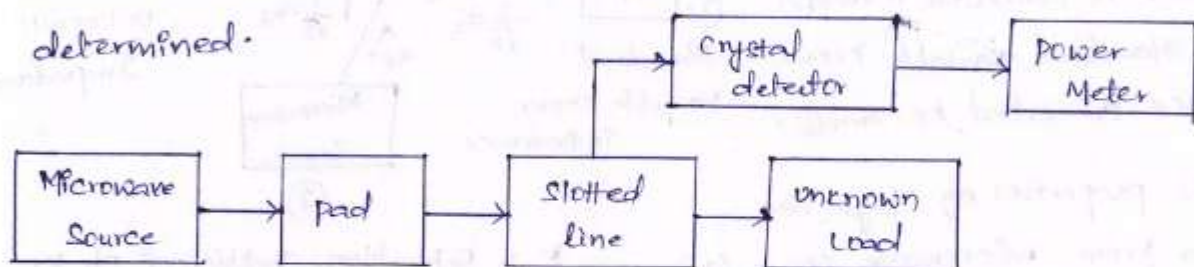
i.e. $R_1 + jX_1 = R_2 + jX_2$

or $\boxed{R_1 = R_2 \text{ and } X_1 = X_2}$

Thus the unknown impedance can be measured by adjusting the standard variable impedance till the bridge is balanced and both impedances become equal.

⑥ Measurement of Impedance using slotted Line:

Incident and Reflected waves will be present proportional to the mismatch of the load under test (whose impedance is to be measured) resulting in standing waves. Using slotted waveguide and with the load Z_L in the circuit, the position of V_{max} and V_{min} can be accurately determined.



Now the load Z_L is replaced by a short circuit and the shift in minimum is measured. If the minimum is shifted to the left, then the impedance is inductive and if it shifts to the right it is capacitive. The unknown impedance can be obtained by the data recorded and a Smith chart. Both impedance and Reflection coefficient can be obtained in magnitude and phase.

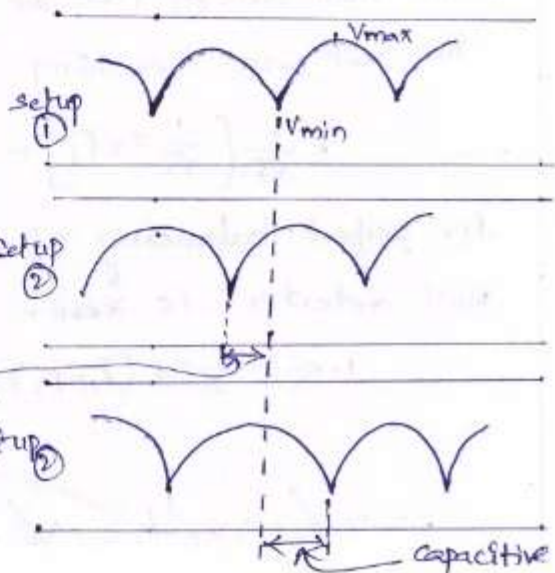
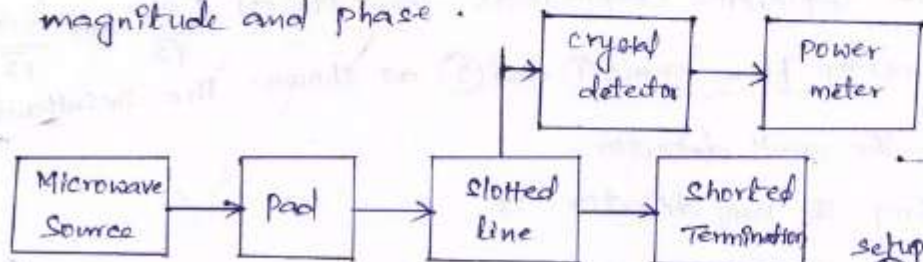
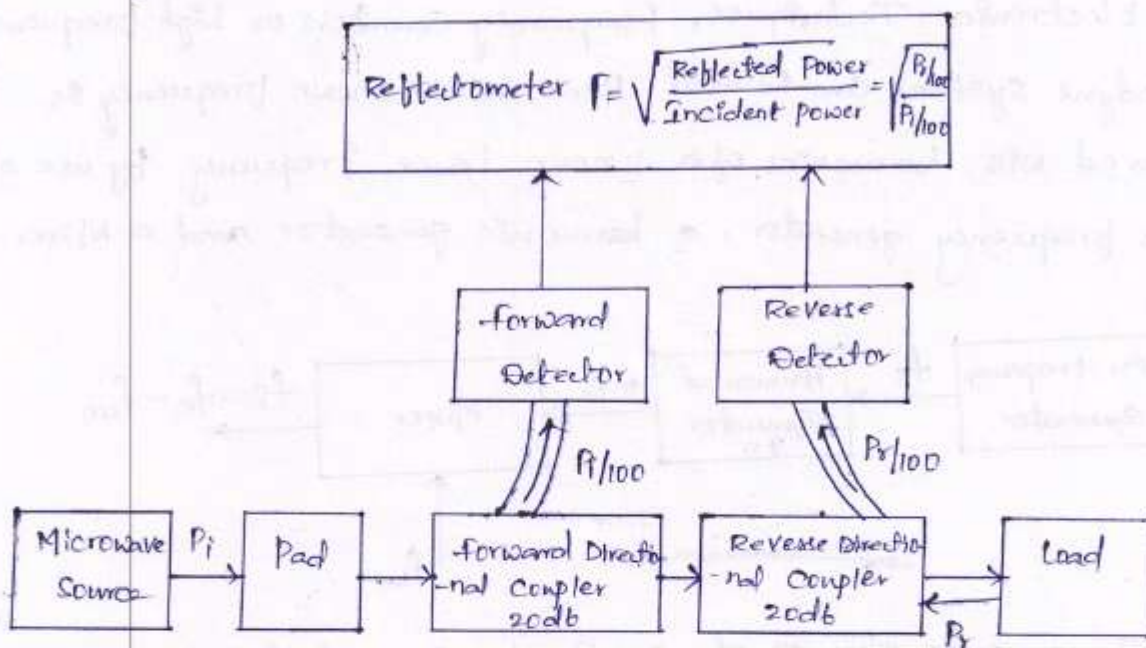


Fig: output standing waves of setup ① and ②.

© Measurement of Impedance using Reflectometer:-
 The Reflectometer indicates magnitude of Impedance but not the phase, angle, whereas slotted line waveguide measurement gives both.
 A typical set-up for Reflectometer technique is shown in figure.



In this, two directional Couplers are used to sample the incident power P_i and the reflected power P_r from the Load. Both the directional couplers are identical (except their direction). The magnitude of the reflection coefficient Γ can be directly obtained on the Reflectometer from which impedance can be calculated.

from the Reflectometer reading we have,

$$\Gamma = \sqrt{P_r/P_i}$$

Since ;

$$S = \frac{1+\Gamma}{1-\Gamma} \quad \text{and} \quad \Gamma = \frac{Z-Z_0}{Z+Z_0}$$

where Z_0 is known wave impedance, $Z \rightarrow$ unknown impedance.

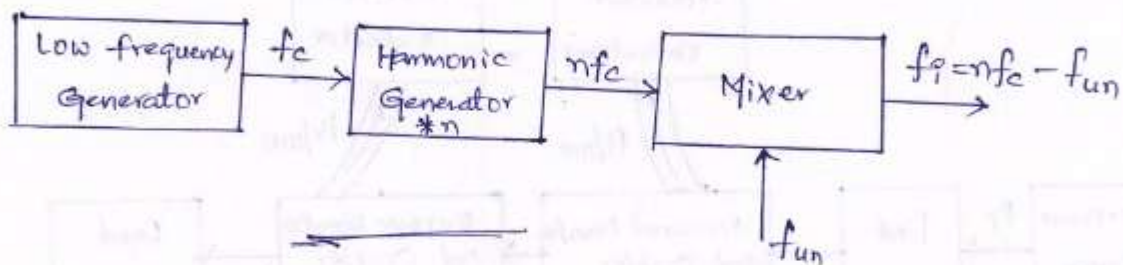
Due to Directional property of the couplers, there will be no interference between forward and Reverse waves. The input power P_i kept to a low level by means of pad.

The Reflectometer accuracy is greatest at low VSWR.

Frequency Measurement : —

Microwave frequency can be measured by either electronic or mechanical Techniques. Among these, Electronic Techniques are more accurate but expensive.

In Electronic Techniques, frequency counters or high frequency heterodyne systems can be used. Here the unknown frequency is compared with harmonics of a known lower frequency by use of a low frequency generator, a harmonic generator and a Mixer.



Measurement of Q of a Cavity Resonator : —

There are several methods for measuring the Q of a cavity resonator.

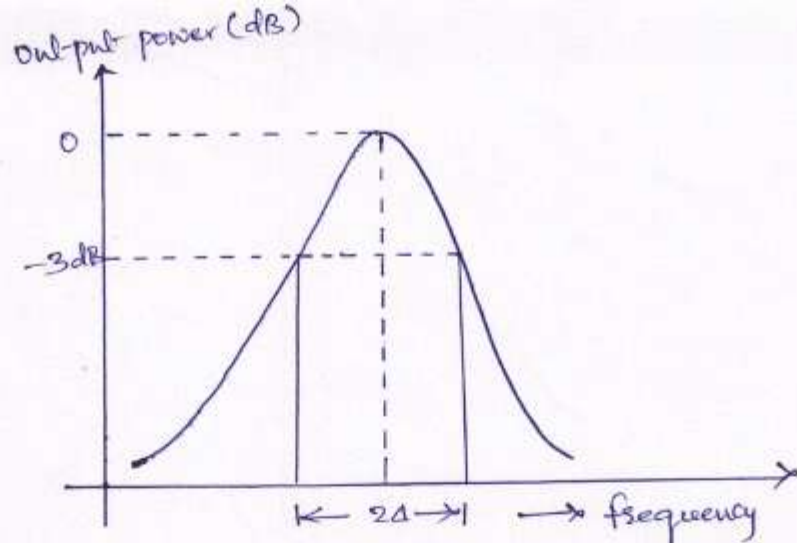
- (1) Transmission Method
- (2) Impedance Measurement
- (3) Transient delay or decrement Method

Among these, transmission method is the simplest.



The setup for transmission method of measuring Q is shown. In this method, the cavity resonator is used as a transmission device and the out-put signal is measured as a function of the frequency resulting in the Resonance Curve.

By varying the frequency of microwave source and keeping signal level constant, the out-put power is measured.



Alternately Cavity can be tuned by keeping both signal level and frequency constant and out-put power measured.

from the resonance curve half power Bandwidth (2Δ) can be obtained

i.e. $2\Delta = \pm \frac{1}{Q_L}$

where Q_L = Loaded value

or $Q_L = \pm \frac{1}{2\Delta} = \pm \frac{\omega}{2(\omega - \omega_0)}$

If the coupling between microwave source and cavity and that between detector and cavity are neglected, $Q_L = Q_0$ (unloaded Q).

However the accuracy of this method is poor, particularly in case of very high Q systems due to narrow band of operation.

