

MICROWAVE ENGINEERING (20A462T)

LECTURE

NOTES B.TECH

III-YEAR& II -SEM

Prepared by:

**Mrs. J. HIMABINDHU, Assistant Professor
Department of Electronics and Communication Engineering**



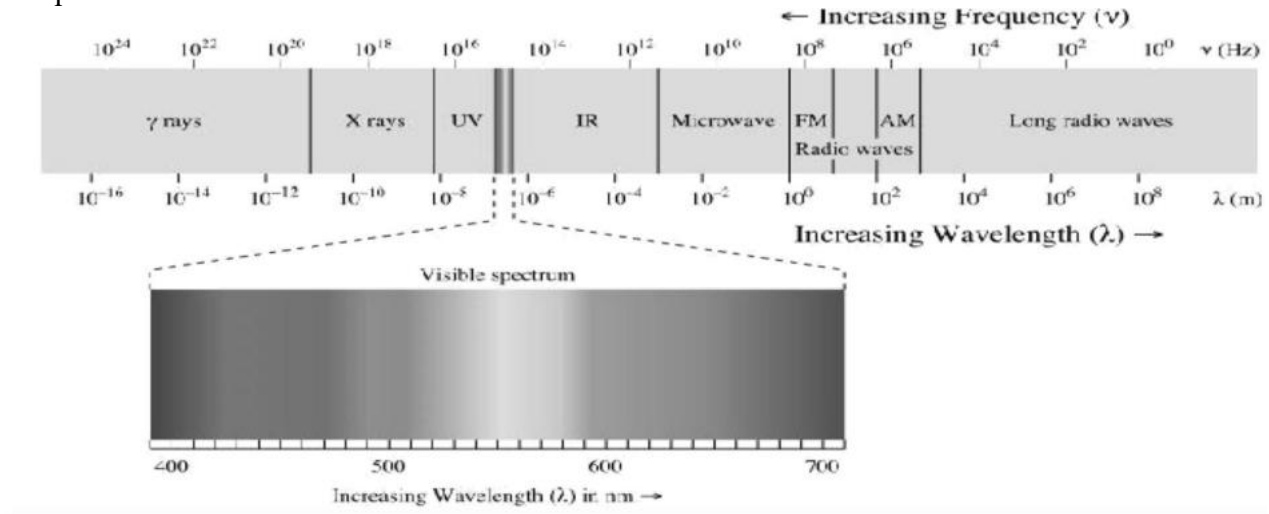
**ANNAMACHARYA INSTITUTE OF
TECHNOLOGY AND SCIENCES
RAJAMPET (An Autonomous Institution)**

UNIT I

Introduction to Microwave Engineering & Wave Guides

INTRODUCTION

Microwaves are electromagnetic waves with wavelengths ranging from 1 mm to 1 m, or frequencies between 300 MHz and 300 GHz.



Apparatus and techniques may be described qualitatively as "microwave" when the wavelengths of signals are roughly the same as the dimensions of the equipment, so that lumped-element circuit theory is inaccurate. As a consequence, practical microwave technique tends to move away from the discrete resistors, capacitors, and inductors used with lower frequency radio waves. Instead, distributed circuit elements and transmission-line theory are more useful methods for design, analysis. Open-wire and coaxial transmission lines give way to waveguides, and lumped-element tuned circuits are replaced by cavity resonators or resonant lines. Effects of reflection, polarization, scattering, diffraction, and atmospheric absorption usually associated with visible light are of practical significance in the study of microwave propagation. The same equations of electromagnetic theory apply at all frequencies.

While the name may suggest a micrometer wavelength, it is better understood as indicating wavelengths very much smaller than those used in radio broadcasting. The boundaries between far infrared light, terahertz radiation, microwaves, and ultra-high-frequency radio waves are fairly arbitrary and are used variously between different fields of study. The term microwave generally refers to "alternating current signals with frequencies between 300 MHz (3×10^8 Hz) and 300 GHz (3×10^{11} Hz)."[1] Both IEC standard 60050 and IEEE standard 100 define "microwave" frequencies starting at 1 GHz (30 cm wavelength).

Electromagnetic waves longer (lower frequency) than microwaves are called "radio waves". Electromagnetic radiation with shorter wavelengths may be called "millimeter waves", terahertz

radiation or even *T-rays*. Definitions differ for millimeter wave band, which the IEEE defines as 110 GHz to 300 GHz.

MICROWAVE FREQUENCY BANDS

The microwave spectrum is usually defined as electromagnetic energy ranging from approximately 1 GHz to 1000 GHz in frequency, but older usage includes lower frequencies. Most common applications are within the 1 to 40 GHz range. Microwave frequency bands, as defined by the Radio Society of Great Britain (RSGB), are shown in the table below: Microwave frequency bands

Designation	Frequency range
L band S	1 to 2 GHz
band C band	2 to 4 GHz
X band Ku	4 to 8 GHz
band K	8 to 12 GHz
band Ka	12 to 18 GHz
band	18 to 26.5 GHz
	26.5 to 40 GHz

Discovery

The existence of electromagnetic waves, of which microwaves are part of the frequency spectrum, was predicted by James Clerk Maxwell in 1864 from his equations. In 1888, Heinrich Hertz was the first to demonstrate the existence of electromagnetic waves by building an apparatus that produced and detected microwaves in the UHF region. The design necessarily used horse-and-buggy materials, including a horse trough, a wrought iron point spark, Leyden jars, and a length of zinc gutter whose parabolic cross-section worked as a reflection antenna. In 1894 J. C. Bose publicly demonstrated radio control of a bell using millimetre wavelengths, and conducted research into the propagation of microwaves.

Plot of the zenith atmospheric transmission on the summit of Mauna Kea throughout the entire gigahertz range of the electromagnetic spectrum at a precipitable water vapor level of 0.001 mm. (simulated)

Frequency range

The microwave range includes ultra-high frequency (UHF) (0.3–3 GHz), super high frequency (SHF) (3–30 GHz), and extremely high frequency (EHF) (30–300 GHz) signals.

Above 300 GHz, the absorption of electromagnetic radiation by Earth's atmosphere is so great that it is effectively opaque, until the atmosphere becomes transparent again in the so-called infrared and optical window frequency ranges.

Microwave Sources

Vacuum tube based devices operate on the ballistic motion of electrons in a vacuum under the influence of controlling electric or magnetic fields, and include the magnetron, klystron, travelling wave tube (TWT), and gyrotron. These devices work in the density modulated mode, rather than the current modulated mode. This means that they work on the basis of clumps of electrons flying ballistically through them, rather than using a continuous stream.

A maser is a device similar to a laser, except that it works at microwave frequencies.

Solid-state sources include the field-effect transistor, at least at lower frequencies, tunnel diodes and Gunn diodes

ADVANTAGES OF MICROWAVES

Communication

- Before the advent of fiber optic transmission, most long distance telephone calls were carried via microwave point-to-point links through sites like the AT&T Long Lines. Starting in the early 1950's, frequency division multiplex was used to send up to 5,400 telephone channels on each microwave radio channel, with as many as ten radio channels combined into one antenna for the *hop* to the next site, up to 70 km away.
 - Wireless LAN protocols, such as Bluetooth and the IEEE 802.11 specifications, also use microwaves in the 2.4 GHz ISM band, although 802.11a uses ISM band and U-NII frequencies in the 5 GHz range. Licensed long-range (up to about 25 km) Wireless Internet Access services can be found in many countries (but not the USA) in the 3.5–4.0 GHz range.
 - Metropolitan Area Networks: MAN protocols, such as WiMAX (Worldwide Interoperability for Microwave Access) based in the IEEE 802.16 specification. The IEEE 802.16 specification was designed to operate between 2 to 11 GHz. The commercial implementations are in the 2.3GHz, 2.5 GHz, 3.5 GHz and 5.8 GHz ranges.
 - Wide Area Mobile Broadband Wireless Access: MBWA protocols based on standards specifications such as IEEE 802.20 or ATIS/ANSI HC-SDMA (e.g. iBurst) are designed to operate between 1.6 and 2.3 GHz to give mobility and in-building penetration characteristics similar to mobile phones but with vastly greater spectral efficiency.
 - Cable TV and Internet access on coaxial cable as well as broadcast television use some of the lower microwave frequencies. Some mobile phone networks, like GSM, also use the lower microwave frequencies.
 - Microwave radio is used in broadcasting and telecommunication transmissions because, due to their short wavelength, highly directive antennas are smaller and therefore more practical than they would be at longer wavelengths (lower frequencies). There is also
-

more bandwidth in the microwave spectrum than in the rest of the radio spectrum; the usable bandwidth below 300 MHz is less than 300 MHz while many GHz can be used above 300 MHz. Typically, microwaves are used in television news to transmit a signal from a remote location to a television station from a specially equipped van.

Remote Sensing

- Radar uses microwave radiation to detect the range, speed, and other characteristics of remote objects. Development of radar was accelerated during World War II due to its great military utility. Now radar is widely used for applications such as air traffic control, navigation of ships, and speed limit enforcement.
- A Gunn diode oscillator and waveguide are used as a motion detector for automatic door openers (although these are being replaced by ultrasonic devices).
- Most radio astronomy uses microwaves.
- Microwave imaging; *see Photoacoustic imaging in biomedicine*

Navigation

Global Navigation Satellite Systems (GNSS) including the American Global Positioning System (GPS) and the Russian (GLONASS) broadcast navigational signals in various bands between about 1.2 GHz and 1.6 GHz.

Power

- A microwave oven passes (non-ionizing) microwave radiation (at a frequency near 2.45 GHz) through food, causing dielectric heating by absorption of energy in the water, fats and sugar contained in the food. Microwave ovens became common kitchen appliances in Western countries in the late 1970s, following development of inexpensive cavity magnetrons.
- Microwave heating is used in industrial processes for drying and curing products.
- Many semiconductor processing techniques use microwaves to generate plasma for such purposes as reactive ion etching and plasma-enhanced chemical vapor deposition (PECVD).
- Microwaves can be used to transmit power over long distances, and post-World War II research was done to examine possibilities. NASA worked in the 1970s and early 1980s to research the possibilities of using Solar power satellite (SPS) systems with large solar arrays that would beam power down to the Earth's surface via microwaves.
- Less-than-lethal weaponry exists that uses millimeter waves to heat a thin layer of human skin to an intolerable temperature so as to make the targeted person move away. A two-second burst of the 95 GHz focused beam heats the skin to a temperature of 130 F (54 C)

at a depth of 1/64th of an inch (0.4 mm). The United States Air Force and Marines are currently using this type of Active Denial System.[2]

APPLICATIONS OF MICROWAVE ENGINEERING

- Antenna gain is proportional to the electrical size of the antenna. At higher frequencies, more antenna gain is therefore possible for a given physical antenna size, which has important consequences for implementing miniaturized microwave systems.
- More bandwidth can be realized at higher frequencies. Bandwidth is critically important because available frequency bands in the electromagnetic spectrum are being rapidly depleted.
- Microwave signals travel by line of sight are not bent by the ionosphere as are lower frequency signals and thus satellite and terrestrial communication links with very high capacities are possible.
- Effective reflection area (radar cross section) of a radar target is proportional to the target's electrical size. Thus generally microwave frequencies are preferred for radar systems.
- Various molecular, atomic, and nuclear resonances occur at microwave frequencies, creating a variety of unique applications in the areas of basic science, remote sensing, medical diagnostics and treatment, and heating methods.
- Today, the majority of applications of microwaves are related to radar and communication systems. Radar systems are used for detecting and locating targets and for air traffic control systems, missile tracking radars, automobile collision avoidance systems, weather prediction, motion detectors, and a wide variety of remote sensing systems.
- Microwave communication systems handle a large fraction of the world's international and other long haul telephone, data and television transmissions.
- Most of the currently developing wireless telecommunications systems, such as direct broadcast satellite (DBS) television, personal communication systems (PCSs), wireless local area networks (WLANS), cellular video (CV) systems, and global positioning satellite (GPS) systems rely heavily on microwave technology.

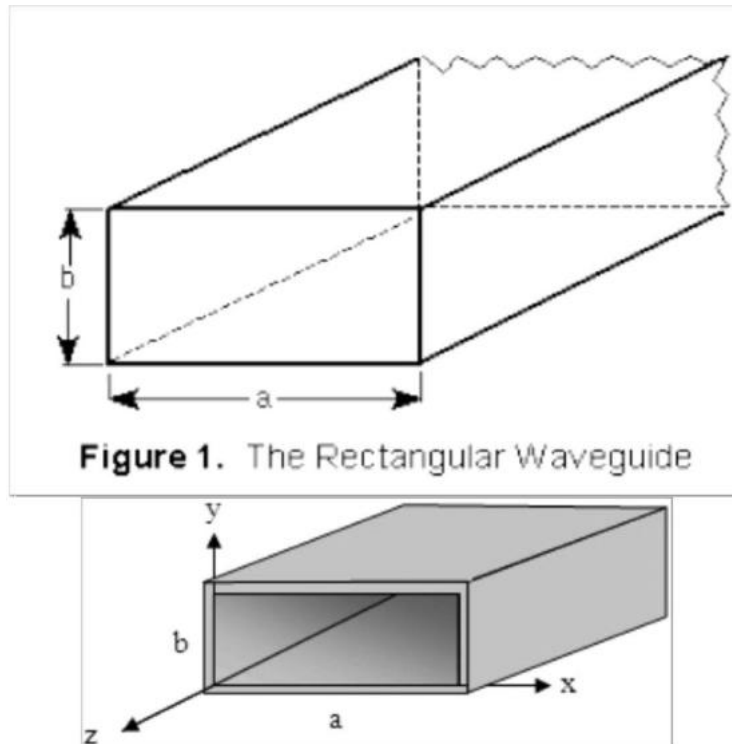
WAVEGUIDE :

The transmission line can't propagate high range of frequencies in GHz due to skin effect. Waveguides are generally used to propagate microwave signal and they always operate beyond certain frequency that is called "*cut off frequency*". so they behaves as high pass filter.

Types of waveguides: -

- (1)rectangular waveguide
- (2)cylindrical waveguide
- (3)elliptical waveguide
- (4)parallel waveguide

RECTANGULAR WAVEGUIDE :



Let us assume that the wave is travelling along z-axis and field variation along z-direction is equal to e^{-Yz} , where z=direction of propagation and Y= propagation constant.

Assume the waveguide is lossless ($\alpha=0$) and walls are perfect conductor ($\sigma=\infty$). According to maxwell's equation:

$$\nabla \times H = J + \frac{\partial D}{\partial t} \text{ and } \nabla \times E = -\frac{\partial B}{\partial t} \text{ . So } \nabla \times H = J + \frac{\partial D}{\partial t} \text{ --- (1. a) ,}$$

$$\nabla \cdot E = \frac{\rho}{\epsilon_0} \text{ --- (1. b)}$$

Expanding equation (1),

$$\begin{bmatrix} \frac{\partial H_z}{\partial y} - \frac{\partial H_y}{\partial z} \\ \frac{\partial H_x}{\partial z} - \frac{\partial H_z}{\partial x} \\ \frac{\partial H_y}{\partial x} - \frac{\partial H_x}{\partial y} \end{bmatrix} = J\omega \epsilon_0 \begin{bmatrix} E_x \\ E_y \\ E_z \end{bmatrix} \in [E_x A_x + E_y A_y + E_z A_z]$$

By equating coefficients of both sides we get,

$$\frac{\partial}{\partial y} H_z - \frac{\partial}{\partial z} H_y = J\omega \epsilon E_x \text{ -----2(a)}$$

$$-\frac{\partial}{\partial x} H_z + \frac{\partial}{\partial z} H_x = J\omega \epsilon E_y \text{ -----2(b)}$$

$$\frac{\partial}{\partial x} H_y - \frac{\partial}{\partial y} H_x = J\omega \epsilon E_z \text{ -----2(c)}$$

As the wave is travelling along z-direction and variation is along -Yz direction

$$\Rightarrow \frac{\partial}{\partial z} (e^{-\gamma z}) = -\gamma e^{-\gamma z}.$$

Comparing above equations, $\gamma = -Y$

So by putting this value of γ in equations 2(a,b,c), we will get

$$\frac{\partial}{\partial y} H_z + Y H_y = j\omega \epsilon E_x \text{ -----3(a)}$$

Similarly from relation $\nabla \times \mathbf{E} = -j\omega \mu \mathbf{H}$ and $-\gamma = -Y$, we will get --3(b)

$$\frac{\partial}{\partial y} E_z + Y E_y = -j\omega \mu H_x \text{ -----4(a)}$$

$$\frac{\partial}{\partial x} E_z + Y E_x = j\omega \mu H_y \text{ -----4(b)}$$

$$\frac{\partial}{\partial x} E_y - \frac{\partial}{\partial y} E_x = -j\omega \mu H_z \text{ -----4(c)}$$

From equation sets of (3),

we will get : $\gamma = -Y$

$$E_x = \frac{1}{j\omega \epsilon} \left[\frac{\partial}{\partial y} H_z + Y H_y \right] \text{ -----(5)}$$

From equation sets of (4), we will get : $\gamma = -Y$

$$E_x = \frac{1}{Y} \left[j\omega \mu H_y - \frac{\partial}{\partial x} E_z \right] \text{ -----(6)}$$

Equating equations (5) and (6), we will get

Similarly we will get by simplifying other equations

$$H_x = -\frac{j\omega\epsilon}{h^2} \frac{\partial}{\partial y} E_z - \frac{Y}{h^2} \frac{\partial}{\partial x} H_z \text{ ----- (8)}$$

$$E_x = -\frac{j\omega\mu}{h^2} \frac{\partial}{\partial y} H_z - \frac{Y}{h^2} \frac{\partial}{\partial x} E_z \text{ ----- (9)}$$

$$E_y = \frac{j\omega\mu}{h^2} \frac{\partial}{\partial x} H_z - \frac{Y}{h^2} \frac{\partial}{\partial y} E_z \text{ ----- (10)}$$

FIELD SOLUTIONS FOR TRANSVERSE MAGNETIC FIELD IN RECTANGULAR WAVEGUIDE :

$H_z=0$ and $E_z \neq 0$

$$E_{xyz} = (C_1 \cos k_x x + C_2 \sin k_x x)(C_3 \cos k_y y + C_4 \sin k_y y) e^{-\gamma z}$$

The values of $C_1, C_2, C_3, C_4, k_x, k_y$ are found out from boundary equations. as we know that the tangential component of E are constants across the boundary, then

$$E = \begin{cases} 0, & x = 0 \text{ and } x = a \\ 0, & y = 0 \text{ and } y = b \end{cases}$$

AT $x=0$ AND $y=0$;

$E = C_1 C_3 e^{-\gamma z} = 0$ but we know that $e^{-\gamma z} \neq 0$ wave is travelling along z -direction.

So either $C_1=0$ or $C_3=0$ otherwise $C_1C_3=0$

AT $x=0$ AND $y=b$;

$$E = C_1(C_3 \cos k_y b + C_4 \sin k_y b) e^{-\gamma z} = 0$$

$$\text{So } C_1 C_3 = 0$$

So equation (g) becomes

$$E_{XYZ} = (C_2 \sin k_x x \times C_4 \sin k_y y) e^{-\gamma z} \text{ -----(i)}$$

Hence for $x=0, E=0$

$$\text{So } (C_2 \sin k_x a \times C_4 \sin k_y y) e^{-\gamma z} = 0$$

$$\Rightarrow \sin k_x a = 0 \Rightarrow k_x = \frac{m\pi}{a}$$

In equation (i) for $y=b \Rightarrow E=0$;

$$\text{So } (C_2 \sin k_x x \times C_4 \sin k_y b) e^{-\gamma z} = 0$$

$$\Rightarrow \sin k_y b = 0 \Rightarrow k_y = \frac{n\pi}{b}$$

So finally solutions for TRANSVERSE MAGNETIC MODE is given by

$$E_z = C \left(\sin\left(\frac{m\pi}{a}\right)x \times \sin\left(\frac{n\pi}{b}\right)y \times e^{-\gamma z} \right)$$

Where $C_2 \times C_4 = C$

CUT-OFF FREQUENCY :

It is the minimum frequency after which propagation occurs inside the waveguide.

As we know that $\Rightarrow K_x^2 + k_y^2 + k_z^2 = k^2$

$$\Rightarrow K_x^2 + k_y^2 = k^2 - k_z^2$$

$$\Rightarrow K_x^2 + k_y^2 = k^2 + 2$$

As we know that $\gamma^2 = k^2 - K_x^2 - k_y^2$

So we will get that : $\Rightarrow K_x^2 + k_y^2 = k^2 + 2 \Rightarrow \gamma^2 = -2$

$$\text{So } \gamma = \sqrt{\left[\left(\frac{m\pi}{a}\right)^2 + \left(\frac{n\pi}{b}\right)^2 - \omega^2 \mu \epsilon\right]}$$

At $f=f_c$ or $\omega=\omega_c$, at cut off frequency propagation is about to start. So $\gamma=0$

where $m=n=0,1,2,3,\dots$

$$\begin{aligned} \text{At free space } f_c &= \frac{1}{2\pi\sqrt{\mu_0\epsilon_0}} \left(\left(\frac{m\pi}{a}\right)^2 + \left(\frac{n\pi}{b}\right)^2 \right)^{1/2} \\ \Rightarrow f_c &= \frac{c}{2} \left(\left(\frac{m\pi}{a}\right)^2 + \left(\frac{n\pi}{b}\right)^2 \right)^{1/2} \end{aligned}$$

CUT – OFF WAVELENGTH:

This is given by

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

DOMINANT MODE :

The mode having lowest cut-off frequency or highest cut-off wavelength is called DOMINANT MODE.

The mode can be TM₀₁, TM₁₀, TM₁₁, But for TM₁₀ and TM₀₁, wave can't exist.

Hence TM₁₁ has lowest cut-off frequency and is the DOMINANT MODE in case of all TM modes only.

PHASE CONSTANT :

As we know that

$$\gamma = \left(\left(\frac{m\pi}{a}\right)^2 + \left(\frac{n\pi}{b}\right)^2 - \omega^2\mu\epsilon \right)^{1/2}$$

So

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

This condition satisfies that only $\omega > \omega_c$

PHASE VELOCITY :

It is given by $V_p = \omega / \beta$

$$V_p = \frac{\omega}{\beta} = \frac{\omega}{\sqrt{\omega^2 \mu \epsilon - \omega^2 c^2 \mu \epsilon}}$$

$$V_p = 1 / [\sqrt{\omega \epsilon (1 - f c^2 / f^2)}]$$

GUIDE WAVELENGTH :

It is given by

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

$$\lambda_g = 1 / \sqrt{(f^2 \mu \epsilon - f c^2 \mu \epsilon)}$$

$$\lambda_g = \frac{c}{f} / (1 - f c^2 / f^2)$$

$$\lambda_g = \frac{\lambda_0}{[1 - (\frac{\lambda_0}{\lambda_c})^2]^{\frac{1}{2}}}$$

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

SOLUTIONS OF TRANSVERSE ELECTRIC MODE :

Here $E_z = 0$ and $H_z \neq 0$

$$H_z = (B_1 \cos k_x x + B_2 \sin k_x x)(B_3 \cos k_y y + B_4 \sin k_y y) -$$

$B_1, B_2, B_3, B_4, k_x, k_y$ are found from boundary conditions.

$$= 0 \quad = 0 \quad =$$

$$= 0 \quad = 0 \quad =$$

At $x=0$ and $y=0$;

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

Here

$$\frac{\partial}{\partial x} H_z = [B_1 * k_x * (-\sin k_x x) + B_2 * k_x * \cos k_x x] (B_3 \cos k_y y + B_4 \sin k_y y) e^{-\gamma z}$$

So

$$E_y = \frac{j\omega\mu}{h^2} [B_1 * k_x * (-\sin k_x x) + B_2 * k_x * \cos k_x x] (B_3 \cos k_y y + B_4 \sin k_y y) e^{-\gamma z}$$

At $x=0$, $\frac{\partial}{\partial x} H_z = 0$

$$0 = [B_2 * k_x] [(B_3 \cos k_y y + B_4 \sin k_y y) e^{-\gamma z}]$$

From this $B_2=0$

so

$$E_x = -\frac{j\omega\mu}{h^2} \frac{\partial}{\partial y} H_z - \frac{\gamma}{h^2} \frac{\partial}{\partial x} E_z \quad \left[\text{as } \frac{\gamma}{h^2} \frac{\partial}{\partial x} E_z = 0 \right];$$

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

Here

$$\frac{\partial}{\partial y} H_z = [B_1 (\cos k_x x) + B_2 \sin k_x x] (-B_3 * k_y * \sin k_y y + B_4 * k_y * \cos k_y y) e^{-\gamma z}$$

$$E_x = -\frac{j\omega\mu}{h^2} [B_1 (\cos k_x x) + B_2 \sin k_x x] (-B_3 * k_y * \sin k_y y + B_4 * k_y * \cos k_y y) e^{-\gamma z}$$

$$[B_1 (\cos k_x x) + B_2 \sin k_x x] (-B_3 * k_y * \sin k_y y + B_4 * k_y * \cos k_y y) e^{-\gamma z} = 0$$

At $y=0$, $\frac{\partial}{\partial y} H_z = 0$

so $\sin = 0 \Rightarrow =$

So the general TRANSVERSE ELECTRIC MODE solution is given by

$$H_z = B \left(\cos \frac{m\pi}{a} x \right) \left(\cos \frac{n\pi}{b} y \right) e^{-\gamma z}$$

Where $B = B_1 B_3$

CUT-OFF FREQUENCY :

The cut-off frequency is given as

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

DOMINANT MODE :

The mode having lowest cut-off frequency or highest cut-off wavelength is called DOMINANT MODE. here TE₀₀ where wave can't exist.

$$\text{So} \quad \begin{aligned} f_c(\text{TE}_{01}) &= c/2b \\ f_c(\text{TE}_{10}) &= c/2a \end{aligned}$$

for rectangular waveguide we know that $a > b$

so TE₁₀ is the dominant mode in all rectangular waveguide.

DEGENERATE MODE :

The modes having same cut-off frequency but different field equations are called degenerate modes.

WAVE IMPEDANCE :

Impedance offered by waveguide either in TE mode or TM mode when wave travels through, it is called wave impedance.

For TE mode

$$\eta_{TE} = \frac{\eta_i}{\sqrt{1 - \frac{f_c^2}{f^2}}}$$

And

$$\eta_{TM} = \eta_i * \sqrt{1 - \frac{f_c^2}{f^2}}$$

Where η =intrinsic impedance=377ohm=120 π

UNIT-II CYLINDRICAL WAVEGUIDES

A circular waveguide is a tubular, circular conductor. A plane wave propagating through a circular waveguide results in transverse electric (TE) or transverse magnetic field(TM) mode. Assume the medium is lossless($\alpha=0$) and the walls of the waveguide is perfect conductor($\sigma=\infty$). The field equations from MAXWELL'S EQUATIONS are:-

$$\nabla \times \mathbf{E} = -j\omega\mu\mathbf{H} \text{----- (1.a)}$$

Taking the first equation,

$$\nabla \times \mathbf{E} = -j\omega\mu\mathbf{H}$$

$$\frac{1}{\rho} \begin{vmatrix} A\rho & \rho A\phi & Az \\ \partial/\partial\rho & \partial/\partial\phi & \partial/\partial z \\ E\rho & \rho E\phi & Ez \end{vmatrix} = -j\omega\mu [H\rho A\rho + H\phi A\phi + HzAz]$$

Equating :-

Similarly expanding $\nabla \times H = j\omega \epsilon E$,

$$\frac{1}{\rho} \begin{vmatrix} A\rho & \rho A\varphi & Az \\ \partial/\partial\rho & \partial/\partial\varphi & \partial/\partial z \\ H\rho & \rho H\varphi & Hz \end{vmatrix} = j\omega\epsilon [E\rho A\rho + E\varphi A\varphi + EzAz]$$

$$\frac{1}{\rho} \left\{ \frac{\partial(Hz)}{\partial\varphi} - \frac{\partial(\rho H\varphi)}{\partial z} \right\} = j\omega\epsilon E\rho \quad \text{----- (3.a)}$$

$$\left\{ \frac{\partial(H\rho)}{\partial z} - \frac{\partial(Hz)}{\partial\rho} \right\} = -j\omega\epsilon E\varphi \quad \text{----- (3.b)}$$

$$\frac{1}{\rho} \left\{ \frac{\partial(\rho H\varphi)}{\partial\rho} - \frac{\partial(H\rho)}{\partial\varphi} \right\} = j\omega\epsilon Ez \quad \text{----- (3.c)}$$

Let us assume that the wave is propagating along z direction. So,

$$Hz = - ;$$

$$\Rightarrow \frac{\partial Hz}{\partial z} = -\gamma e^{-\gamma z}$$

$$\boxed{\frac{\partial}{\partial z} = -\gamma}$$

Putting in equation 2 and 3:

$$\left\{ \frac{\partial(Ez)}{\partial\varphi} + \gamma\rho E\varphi \right\} = -j\omega\mu\rho H\rho \quad \text{----- (4.a)}$$

$$\left\{ \gamma E\rho + \frac{\partial(Ez)}{\partial\rho} \right\} = j\omega\mu\rho H\varphi \quad \text{----- (4.b)}$$

$$\left\{ \frac{\partial(\rho E\varphi)}{\partial\rho} - \frac{\partial(E\rho)}{\partial\varphi} \right\} = -j\omega\mu\rho Hz \quad \text{----- (4.c)}$$

And

Now from eq(4.a) and eq(5.b),we get

$$\begin{aligned}
 H_\rho &= \frac{1}{j\omega\mu\rho} \left\{ \frac{\partial(E_z)}{\partial\varphi} + \gamma\rho E\varphi \right\}; \quad H_\rho = \frac{1}{\gamma} \left[-j\omega\epsilon E\varphi - \frac{\partial(H_z)}{\partial\rho} \right] \\
 \therefore \frac{1}{-j\omega\mu\rho} \left\{ \frac{\partial(E_z)}{\partial\varphi} + \gamma\rho E\varphi \right\} &= \frac{1}{\gamma} \left[-j\omega\epsilon E\varphi - \frac{\partial(H_z)}{\partial\rho} \right] \\
 \rightarrow \frac{1}{\rho} \frac{\partial(E_z)}{\partial\varphi} + \gamma E\varphi &= -\frac{\omega^2\mu\epsilon}{\gamma} E\varphi + \frac{j\omega\mu}{\gamma} \frac{\partial H_z}{\partial\rho} \\
 \Rightarrow \left(\frac{\gamma^2 + \omega^2\mu\epsilon}{\gamma} \right) E\varphi &= \frac{j\omega\mu}{\gamma} \frac{\partial H_z}{\partial\rho} - \frac{1}{\rho} \frac{\partial H_z}{\partial\varphi}
 \end{aligned}$$

Let $(\gamma^2 + \omega^2\mu\epsilon) = h^2 = Kc^2$;

For lossless medium $\alpha=0, \gamma=j\beta$;

Now the final equation for $E\varphi$ is

$$E\varphi = \frac{-j}{Kc^2} \left(\frac{\beta}{\rho} \frac{\partial E_z}{\partial\varphi} - \omega\mu \frac{\partial H_z}{\partial\rho} \right) \quad \text{----- (6.a)}$$

$$H_\varphi = \frac{-j}{Kc^2} \left(\omega\epsilon \frac{\partial E_z}{\partial\rho} + \frac{\beta}{\rho} \frac{\partial H_z}{\partial\varphi} \right) \quad \text{----- (6.b)}$$

$$E_\rho = \frac{-j}{Kc^2} \left(\frac{\omega\mu}{\rho} \frac{\partial H_z}{\partial\varphi} + \beta \frac{\partial E_z}{\partial\rho} \right) \quad \text{----- (6.c)}$$

$$H_\rho = \frac{j}{Kc^2} \left(\frac{\omega\epsilon}{\rho} \frac{\partial E_z}{\partial\varphi} - \beta \frac{\partial H_z}{\partial\rho} \right) \quad \text{----- (6.d)}$$

Equations (6.a),(6.b),(6.c),(6.d) are the field equations for cylindrical waveguides.

TE MODE IN CYLINDRICAL WAVEGUIDE :-

For TE mode, $E_z=0$, $H_z \neq 0$.

As the wave travels along z-direction, $e^{-\gamma z}$ is the solution along z-direction.

$$\text{As, } \gamma^2 + \omega^2 \mu \epsilon = h^2; \quad ;$$

$$-\beta^2 + \omega^2 \mu \epsilon = Kc^2;$$

$$-\beta^2 + K^2 = Kc^2 \quad (\text{as } K^2 = \omega^2 \mu \epsilon)$$

According to maxwell's equation, the laplacian of H_z :

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

$$\nabla^2 H_z + \omega^2 \mu \epsilon H_z = 0;$$

$$\nabla^2 H_z + K^2 H_z = 0$$

Expanding the above equation, we get:

$$\frac{\partial^2 H_z}{\partial \rho^2} + \frac{1}{\rho} \frac{\partial H_z}{\partial \rho} + \frac{1}{\rho^2} \frac{\partial^2 H_z}{\partial \phi^2} + \frac{\partial^2 H_z}{\partial z^2} + K^2 H_z = 0;$$

Now

$$H_z = e^{-\gamma z}; \quad \frac{\partial^2 H_z}{\partial z^2} = (-\gamma)^2 e^{-\gamma z}; \quad \frac{\partial^2 H_z}{\partial z^2} = -\beta^2 e^{-\gamma z};$$

$$\frac{\partial^2 H_z}{\partial z^2} = -\beta^2 H_z \quad \Rightarrow \quad \frac{\partial^2}{\partial z^2} = -\beta^2;$$

Putting this value in the above equations, we get:-

$$\frac{\partial^2 H_z}{\partial \rho^2} + \frac{1}{\rho} \frac{\partial H_z}{\partial \rho} + \frac{1}{\rho^2} \frac{\partial^2 H_z}{\partial \phi^2} - \beta^2 H_z + K^2 H_z = 0;$$

$$\frac{\partial^2 H_z}{\partial \rho^2} + \frac{1}{\rho} \frac{\partial H_z}{\partial \rho} + \frac{1}{\rho^2} \frac{\partial^2 H_z}{\partial \phi^2} + Kc^2 H_z = 0; \quad [\text{as } -\beta^2 + K^2 = Kc^2]$$

$$\frac{\partial^2 H_z}{\partial \rho^2} + \frac{1}{\rho} \frac{\partial H_z}{\partial \rho} + Kc^2 H_z = -\frac{1}{\rho^2} \frac{\partial^2 H_z}{\partial \phi^2};$$

The partial differential with respect to ρ and ϕ in the above equation are equal only when the individuals are constant (Let it be Ko^2).

Solutions to the above differential equation is:-

$$H_z = B_1 \sin(K_o \phi) + B_2 \cos(K_o \phi) \text{ ---- } \{\text{solution along } \phi \text{ direction}\}.$$

Now,

$$\frac{\partial^2 H_z}{\partial \rho^2} + \frac{1}{\rho} \frac{\partial H_z}{\partial \rho} + (K_c^2 H_z - K_c^2) = 0$$

This equation is similar to BESSEL'S EQUATION, so the solution of this equation is $H_z = C_n J_n(K_c \rho)$ ----- {solution along ρ -direction}

Hence,

$$H_z = H_z(\rho) H_z(\phi) e^{-\gamma z};$$

So, the final solution is,

$$H_z = C_n J_n(K_c \rho) [B_1 \sin(K_o \phi) + B_2 \cos(K_o \phi)] e^{-\gamma z}$$

Applying boundary conditions:-

$$\text{At } \rho = a, E_\phi = 0 \Rightarrow \partial H_z / \partial \rho = 0,$$

$$\Rightarrow J_n(K_c \rho) = 0$$

$$\Rightarrow J_n(K_c a) = 0$$

If the roots of above equation are defined as P_{mn} , then

$$K_c = P_{mn}/a;$$

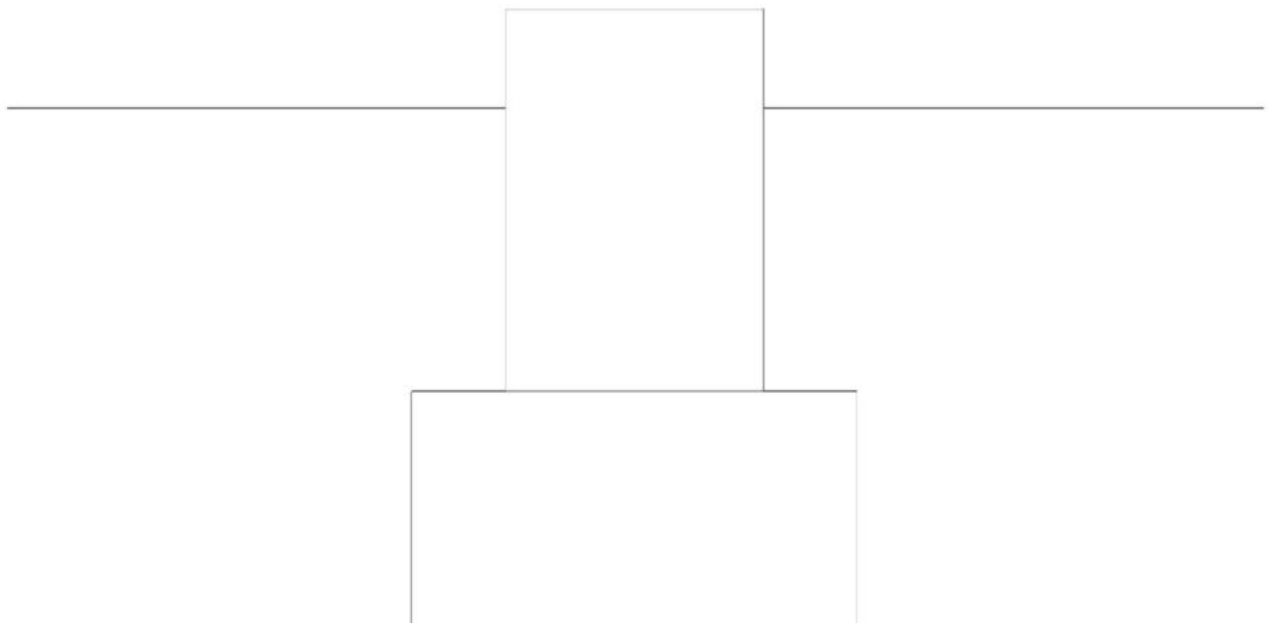
$$H_z = C_n J_n\left(\frac{P_{mn}}{a} \rho\right) [B_1 \sin(K_o \phi) + B_2 \cos(K_o \phi)] e^{-\gamma z}$$

CUT-OFF FREQUENCY :- It is the minimum frequency after which the propagation occurs inside the cavity.

$$\therefore \gamma = 0;$$

But we know that

$$\gamma^2 + \omega^2 \mu \epsilon = K_c^2;$$



CUTOFF WAVELENGTH :-

$$\frac{c}{f_c} = \frac{1}{\frac{P_{mn}}{2\pi a \sqrt{\mu\epsilon}}} = \frac{2\pi a}{P_{mn}}$$

The experimental values of P_{mn} are:-

n \ m	1	2	3
0	3.832	7.016	10.173
1	1.841	5.331	8.536
2	3.054	6.706	9.969
3	3.054	6.706	9.970

As seen from this table, TE₁₁ mode has the lowest cut off frequency, hence **TE₁₁** is the dominating mode.

MICROWAVE COMPONENTS

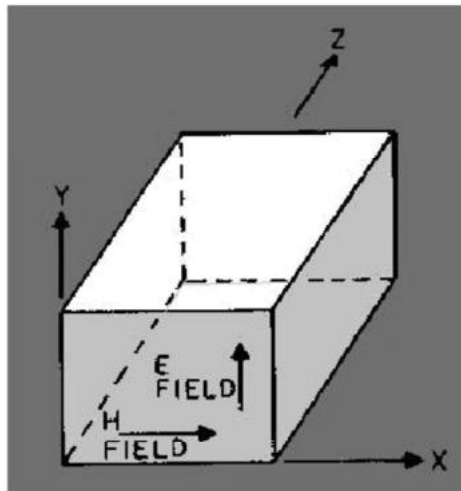
MICROWAVE RESONATOR:

- They are used in many applications such as oscillators, filters, frequency meters, tuned amplifiers and the like.
- A microwave resonator is a metallic enclosure that confines electromagnetic energy and stores it inside a cavity that determines its equivalent capacitance and inductance and from the energy dissipated due to finite conductive walls we can determine the equivalent resistance.
- The resonator has finite number of resonating modes and each mode corresponds to a particular resonant frequency.
- When the frequency of input signal equals to the resonant frequency, maximum amplitude of standing wave occurs and the peak energy stored in the electric and magnetic field are calculated.

RECTANGULAR WAVEGUIDE CAVITY RESONATOR:

Resonator can be constructed from closed section of waveguide by shorting both ends thus forming a closed box or cavity which store the electromagnetic energy and the power can be dissipated in the metallic walls as well as the dielectric medium

DIAGRAM:



The geometry of rectangular cavity resonator spreads as

$$0 \leq x \leq a;$$

$$0 \leq y \leq b;$$

$$0 \leq z \leq d$$

Hence the expression for cut-off frequency will be

This is the expression for resonant frequency of cavity resonator.

The mode having lowest resonant frequency is called DOMINANT MODE and for TE AND TM the dominant modes are TE-101 and TM-110 respectively.

QUALITY FACTOR OF CAVITY RESONATOR:

$$Q = 2\pi \times \frac{\text{maximum energy stored per cycle}}{\text{energy dissipated per cycle}}$$

FACTORS AFFECTING THE QUALITY FACTOR:

Quality factor depends upon 2 factors:

Lossy conducting walls

Lossy dielectric medium of a waveguide

1) LOSSY CONDUCTING WALL:

The Q-factor of a cavity with lossy conducting walls but lossless dielectric medium i.e. $\sigma_c \neq \infty$ and $\sigma = 0$

Then $Q_c = (2\omega_0 W_e / P_c)$

Where ω_0 -resonant frequency

W_e -stored electrical energy

P_c -power loss in conducting walls

2) LOSSY DIELECTRIC MEDIUM:

The Q-factor of a cavity with lossy dielectric medium but lossless conducting walls i.e. $\sigma_c = \infty$ and $\sigma \neq 0$

$$Q_d = 2\omega_0 W_e / P_d = (1 / \tan \delta)$$

Where $\tan \delta =$

P_d = power loss in dielectric medium

When both the conducting walls and the dielectric medium are lossy in nature then

Total power loss = $P_c + P_d$

$$\frac{1}{Q_{\text{total}}} = \frac{1}{Q_c} + \frac{1}{Q_d}$$

or $Q_{\text{total}} = 1 / (\frac{1}{Q_c} + \frac{1}{Q_d})$

MICROSTRIP

Microstrip transmission line is a kind of "high grade" printed circuit construction, consisting of a track of copper or other conductor on an insulating substrate. There is a "backplane" on the other side of the insulating substrate, formed from similar conductor. A picture (37kB). Looked at end on, there is a "hot" conductor which is the track on the top, and a "return" conductor which is the backplane on the bottom. Microstrip is therefore a variant of 2-wire transmission line. If one solves the electromagnetic equations to find the field distributions, one finds very nearly a completely TEM (transverse electromagnetic) pattern. This means that there are only a few regions in which there is a component of electric or magnetic field in the direction of wave propagation. There is a picture of these field patterns (incomplete) in T C Edwards "Foundations for Microstrip Circuit Design" edition 2 page 45. See the booklist for further bibliographic details. The field pattern is commonly referred to as a Quasi TEM pattern.

Under some conditions one has to take account of the effects due to longitudinal fields. An example is geometrical dispersion, where different wave frequencies travel at different phase velocities, and the group and phase velocities are different. The quasi TEM pattern arises because of the interface between the dielectric substrate and the surrounding air. The electric field lines have a discontinuity in direction at the interface.

The boundary conditions for electric field are that the normal component (ie the component at right angles to the surface) of the electric field times the dielectric constant is continuous across the boundary; thus in the dielectric which may have dielectric constant 10, the electric field suddenly drops to 1/10 of its value in air. On the other hand, the tangential component (parallel to the interface) of the electric field is continuous across the boundary.

In general then we observe a sudden change of direction of electric field lines at the interface, which gives rise to a longitudinal magnetic field component from the second Maxwell's equation, $\text{curl } E = - dB/dt$. Since some of the electric energy is stored in the air and some in the dielectric, the effective dielectric constant for the waves on the transmission line will lie somewhere between that of the air and that of the dielectric. Typically the effective dielectric

constant will be 50-85% of the substrate dielectric constant. As an example, in (notionally) air spaced microstrip the velocity of waves would be $c = 3 * 10^8$ metres per second. We have to divide this figure by the square root of the effective dielectric constant to find the actual wave velocity for the real microstrip line.

At 10 GHz the wavelength on notionally air spaced microstrip is therefore 3 cms; however on a substrate with effective dielectric constant of 7 the wavelength is $3/(\sqrt{7}) = 1.13\text{cms}$.

WAVEGUIDE CUTOFF FREQUENCY:

-waveguide cutoff frequency is an essential parameter for any waveguide - it does not propagate signals below this frequency. It is easy to understand and calculate with our equations.

The cutoff frequency is the frequency below which the waveguide will not operate.

Accordingly it is essential that any signals required to pass through the waveguide do not extend close to or below the cutoff frequency.

The waveguide cutoff frequency is therefore one of the major specifications associated with any waveguide product.

Waveguide cutoff frequency basics

Waveguides will only carry or propagate signals above a certain frequency, known as the cut-off frequency. Below this the waveguide is not able to carry the signals. The cut-off frequency of the waveguide depends upon its dimensions. In view of the mechanical constraints this means that waveguides are only used for microwave frequencies. Although it is theoretically possible to build waveguides for lower frequencies the size would not make them viable to contain within normal dimensions and their cost would be prohibitive.

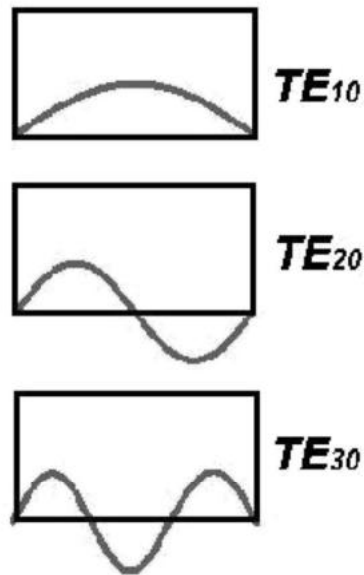
As a very rough guide to the dimensions required for a waveguide, the width of a waveguide needs to be of the same order of magnitude as the wavelength of the signal being carried. As a result, there is a number of standard sizes used for waveguides as detailed in another page of this tutorial. Also other forms of waveguide may be specifically designed to operate on a given band of frequencies

What is waveguide cutoff frequency? - the concept

Although the exact mechanics for the cutoff frequency of a waveguide vary according to whether it is rectangular, circular, etc, a good visualisation can be gained from the example of a rectangular waveguide. This is also the most widely used form.

Signals can progress along a waveguide using a number of modes. However the dominant mode is the one that has the lowest cutoff frequency. For a rectangular waveguide, this is the TE₁₀ mode.

The TE means transverse electric and indicates that the electric field is transverse to the direction of propagation.



TE modes for a rectangular waveguide

The diagram shows the electric field across the cross section of the waveguide. The lowest frequency that can be propagated by a mode equates to that where the wave can "fit into" the waveguide.

As seen by the diagram, it is possible for a number of modes to be active and this can cause significant problems and issues. All the modes propagate in slightly different ways and therefore if a number of modes are active, signal issues occur.

It is therefore best to select the waveguide dimensions so that, for a given input signal, only the energy of the dominant mode can be transmitted by the waveguide. For example: for a given frequency, the width of a rectangular guide may be too large: this would cause the TE₂₀ mode to propagate.

As a result, for low aspect ratio rectangular waveguides the TE₂₀ mode is the next higher order mode and it is harmonically related to the cutoff frequency of the TE₁₀ mode. This relationship and attenuation and propagation characteristics that determine the normal operating frequency range of rectangular waveguide.

Rectangular waveguide cutoff frequency

Although waveguides can support many modes of transmission, the one that is used, virtually exclusively is the TE₁₀ mode. If this assumption is made, then the calculation for the lower cutoff point becomes very simple:

$$f_c = \frac{c}{2a}$$

Where

f_c = rectangular waveguide cutoff frequency in Hz

c = speed of light within the waveguide in metres per second a =
the large internal dimension of the waveguide in metres

It is worth noting that the cutoff frequency is independent of the other dimension of the waveguide. This is because the major dimension governs the lowest frequency at which the waveguide can propagate a signal.

Circular waveguide cutoff frequency

the equation for a circular waveguide is a little more complicated (but not a lot).

$$f_c = \frac{1.8412 c}{2 \pi a}$$

where:

f_c = circular waveguide cutoff frequency in Hz

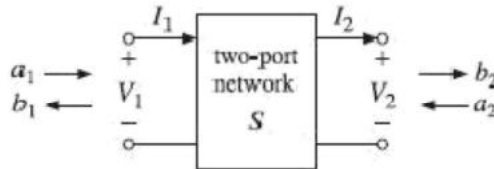
c = speed of light within the waveguide in metres per second a =
the internal radius for the circular waveguide in metres

Although it is possible to provide more generic waveguide cutoff frequency formulae, these ones are simple, easy to use and accommodate, by far the majority of calculations needed.

UNIT III MICROWAVE COMPONENTS

SCATTERING PARAMETERS

Linear two-port (and multi-port) networks are characterized by a number of equivalent circuit parameters, such as their transfer matrix, impedance matrix, admittance matrix, and scattering matrix. Fig. shows a typical two-port network.



The transfer matrix, also known as the ABCD matrix, relates the voltage and current at port 1 to those at port 2, whereas the impedance matrix relates the two voltages V_1, V_2 to the two currents I_1, I_2 :

$$\begin{bmatrix} V_1 \\ I_1 \end{bmatrix} = \begin{bmatrix} A & B \\ C & D \end{bmatrix} \begin{bmatrix} V_2 \\ I_2 \end{bmatrix} \quad (\text{transfer matrix})$$

$$\begin{bmatrix} V_1 \\ V_2 \end{bmatrix} = \begin{bmatrix} Z_{11} & Z_{12} \\ Z_{21} & Z_{22} \end{bmatrix} \begin{bmatrix} I_1 \\ -I_2 \end{bmatrix} \quad (\text{impedance matrix})$$

Thus, the transfer and impedance matrices are the 2×2 matrices:

$$T = \begin{bmatrix} A & B \\ C & D \end{bmatrix}, \quad Z = \begin{bmatrix} Z_{11} & Z_{12} \\ Z_{21} & Z_{22} \end{bmatrix}$$

The admittance matrix is simply the inverse of the impedance matrix, $Y = Z^{-1}$. The scattering matrix relates the outgoing waves b_1, b_2 to the incoming waves a_1, a_2 that are incident on the two-port:

$$\begin{bmatrix} b_1 \\ b_2 \end{bmatrix} = \begin{bmatrix} S_{11} & S_{12} \\ S_{21} & S_{22} \end{bmatrix} \begin{bmatrix} a_1 \\ a_2 \end{bmatrix}, \quad S = \begin{bmatrix} S_{11} & S_{12} \\ S_{21} & S_{22} \end{bmatrix} \quad (\text{scattering matrix})$$

The matrix elements $S_{11}, S_{12}, S_{21}, S_{22}$ are referred to as the scattering parameters or the S-parameters. The parameters S_{11}, S_{22} have the meaning of reflection coefficients, and S_{21}, S_{12} , the meaning of transmission coefficients.

THE SCATTERING MATRIX

The scattering matrix is defined as the relationship between the forward and backward moving waves. For a two-port network, like any other set of two-port parameters, the scattering matrix is a 2×2 matrix.

$$\begin{bmatrix} V_1^- \\ V_2^- \end{bmatrix} = \begin{bmatrix} S_{11} & S_{12} \\ S_{21} & S_{22} \end{bmatrix} \begin{bmatrix} V_1^+ \\ V_2^+ \end{bmatrix}$$

PROPERTIES OF SMATRIX:

In general the scattering parameters are complex quantities having the following Properties:

Property (1)

When any Z port is perfectly matched to the junction, then there are no reflections from that Thus S = 0. If all the ports are perfectly matched, then the leading diagonal II elements will all be zero.

Property (2)

Symmetric Property of S-matrix: If a microwave junction satisfies reciprocity condition and if there are no active devices, then S parameters are equal to their corresponding transposes.

$$i.e., S_{ij} = S_{ji}$$

Property (3)

Unitary property for a lossless junction - This property states that for any lossless network, the sum of the products of each term of anyone row or a nyone column of the [SJ matrix with its complex conjugate is unity

Property (4) :

Phase - Shift Property:

Complex S-parameters of a network are defined with respect to the positions of the port or reference planes. For a two-port network with unprimed reference planes 1 and 2 as shown in figure 4.6, the S-parameters have definite values.

COUPLING MECHANISMS:

PROBE, LOOP, APERTURE TYPES>

The three devices used to inject or remove energy from waveguides are PROBES, LOOPS, and SLOTS. Slots may also be called APERTURES or WINDOWS.

As previously discussed, when a small probe is inserted into a waveguide and supplied with microwave energy, it acts as a quarter-wave antenna. Current flows in the probe and sets up an E field such as the one shown in figure 1-39, view (A). The E lines detach themselves from the probe. When the probe is located at the point of highest efficiency, the E lines set up an E field of considerable intensity.

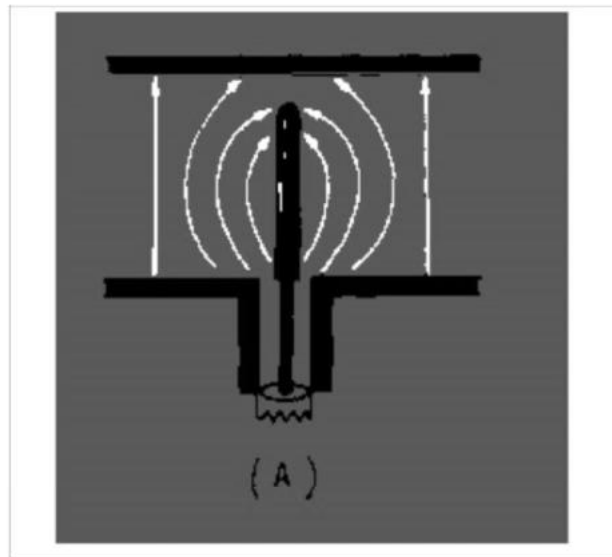


Fig: **Probe coupling** in a rectangular waveguide.

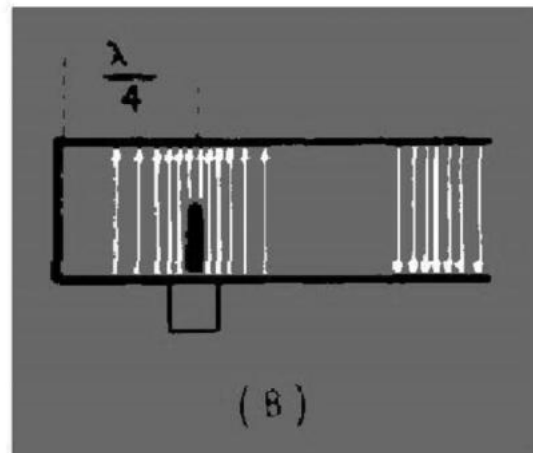


Figure. Probe coupling in a rectangular waveguide.

The most efficient place to locate the probe is in the center of the "a" wall, parallel to the "b" wall, and one quarter-wavelength from the shorted end of the waveguide, as shown in figure 1-39, views (B) and (C). This is the point at which the E field is maximum in the dominant mode. Therefore, energy transfer (coupling) is maximum at this point. Note that the quarter-wavelength spacing is at the frequency required to propagate the dominant mode.

In many applications a lesser degree of energy transfer, called loose coupling, is desirable. The amount of energy transfer can be reduced by decreasing the length of the probe, by moving it out of the center of the E field, or by shielding it. Where the degree of coupling must be varied frequently, the probe is made retractable so the length can be easily changed.

The size and shape of the probe determines its frequency, bandwidth, and power-handling capability. As the diameter of a probe increases, the bandwidth increases. A probe similar in shape to a door knob is capable of handling much higher power and a larger bandwidth than a conventional probe. The greater power-handling capability is directly related to the increased surface area. Two examples of broad-bandwidth probes are illustrated in figure 1-39, view (D). Removal of energy from a waveguide is simply a reversal of the injection process using the same type of probe.

Another way of injecting energy into a waveguide is by setting up an H field in the waveguide. This can be accomplished by inserting a small loop which carries a high current into the waveguide, as shown in figure 1-40, view (A). A magnetic field builds up around the loop and expands to fit the waveguide, as shown in view (B). If the frequency of the current in the loop is within the bandwidth of the waveguide, energy will be transferred to the waveguide.

For the most efficient coupling to the waveguide, the loop is inserted at one of several points where the magnetic field will be of greatest strength. Four of those points are shown in figure 1-40, view (C).

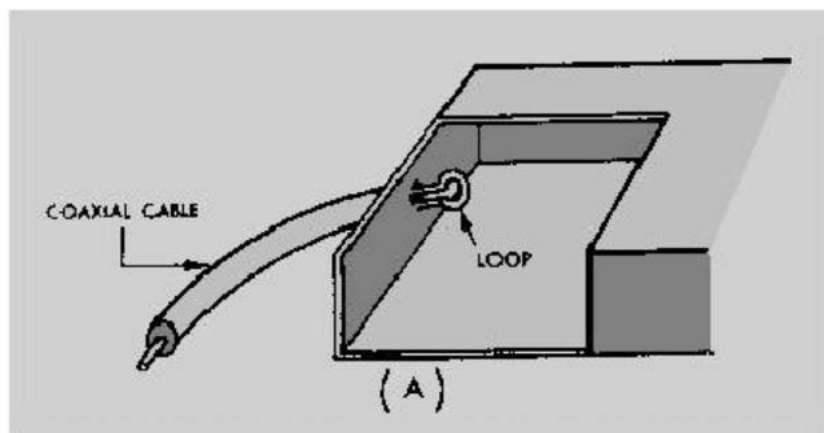
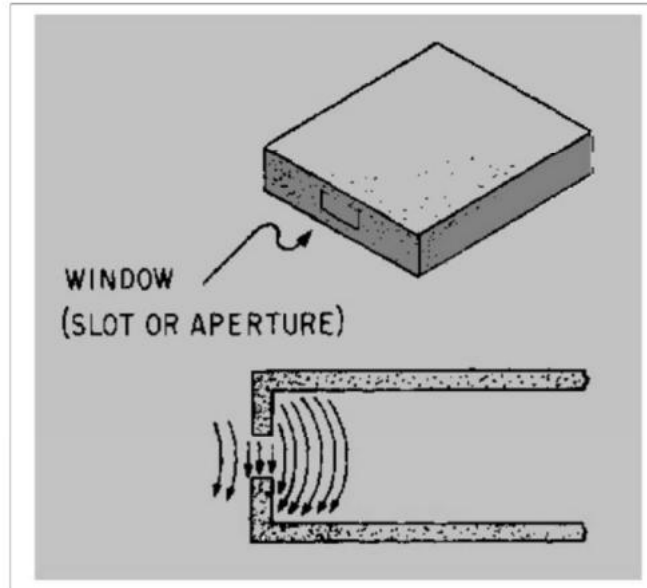


Figure - Loop coupling in a rectangular waveguide.

When less efficient coupling is desired, you can rotate or move the loop until it encircles a smaller number of H lines. When the diameter of the loop is increased, its power-handling capability also increases. The bandwidth can be increased by increasing the size of the wire used to make the loop.

When a loop is introduced into a waveguide in which an H field is present, a current is induced in the loop. When this condition exists, energy is removed from the waveguide.

Slots or apertures are sometimes used when very loose (inefficient) coupling is desired, as shown in figure 1-41. In this method energy enters through a small slot in the waveguide and the E field expands into the waveguide. The E lines expand first across the slot and then across the interior of the waveguide. Minimum reflections occur when energy is injected or removed if the size of the slot is properly proportioned to the frequency of the energy.



After learning how energy is coupled into and out of a waveguide with slots, you might think that leaving the end open is the most simple way of injecting or removing energy in a waveguide. This is not the case, however, because when energy leaves a waveguide, fields form around the end of the waveguide. These fields cause an impedance mismatch which, in turn, causes the development of standing waves and a drastic loss in efficiency.

Impedance matching using a waveguide iris

Irises are effectively obstructions within the waveguide that provide a capacitive or inductive element within the waveguide to provide the impedance matching.

The obstruction or waveguide iris is located in either the transverse plane of the magnetic or electric field. A waveguide iris places a shunt capacitance or inductance across the waveguide and it is directly proportional to the size of the waveguide iris.

An inductive waveguide iris is placed within the magnetic field, and a capacitive waveguide iris is placed within the electric field. These can be susceptible to breakdown under high power conditions - particularly the electric plane irises as they concentrate the electric field. Accordingly the use of a waveguide iris or screw / post can limit the power handling capacity.

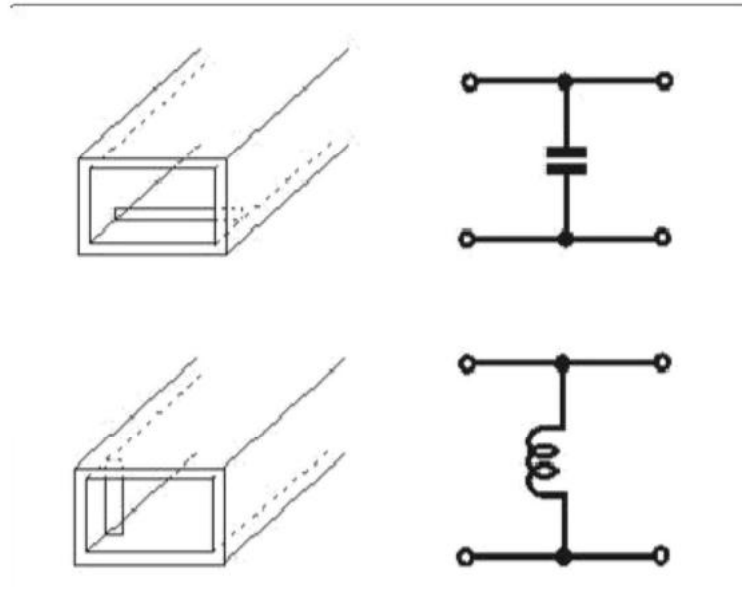


Fig: Impedance matching using a waveguide iris

The waveguide iris may either be on only one side of the waveguide, or there may be a waveguide iris on both sides to balance the system. A single waveguide iris is often referred to as an asymmetrical waveguide iris or diaphragm and one where there are two, one either side is known as a symmetrical waveguide iris.

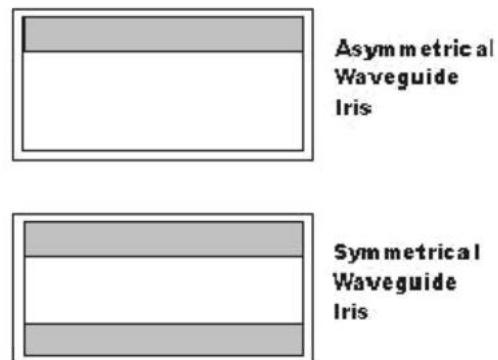


Fig :Symmetrical and asymmetrical waveguide iris implementations

A combination of both E and H plane waveguide irises can be used to provide both inductive and capacitive reactance. This forms a tuned circuit. At resonance, the iris acts as a high impedance shunt. Above or below resonance, the iris acts as a capacitive or inductive reactance.

Impedance matching using a waveguide post or screw

In addition to using a waveguide iris, post or screw can also be used to give a similar effect and thereby provide waveguide impedance matching.

The waveguide post or screw is made from a conductive material. To make the post or screw inductive, it should extend through the waveguide completely making contact with both top and bottom walls. For a capacitive reactance the post or screw should only extend part of the way through.

When a screw is used, the level can be varied to adjust the waveguide to the right conditions.

The diagram shows the electric field across the cross section of the waveguide. The lowest frequency that can be propagated by a mode equates to that where the wave can "fit into" the waveguide.

As seen by the diagram, it is possible for a number of modes to be active and this can cause significant problems and issues. All the modes propagate in slightly different ways and therefore if a number of modes are active, signal issues occur.

It is therefore best to select the waveguide dimensions so that, for a given input signal, only the energy of the dominant mode can be transmitted by the waveguide. For example: for a given frequency, the width of a rectangular guide may be too large: this would cause the TE₂₀ mode to propagate.

As a result, for low aspect ratio rectangular waveguides the TE₂₀ mode is the next higher order mode and it is harmonically related to the cutoff frequency of the TE₁₀ mode. This relationship and attenuation and propagation characteristics that determine the normal operating frequency range of rectangular waveguide

Waveguide Bends

Details of RF waveguide bends allowing changes in the direction of the transmission line waveguide E bend and waveguide H bend.

Waveguide is normally rigid, except for flexible waveguide, and therefore it is often necessary to direct the waveguide in a particular direction. Using waveguide bends and twists it is possible to arrange the waveguide into the positions required.

When using waveguide bends and waveguide twists, it is necessary to ensure the bending and twisting is accomplished in the correct manner otherwise the electric and magnetic fields will be unduly distorted and the signal will not propagate in the manner required causing loss and reflections. Accordingly waveguide bend and waveguide twist sections are manufactured specifically to allow the waveguide direction to be altered without unduly destroying the field patterns and introducing loss.

Types of waveguide bend

There are several ways in which waveguide bends can be accomplished. They may be used according to the applications and the requirements.

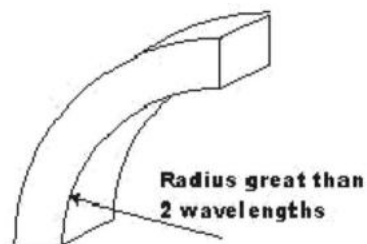
- Waveguide E bend
-

- Waveguide H bend
- Waveguide sharp E bend
- Waveguide sharp H bend

Each type of bend is achieved in a way that enables the signal to propagate correctly and with the minimum of disruption to the fields and hence to the overall signal.

Ideally the waveguide should be bent very gradually, but this is normally not viable and therefore specific waveguide bends are used.

Most proprietary waveguide bends are common angles - 90° waveguide bends are the most common by far.

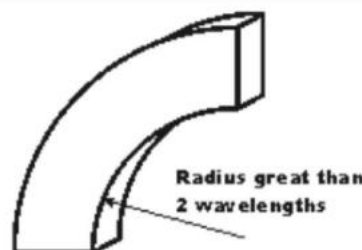


Waveguide E bend

o prevent reflections this waveguide bend must have a radius greater than two wavelengths.

Waveguide H bend

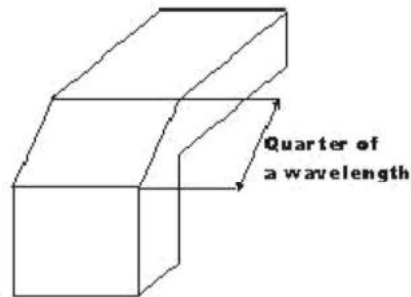
This form of waveguide bend is very similar to the E bend, except that it distorts the H or magnetic field. It creates the bend around the thinner side of the waveguide



Waveguide H bend

Waveguide sharp E bend

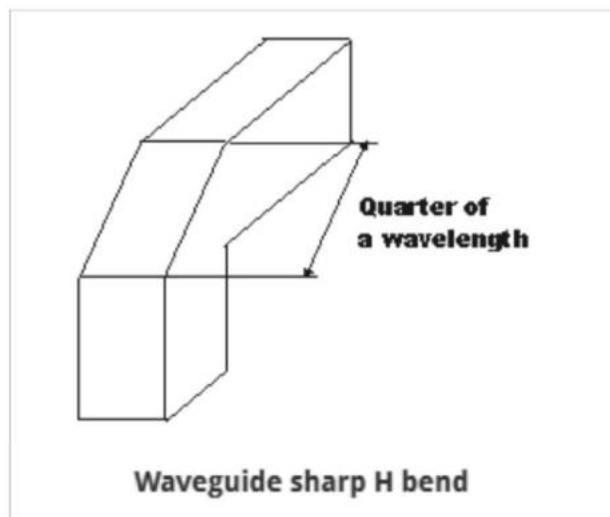
In some circumstances a much shorter or sharper bend may be required. This can be accomplished in a slightly different manner. The technique is to use a 45° bend in the waveguide. Effectively the signal is reflected, and using a 45° surface the reflections occur in such a way that the fields are left undisturbed, although the phase is inverted and in some applications this may need accounting for or correcting.



Waveguide sharp E bend

Waveguide sharp H bend

This form of waveguide bend is the same as the sharp E bend, except that the waveguide bend affects the H field rather than the E field.



Waveguide sharp H bend

WAVEGUIDE TWISTS

There are also instances where the waveguide may require twisting. This too, can be accomplished. A gradual twist in the waveguide is used to turn the polarisation of the waveguide and hence the waveform.

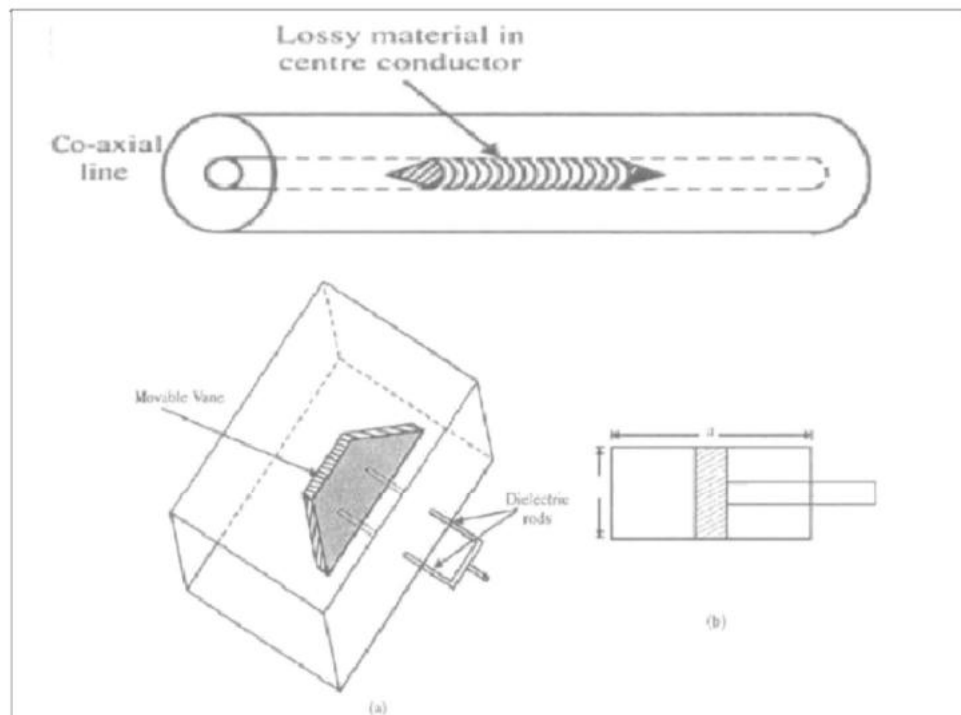
In order to prevent undue distortion on the waveform a 90° twist should be undertaken over a distance greater than two wavelengths of the frequency in use. If a complete inversion is required, e.g. for phasing requirements, the overall inversion or 180° twist should be undertaken over a four wavelength distance.

Waveguide bends and waveguide twists are very useful items to have when building a waveguide system. Using waveguide E bends and waveguide H bends and their snap bend counterparts allows the waveguide to be turned through the required angle to meet the mechanical constraints of the overall waveguide system. Waveguide twists are also useful in many applications to ensure the polarisation is correct.

ATTENUATORS:

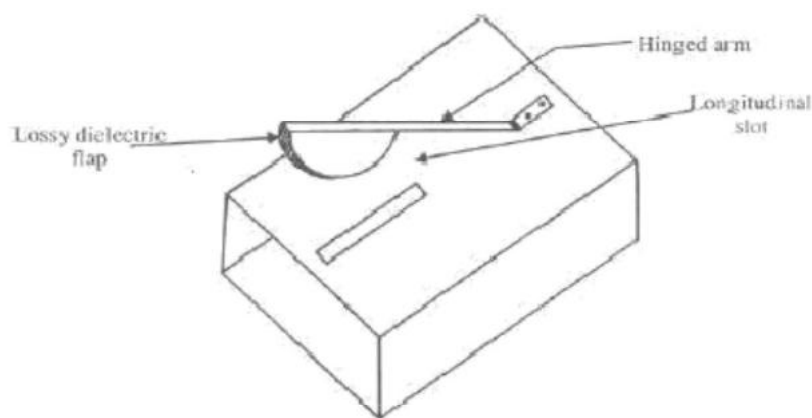
In order to control power levels in a microwave system by partially absorbing the transmitted microwave signal, attenuators are employed. Resistive films (dielectric glass slab coated with aquadag) are used in the design of both fixed and variable attenuators.

A co-axial fixed attenuator uses the dielectric lossy material inside the centre conductor of the co-axial line to absorb some of the centre conductor microwave power propagating through it. The dielectric rod decides the amount of attenuation introduced. The microwave power absorbed by the lossy material is dissipated as heat.



In waveguides, the dielectric slab coated with aduadag is placed at the centre of the waveguide parallel to the maximum E-field for dominant TE₁₀ mode. Induced current on the lossy material due to incoming microwave signal, results in power dissipation, leading to attenuation of the signal. The dielectric slab is tapered at both ends upto a length of more than half wavelength to reduce reflections as shown in figure 5.7. The dielectric slab may be made movable along the breadth of the waveguide by supporting it with two dielectric rods separated by an odd multiple of quarter guidewavelength and perpendicular to electric field.

When the slab is at the centre, then the attenuation is maximum (since the electric field is concentrated at the centre for TE₁₀ mode) and when it is moved towards one side-wall, the attenuation goes on decreasing thereby controlling the microwave power coming out of the other port.

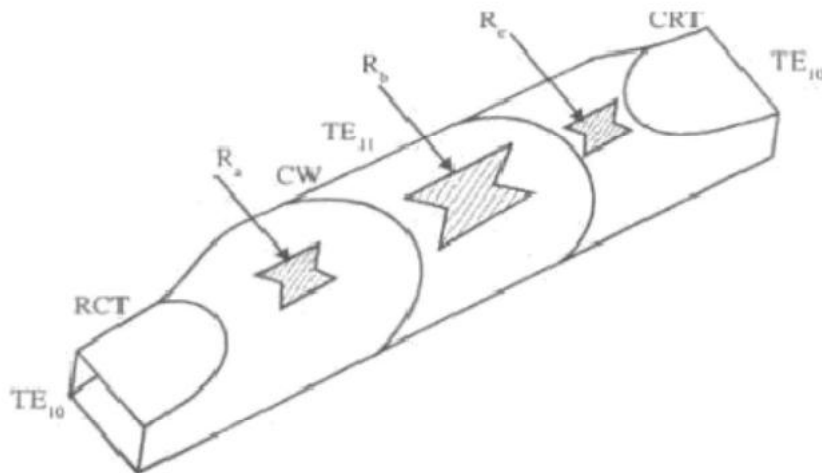


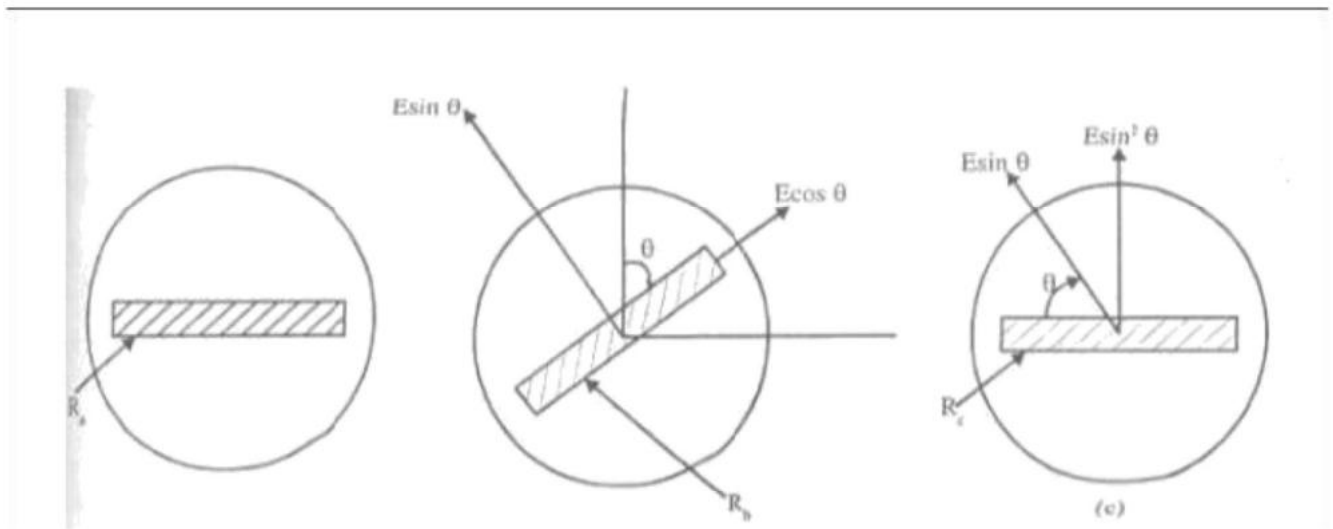
Above figure shows a flap attenuator which is also a variable attenuator. A semi-circular flap made of lossy dielectric is made to descend into the longitudinal slot cut at the center of the top wall of rectangular waveguide. When the flap is completely outside the slot, then the attenuation is zero and when it is completely inside, the attenuation is maximum. A maximum direction of 90 dB attenuation is possible with this attenuator with a VSWR of 1.05. The dielectric slab can be properly shaped according to convenience to get a linear variation of attenuation within the depth of insertion.

A precision type variable attenuator consists of a rectangular to circular transition (ReT), a piece of circular waveguide (CW) and a circular-to-rectangular transition (CRT) as shown in below figure. Resistive cards R_a , R_b and R_c are placed inside these sections as shown. The centre circular section containing the resistive card R_b can be precisely rotated by 360° with respect to the two fixed resistive cards. The induced current on the resistive card R due to the incident signal is dissipated as heat producing attenuation of the transmitted signal. TE mode in RCT is converted into TE in circular waveguide. The resistive cards R and R_a kept perpendicular to the electric field of TE₁₀ mode so that it does not absorb the energy. But any component parallel to its plane will be readily absorbed. Hence, pure TE mode is excited in circular waveguide section II.

If the resistive card in the centre section is kept at an angle θ relative to the E-field direction of the TE₁₀ mode, the component $E \cos(\theta)$ parallel to the card gets absorbed while the component $E \sin \theta$ is transmitted without attenuation. This component finally comes out as $E \sin 2\theta$

as shown in figure below.





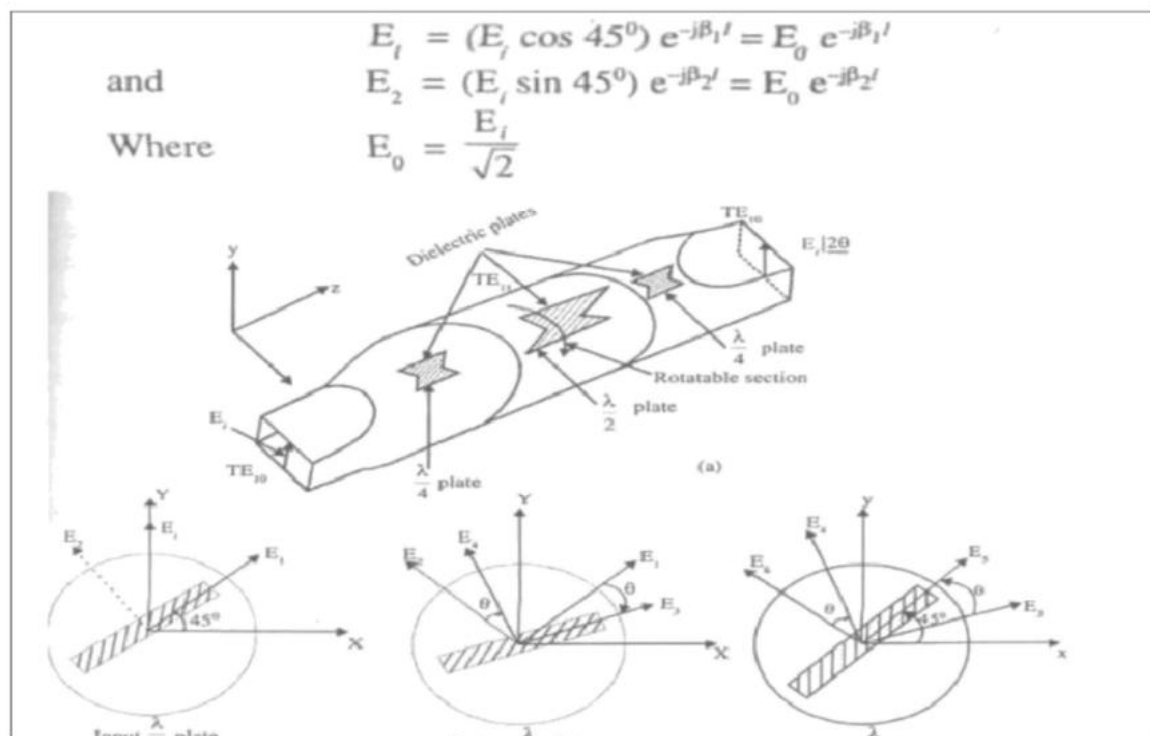
PHASE SHIFTERS:

A microwave phase shifter is a two port device which produces a variable shift in phase of the incoming microwave signal. A lossless dielectric slab when placed inside the rectangular waveguide produces a phase shift.

PRECISION PHASE SHIFTER

The rotary type of precision phase shifter is shown in figure below which consists of a circular waveguide containing a lossless dielectric plate of length $2l$ called "half-wave section", a section of rectangular-to-circular transition containing a lossless dielectric plate of length l , called "quarter-wave section", oriented at an angle of 45° to the broader wall of the rectangular waveguide and a circular-to-rectangular transition again containing a lossless dielectric plate of same length l (quarter wave section) oriented at an angle 45° .

The incident TE_{10} mode becomes TE_{11} mode in circular waveguide section. The half-wave section produces a phase shift equal to twice that produced by the quarter wave section. The dielectric plates are tapered at both ends to reduce reflections due to discontinuity.



When TE₁₀ mode is propagated through the input rectangular waveguide of the rectangular to circular transition, then it is converted into TE₁₁ in the circular waveguide section. Let E_i be the maximum electric field strength of this mode which is resolved into components, E_1 parallel to the plate and E_2 perpendicular to E_1 as shown in figure 5.12 (b). After propagation through the plate these components are given by

The length l is adjusted such that these two components E_1 and E_2 have equal amplitude but differing in phase by $= 90^\circ$

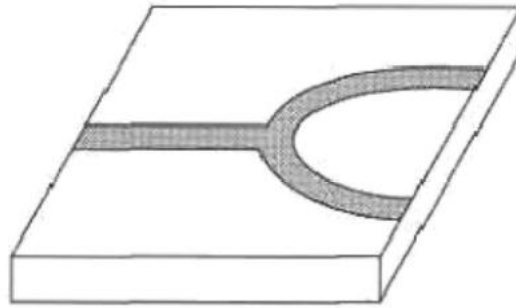
The quarter wave sections convert a linearly polarized TE₁₁ wave into a circularly polarized wave and vice-versa. After emerging out of the half-wave section, the electric field components parallel and perpendicular

to the half-wave plate. After emerging out of the half-wave section, the field components E_3 and E_4 as given in above equations, may again be resolved into two TE₁₁ modes, polarized parallel and perpendicular to the output quarter-wave plate. At the output end of this quarter-wave plate, the field components parallel and perpendicular to the quarter wave plate, by referring to figure above.

WAVEGUIDE MULTIPORT JUNCTIONS:

T-JUNCTION POWER DIVIDER USING WAVEGUIDE:

The T-junction power divider is a 3-port network that can be constructed either from a transmission line or from the waveguide depending upon the frequency of operation.

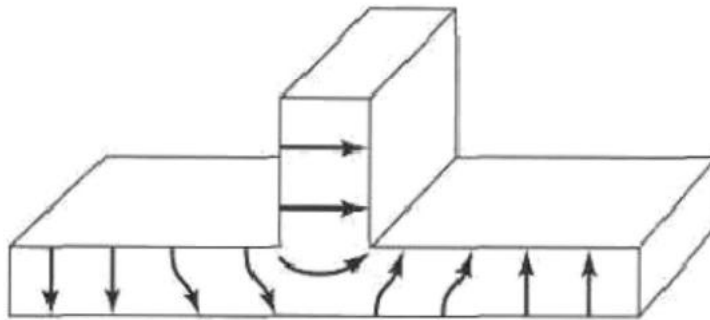


For very high frequency, power divider using waveguide is of 4 types

- E-Plane Tee
- H-Plane Tee
- E-H Plane Tee/Magic Tee
- Rat Race Tee

E-PLANE TEE:

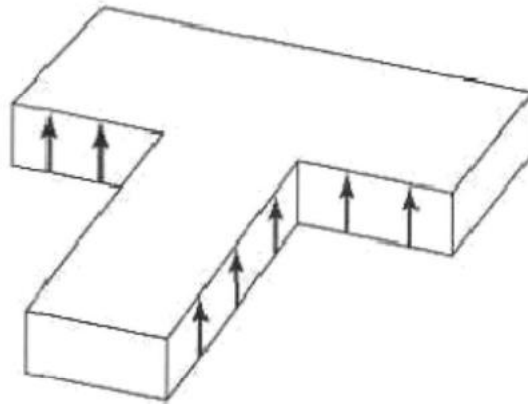
Diagram



- It can be constructed by making a rectangular slot along the wide dimension of the main waveguide and inserting another auxiliary waveguide along the direction so that it becomes a 3-port network.
- Port-1 and Port-2 are called collinear ports and Port-3 is called the E-arm.
- E-arm is parallel to the electric field of the main waveguide.
- If the wave is entering into the junction from E-arm it splits or gets divided into Port-1 and Port-2 with equal magnitude but opposite in phase
- If the wave is entering through Port-1 and Port-2 then the resulting field through Port-3 is proportional to the difference between the instantaneous field from Port-1 and Port-2

H-PLANE TEE:

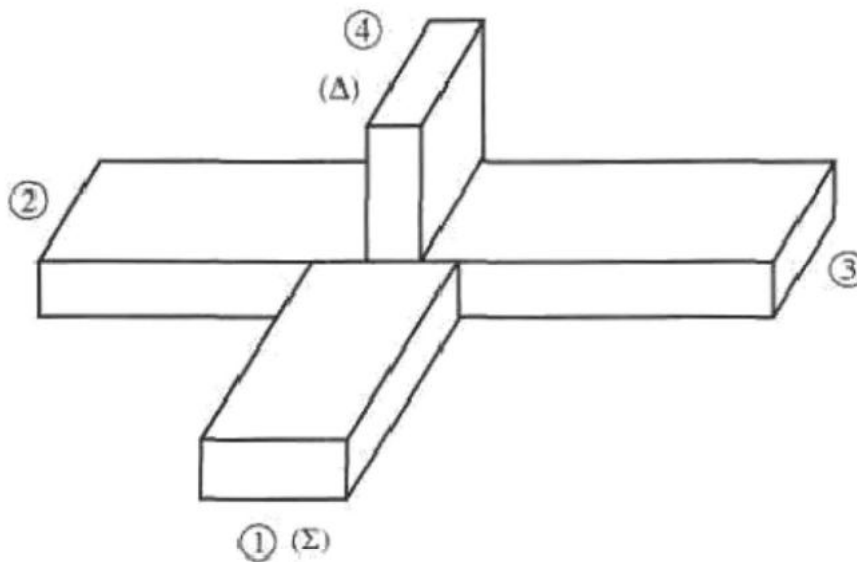
Diagram:



- An H-plane tee is formed by making a rectangular slot along the width of the main waveguide and inserting an auxiliary waveguide along this direction.
- In this case, the axis of the H-arm is parallel to the plane of the main waveguide.
- The wave entering through H-arm splits up through Port-1 and Port-2 with equal magnitude and same phase
- If the wave enters through Port-1 and Port-2 then the power through Port-3 is the phasor sum of those at Port-1 and Port-2.
- E-Plane tee is called PHASE DELAY and H-Plane tee is called PHASE ADVANCE.

E-H PLANE TEE/MAGIC TEE:

Diagram:

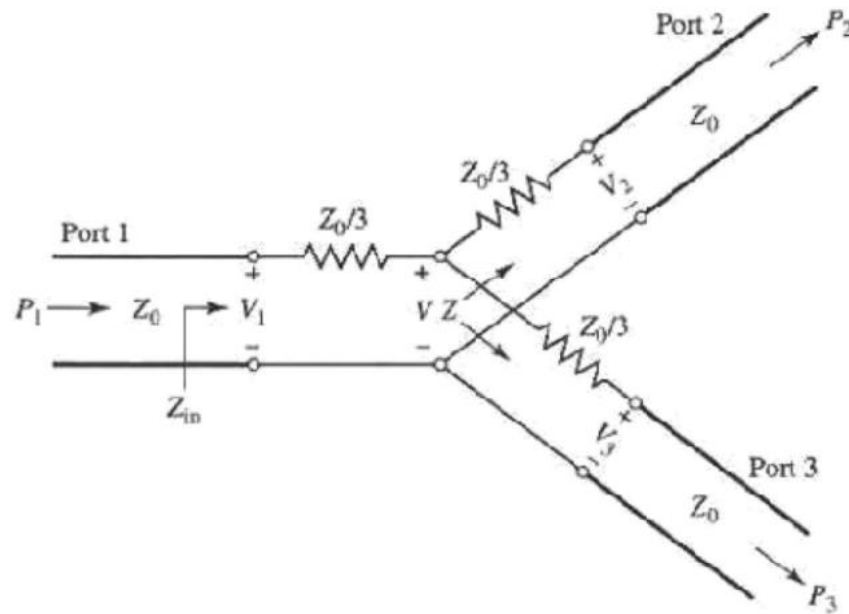


- 4 It is a combination of E-Plane tee and H-Plane tee.

- If two waves of equal magnitude and the same phase are fed into Port-1 and Port-2, the output will be zero at Port-3 and additive at Port-4.
- If a wave is fed into Port-4 (H-arm) then it will be divided equally between Port-1 and Port-2 of collinear arms (same in phase) and will not appear at Port-3 or E-arm.
- If a wave is fed in Port-3 then it will produce an output of equal magnitude and opposite phase at Port-1 and Port-2 and the output at Port-4 will be zero.
- If a wave is fed in any one of the collinear arms at Port-1 or Port-2, it will not appear in the other collinear arm because the E-arm causes a phase delay and the H-arm causes phase advance.

T-JUNCTION POWER DIVIDER USING TRANSMISSION LINE:

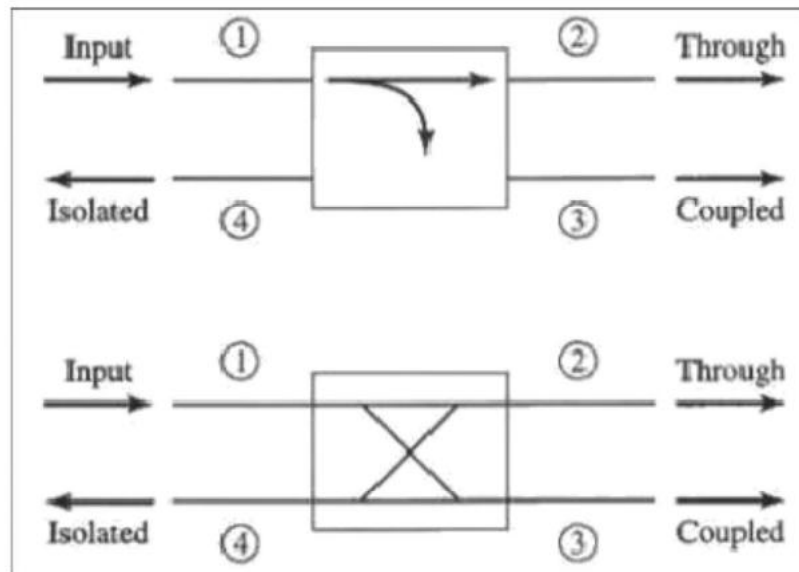
Diagram:



- It is a junction of 3 transmission lines
- In this case, if P_1 is the input port power then P_2 and P_3 are the power of output Port-2 and Port-3 respectively.
- To transfer maximum power from port-1 to port-2 and port-3 the impedance must match at the junction.

DIRECTIONAL COUPLER:

Diagram:



- It is a 4- port waveguide junction consisting of a primary waveguide 1-2 and a secondary waveguide 3-4.
- When all the ports are terminated in their characteristic impedance there is free transmission of power without reflection between port-1 and port-2 and no power transmission takes place between port-1 and port-3 or port-2 and port-4 a no coupling exists.
- The characteristic of a directional coupler is expressed in terms of its coupling factor and directivity.
- The coupling factor is the measure of ratio of power levels in primary and secondary lines.
- Directivity is the measure of how well the forward travelling wave in the primary waveguide couples only to a specific port of the secondary waveguide.
- In ideal case, directivity is infinite i.e. power at port-3 =0 because port-2 and port-4 are perfectly matched.
- Let wave propagates from port-1 to port-2 in primary line then:

$$\text{Coupling factor (dB)} = 10 \log_{10} (P_1/P_4)$$

$$\text{Directivity (dB)} = 10 \log_{10} (P_4/P_3)$$

Where P_1 = power input to port-1

P_3 = power output from port-3 and P_4 = power output from port-4

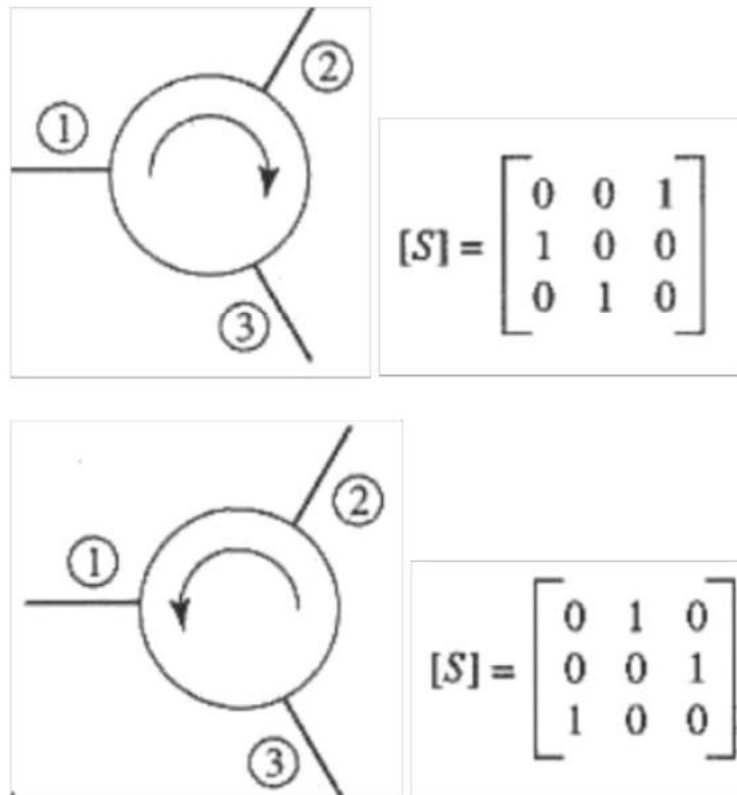
CIRCULATORS AND ISOLATORS:

Both microwave circulators and microwave isolators are non-reciprocal transmission devices that use Faraday rotation in the ferrite material.

CIRCULATOR:

- A microwave circulator is a multiport waveguide junction in which the wave can flow only in one direction i.e. from the n th port to the $(n+1)$ th port.
- It has no restriction on the number of ports
- 4-port microwave circulator is most common.
- One of its types is a combination of two 3-dB side hole directional couplers and a rectangular waveguide with two non reciprocal phase shifters.

Diagram:



- Each of the two 3db couplers introduce phase shift of 90 degrees
- Each of the two phase shifters produce a fixed phase change in a certain direction.
- Wave incident to port-1 splits into 2 components by coupler-1.
- The wave in primary guide arrives at port-2 with 180 degrees phase shift.
- The second wave propagates through two couplers and secondary guide and arrives at port-2 with a relative phase shift of 180 degrees.
- But at port-4 the wave travelling through primary guide phase shifter and coupler-2 arrives with 270 degrees phase change.
- Wave from coupler-1 and secondary guide arrives at port-4 with phase shift of 90 degrees.

Power transmission from port-1 to port-4 =0 as the two waves reaching at port-4 are out of phase by 180 degrees.

$$w_1 - w_3 = (2m+1) \pi \text{ rad/s}$$

$$w_2 - w_4 = 2n\pi \text{ rad/s}$$

Power flow sequence: 1-> 2 -> 3 -> 4-> 1

MICROWAVE ISOLATOR:

- A non reciprocal transmission device used to isolate one component from reflections of other components in the transmission line.
- Ideally complete absorption of power takes place in one direction and lossless transmission is provided in the opposite direction
- Also called UNILINE, it is used to improve the frequency stability of microwave generators like klystrons and magnetrons in which reflections from the load affects the generated frequency.
- It can be made by terminating ports 3 and 4 of a 4-port circulator with matched loads.
- Additionally it can be made by inserting a ferrite rod along the axis of a rectangular waveguide.

DIRECTIONAL COUPLER (DC):

Directional coupler is a 4 port wave guide junction. It consists of a primary wave guide and a secondary wave guide connects together through apertures. These are uni directional devices. Directional couplers are required to satisfy (1) reciprocity (2) conservation of energy (3) all ports matched terminated.

The characteristics of a **DC** can be expressed in terms of its:

1) Coupling factor: The ratio, in dB, of the power incident and the power coupled in auxiliary arm in forward direction.

Where P_i = Incident power; P_c = Coupled Power

2) Directivity: The ratio expressed in decibels, of the power coupled in the forward direction to the power coupled in the backward direction of the auxiliary arm with unused terminals matched terminated.

$$D = 10 \log_{10} (P_c / P_r) \text{ dB}$$

Where P_r = Reverse Power

P_c = Coupled Power

3) Insertion loss: The Ratio, expressed in decibels, of the power incident to the power transmitted in the main line of the coupler when auxiliary arms are matched terminated.

$$I = 10 \log_{10} (P_i / P_{i1})$$

Where P_{i1} = Received power at the transmitted port

4) Isolation: The ratio, expressed in decibels of the power incident in the main arm to the backward power coupled in the auxiliary arm, with other ports matched terminated.

$$L = 10 \log_{10} (P_i/P_r) \text{ Db}$$

For an ideal coupler D & I are infinite while C & L are Zero

Several types of directional couplers exist,
such as

1. Two hole directional coupler, Schwinger
2. directional coupler and Bethi - hole directional coupler.

Directional couplers are very good power samplers.

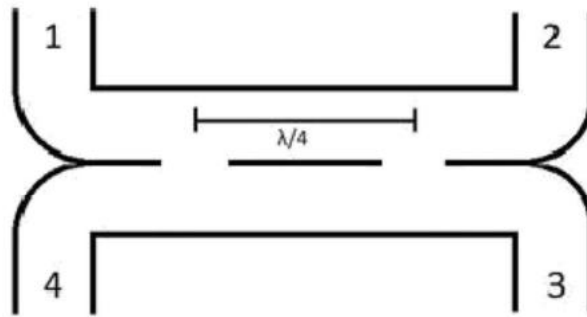
Bethe-hole coupler

Bethe-hole is a waveguide directional coupler, using a single hole, and it works over a narrow band. The Bethe-hole is a reverse coupler, as opposed to most waveguide couplers that use multi-hole and are forward couplers.

The origin of the name comes from a paper published by H A Bethe, titled "Theory of Diffraction by Small Holes", published in the Physical Review, back in 1942. If you google it you might find it, even though it is probably subject to copyright protection. This is a tough read, unless you like to ponder equations....

Multi-hole coupler

In waveguide, a two-hole coupler, two waveguides share a broad wall. The holes are $1/4$ wave apart. In the forward case the coupled signals add, in the reverse they subtract (180° apart) and disappear. Coupling factor is controlled by hole size. The "holes" are often x-shaped, or perhaps other proprietary shapes. It is possible to provide very flat coupling over an entire waveguide band if you know what you are doing (think "Chebychev"...)

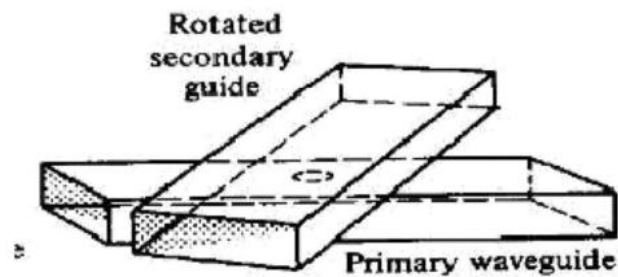


TWO HOLE DIRECTIONAL COUPLERS:

A two hole directional coupler with traveling wave propagating in it is illustrated. The spacing between the centers of two holes is

$$L = (2n + 1) \frac{\lambda_g}{4}$$

A fraction of the wave energy entered into port 1 passes through the holes and is radiated into the secondary guide as the holes act as slot antennas. The forward waves in the secondary guide are in same phase, regardless of the hole space and are added at port 4. the backward waves in the secondary guide are out of phase and are cancelled in port 3.



UNIT IV

MICROWAVE SOURCES :Klystrons, TWT's, Magnetrons

Limitations and losses of conventional tubes:

The Efficiency of Conventional Microwave Tube is largely independent of Frequency upto a certain limit, when frequency increases beyond a certain limit efficiency drastically decreases.

Conventional low frequency tubes like triodes fail to operate at microwave frequencies (MF) because the electron transit time from cathode to grid becomes do large that it cannot produce microwave oscillations. In order for an amplifier to work efficiently at the desired frequency the propagation times must be insignificant. And we see conventional tubes have a significant propagation times and hence cannot be used at microwave frequencies.

The device parameters for this tubes starts taking a dominating part in circuit and hence successful oscillations aren't met. There are also other limitations attached to them:

There are few important points that need to be noted when Microwave frequency is increased

- ☐ Interelectrode capacitance
- ☐ Dielectric losses
- ☐ Lead inductance effect
- ☐ Effects due to radiation losses and radio frequency(RF) losses
- ☐ Skin effect(*which is is the tendency of an alternating current (AC) to become distributed within a conductor such that the current density is largest near the surface of the conductor, and decreases with greater depths in the conductor*)
- ☐ Gain-bandwidth limitations

Interelectrode Capacitance

The Interelectrode capacitance in vacuum tubes at low or Medium frequency produce large Capacitive reactance with no serious effect.

The Capacitive reactance become so small when the frequency drastically increased

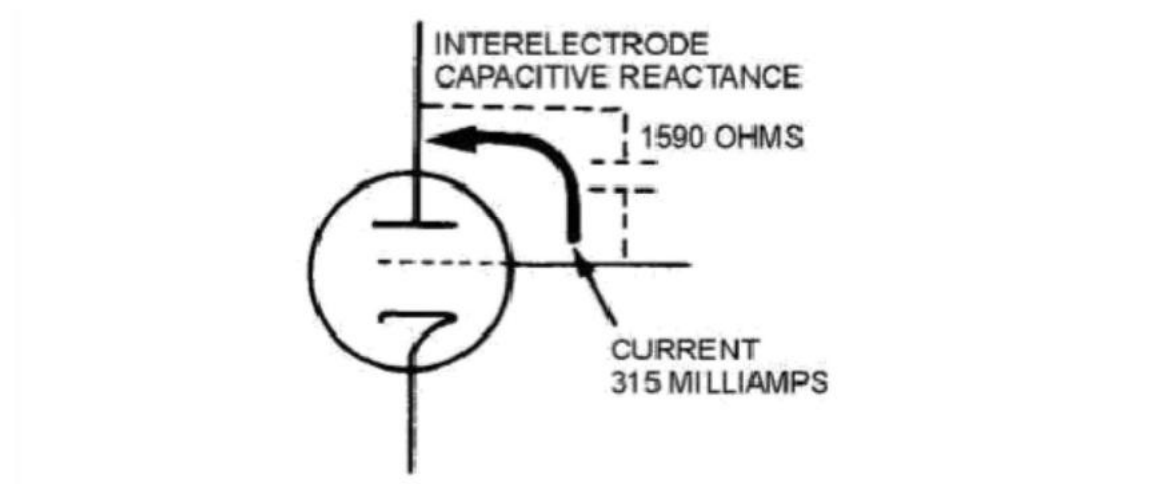


Figure —Interelectrode capacitance in a vacuum tube at 100 MHz

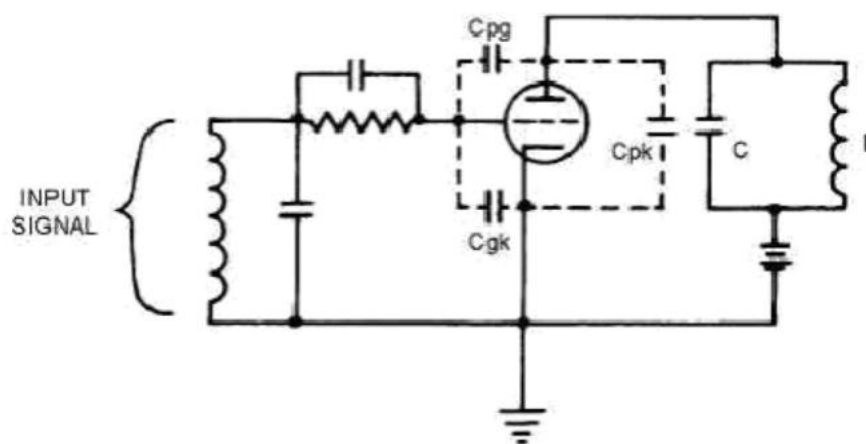


Figure —Interelectrode capacitance in a vacuum tube.

A good point to remember is that the higher the frequency, or the larger the interelectrode capacitance, the higher will be the current through this capacitance. The circuit in figure 2-1C, shows the interelectrode capacitance between the grid and the cathode (C_{gk}) in parallel with the signal source. As the frequency of the input signal increases, the effective grid-to-cathode impedance of the tube decreases because of a decrease in the reactance of the interelectrode capacitance. If the signal frequency is 100 megahertz or greater, the reactance of the grid-to-cathode capacitance is so small that much of the signal is short-circuited within the tube. Since the interelectrode capacitances are effectively in parallel with the tuned circuits, as shown in figures

Lead Inductance

Another frequency-limiting factor is the LEAD INDUCTANCE of the tube elements. Since the lead inductances within a tube are effectively in parallel with the interelectrode capacitance, the

net effect is to raise the frequency limit. However, the inductance of the cathode lead is common to both the grid and plate circuits. This provides a path for degenerative feedback which reduces overall circuit efficiency.

Lead Inductance within the tube are effectively in parallel with the interelectrode capacitance, the net effect is to raise the frequency limit, However the inductance of the positive cathode lead is common to both the Grid plate Circuit.

This provide the path for Regenerative Feed back which reduces overall Circuit efficiency.

Transient Time

Transient time is the time required for electrons to travel from the Cathode to the plate, this time time is insignificant at lower frequency .

However at the higher frequency the transient time become and appreciable portion of signal cycle and begins to hinder effeciency.

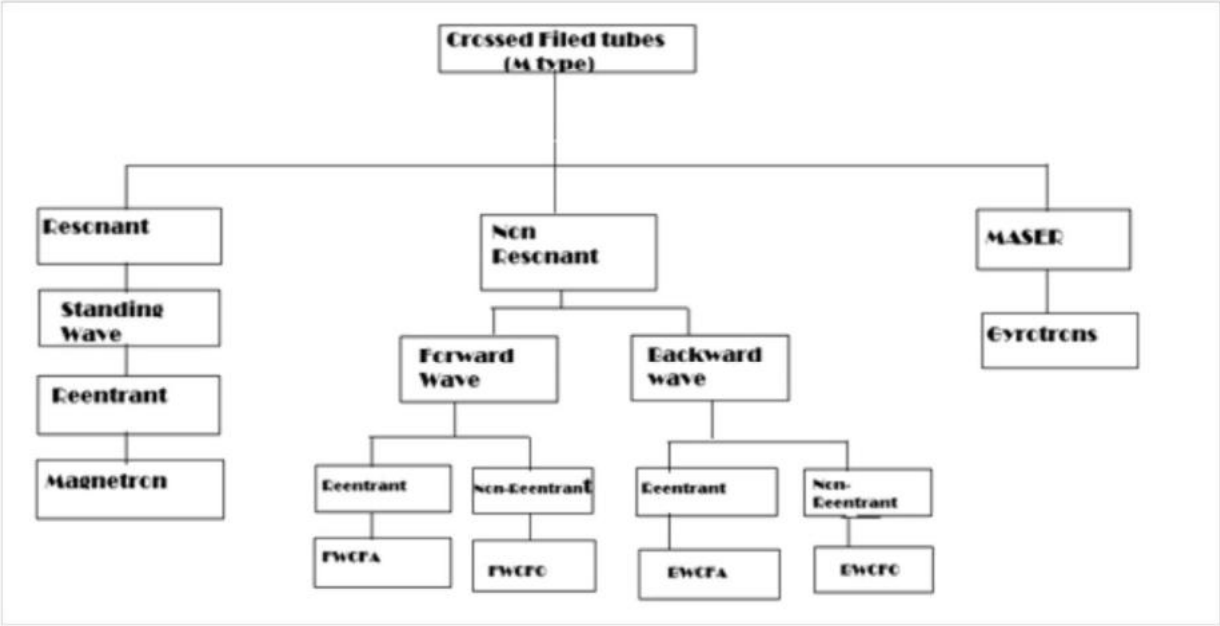
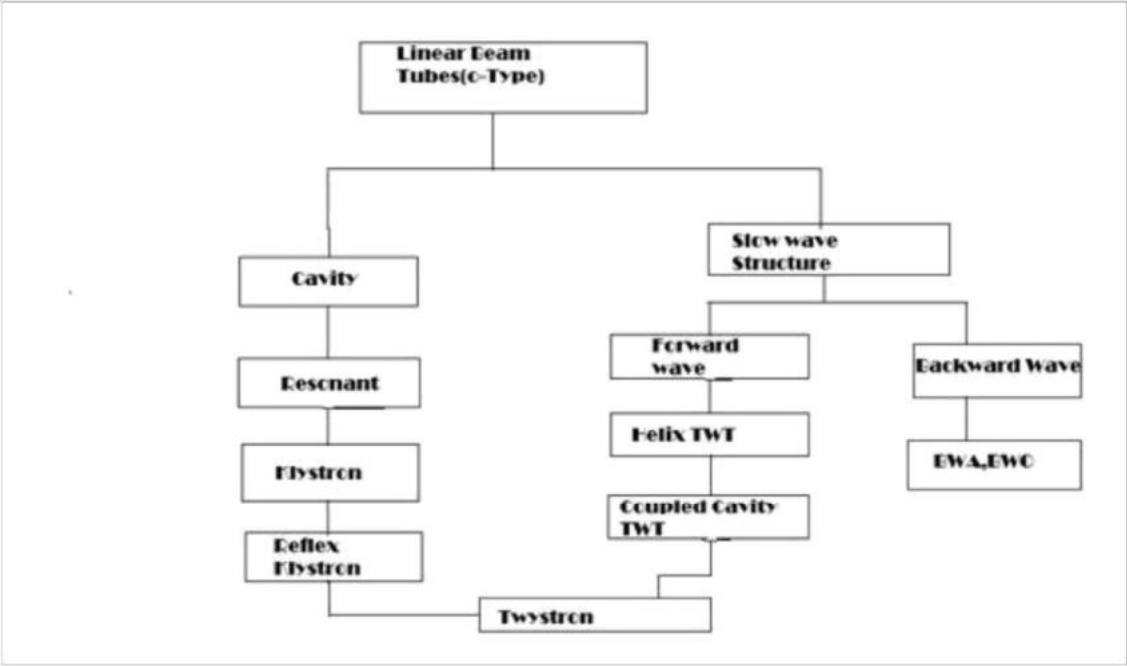
CLASSIFICATION OF MICROWAVE TUBES:

Microwave tubes can be broadly classifies into two categories

1.O-TYPE L inear Tubes (*Travelling tube amplifiers,Klystrons*)

In O-Type tube , a magnetic field whose axis coincides with the electron beam is used to hold the beam togetheras it travels the length of the tube

2.M-TYPE Tubes (*Magnetrons and cross field devices*)



Basically there are only main two types of microwave tubes

- Tubes with electromagnetic cavity(*klystrons and magnetrons*)
- Tubes with slow wave circuits(*traveling wave tubes*)

KLYSTRON

A **klystron** is a specialized linear-beam vacuum tube (evacuated electron tube). The pseudo-Greek word *klystron* comes from the stem form κλυσ- (*klys*) of a Greek verb referring to the action of waves breaking against a shore, and the end of the word *electron*.

The brothers Russell and Sigurd Varian of Stanford University are generally considered to be the inventors of the klystron. Their prototype was completed in August 1937. Upon publication in 1939,[1] news of the klystron immediately influenced the work of US and UK researchers working on radar equipment. The Varians went on to found Varian Associates to commercialize the technology (for example to make small linear accelerators to generate photons for external beam radiation therapy). In their 1939 paper, they acknowledged the contribution of A. Arsenjewa-Heil and O. Heil (wife and husband) for their velocity modulation theory in 1935.[2]

During the second World War, the Axis powers relied mostly on (then low-powered) klystron technology for their radar system microwave generation, while the Allies used the far more powerful but frequency-drifting technology of the cavity magnetron for microwave generation. Klystron tube technologies for very high-power applications, such as synchrotrons and radar systems, have since been developed.

The klystron is a linear-beam device that overcomes the transit-time limitations of a grid-controlled tube by accelerating an electron stream to a high velocity before it is modulated. Modulation is accomplished by varying the velocity of the beam, which causes the drifting of electrons into bunches to produce RF space current. One or more cavities reinforce this action at the operating frequency. The output cavity acts as a transformer to couple the high-impedance beam to a low-impedance transmission line. The frequency response of a klystron is limited by the impedance-bandwidth product of the cavities, but may be extended through stagger tuning or the use of multiple-resonance filter-type cavities.

The klystron is one of the primary means of generating high power at UHF and above. Output powers for multicavity devices range from a few thousand watts to 10MW or more. The klystron provides high gain and requires little external support circuitry. Mechanically, the klystron is relatively simple. It offers long life and requires minimal routine maintenance.

INTRODUCTION

Klystrons are used as an oscillator (such as the reflex klystron) or amplifier at microwave and radio frequencies to produce both low-power reference signals for superheterodyne radar receivers and to produce high-power carrier waves for communications and the driving force for linear accelerators.

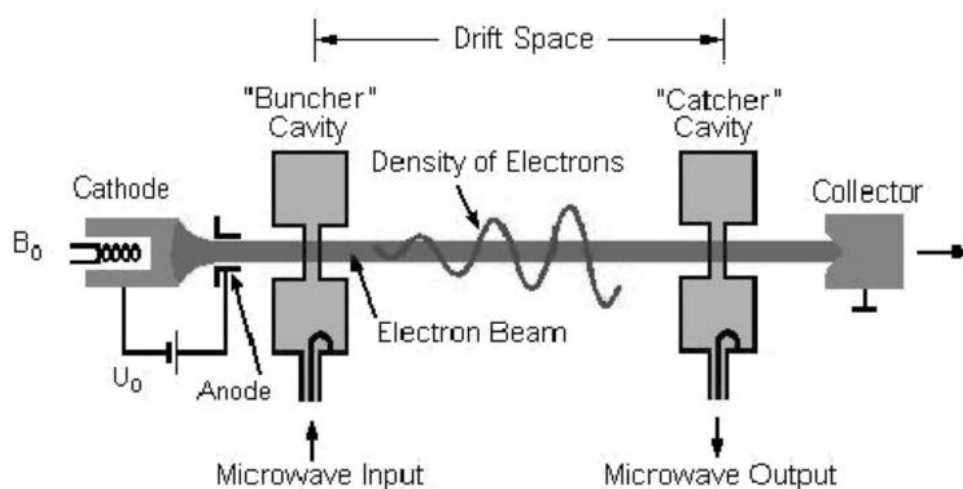
All modern klystrons are amplifiers, since reflex klystrons have been surpassed by alternative technologies. Klystron amplifiers have the advantage (over the magnetron) of coherently amplifying a reference signal so its output may be precisely controlled in amplitude, frequency and phase.

Many klystrons have a waveguide for coupling microwave energy into and out of the device, although it is also quite common for lower power and lower frequency klystrons to use coaxial couplings instead. In some cases a coupling probe is used to couple the microwave energy from a klystron into a separate external waveguide.

Explanation

Klystrons amplify RF signals by extracting energy from a DC electron beam. A beam of electrons is produced by a thermionic cathode (a heated pellet of low work function material), and accelerated to high voltage (typically in the tens of kilovolts). This beam is then passed through an input cavity. RF energy is fed into the input cavity at, or near, its natural frequency to produce a voltage which acts on the electron beam. The electric field causes the electrons to bunch: electrons that pass through during an opposing electric field are accelerated and later electrons are slowed, causing the previously continuous electron beam to form bunches at the input frequency. To reinforce the bunching, a klystron may contain additional "buncher" cavities. The electron bunches excite a voltage on the output cavity, and the RF energy developed flows out through a waveguide. The spent electron beam, which now contains less energy than it started with, is destroyed in a collector.

TWO-CAVITY KLYSTRON AMPLIFIER:



In the two-chamber klystron, the electron beam is injected into a resonant cavity. The electron beam, accelerated by a positive potential, is constrained to travel through a cylindrical *drift tube* in a straight path by an axial magnetic field. While passing through the first cavity, the electron beam is velocity modulated by the weak RF signal. In the moving frame of the electron beam, the velocity modulation is equivalent to a plasma oscillation, so in a quarter of one period of the plasma frequency, the velocity modulation is converted to density modulation, i.e. bunches of electrons. As the bunched electrons enter the second chamber they induce standing waves at the same frequency as the input signal. The signal induced in the second chamber is much stronger than that in the first.

TWO-CAVITY KLYSTRON OSCILLATOR:

The two-cavity amplifier klystron is readily turned into an oscillator klystron by providing a feedback loop between the input and output cavities. Two-cavity oscillator klystrons have the advantage of being among the lowest-noise microwave sources available, and for that reason have often been used in the illuminator systems of missile targeting radars.

The two-cavity oscillator klystron normally generates more power than the reflex klystron—typically watts of output rather than milliwatts. Since there is no reflector, only one high-voltage supply is to cause the tube to oscillate, the voltage must be adjusted to a particular value.

This is because the electron beam must produce the bunched electrons in the second cavity in order to generate output power. Voltage must be adjusted by varying the velocity of the electron beam to a suitable level due to the fixed physical separation between the two cavities. Often several "modes" of oscillation can be observed in a given klystron.

REFLEX KLYSTRON :

The reflex klystron is a single cavity variable frequency time-base generator of low power and load efficiency.

The reflex klystron uses a single-cavity resonator to modulate the RF beam and extract energy from it. The construction of a reflex klystron is shown in Figure In its basic form, the tube consists of the following elements:

- A cathode
- Focusing electrode at cathode potential
- Coaxial line or reentrant-type cavity resonator, which also serves as an anode
- Repeller or
- reflector electrode, which is operated at a moderately negative potential with respect to the cathode.

APPLICATION:

- ☐ It is widely used as in radar receiver
 - ☐ Local oscillators in microwave receiver
 - ☐ Portable microwave rings
 - ☐ Pump oscillator in parametric amplifier
-

CONSTRUCTION:

Reflex cavity klystron consists of an electron gun , filament surrounded by cathode and a floating electron at cathode potential

The cathode is so shaped that, in relation to the focusing electrode and anode, an electron beam is formed that passes through a gap in the resonator, as shown in the figure, and travels toward the repeller. Because the repeller has a negative potential with respect to the cathode, it turns the electrons back toward the anode, where they pass through the anode gap a second time.

Electron gun emits electron with constant velocity

OPERATION:

The electron that are emitted from cathode with constant velocity enter the cavity where the velocity of electrons is changed or modified depending upon the cavity voltage.

The oscillations is started by the device due to high quality factor and to make it sustained we have to apply the feedback.

Hence there are the electrons which will bunch together to deliver the energy act a time to the RF signal.

Inside the cavity velocity modulation takes place. Velocity modulation is the process in which the velocity of the emitted electrons are modified or change with respect to cavity voltage. The exit velocity or velocity of the electrons after the cavity is given as

$$V' = \sqrt{\frac{2q(Va + Vi \sin(\omega t))}{m}}$$

In the cavity gap the electrons beams get velocity modulated and get bunched to the drift space existing between cavity and repellar.

Bunching is a process by which the electrons take the energy from the cavity at a different time and deliver to the cavity at the same time .

Bunching continuously takes place for every negative going half cycle and the most appropriate time for the electrons to return back to the cavity ,when the cavity has positive peak .So that it can give maximum retardation force to electron.

It is found that when the electrons return to the cavity in the second positive peak that is 1 whole $\frac{3}{4}$ cycle.($n=1\pi$).It is obtained max power and hence it is called dominant mode.

The electrons are emitted from cathode with constant anode voltage Va , hence the initial entrance velocity of electrons is

$$V = \sqrt{\frac{2qVa}{m}} \text{ m/s}$$

$$V = \sqrt{\frac{2q(Va + Vi \sin(\omega t))}{m}}$$

$$V = \sqrt{\frac{2qVa}{m} + \frac{2qVaKiVi \sin \omega t}{Vam}}$$

$$V = \sqrt{Vo^2 + \frac{Vo^2 KiVi \sin \omega t}{Va}}$$

$$V = Vo \sqrt{1 + \frac{KiVi \sin \omega t}{Va}}$$

$$V = \sqrt{1 + \frac{Vo^2 KiVi \sin \omega t}{Va}}$$

$$\text{When, } \frac{KiVi}{Va} = \text{Depth of velocity}$$

$$V = V_0 \left(1 + \frac{V_0^2 i V_i \sin \omega t}{2 V_a} \right)$$

TRANSIENT TIME:

Transit time is defined as the time spent by the electrons in the cavity space or, time taken by the electrons to leave the cavity and again return to the cavity.

If 1 is the time at which electrons leave the cavity and 2 is the time at which electrons bunch in the cavity then, transit time

$$= t_1 - t_2$$

During this time the net displacement by electrons is zero. That the potential of two point A and B is V_A and V_B (plate) as known in figure, then,

$$\nabla + \sin +$$

$$= + + \sin$$

Neglecting the AC component,

$$= +$$

$$E = \frac{\partial V_{AB}}{\partial x} = \frac{-V_{AB}}{s} = \frac{-(V_a + V_R)}{s} \text{ -----(a)}$$

The force experienced on an electron

From equations a and b we get

$$F = qe = \frac{-q(V_a + V_R)}{s}$$

$$F = \frac{m \partial^2 x}{\partial t^2} \text{ -----(b)}$$

$$\frac{m \partial^2 x}{\partial t^2} = \frac{-q(V_a + V_R)}{s}$$

Integrating both sides we get,

$$\int \frac{m \partial^2 x}{\partial t^2} = \frac{-q(V_a + V_R)}{s}$$

$$\int m \partial x^2 = \frac{-q(V_a + V_R)}{s} \left\{ t_{t1} \partial t \right.$$

$$m \frac{\partial x}{\partial t} = \frac{-q(V_a + V_R)}{s} t_{t1} t$$

$$\frac{\partial x}{\partial t} = \frac{-q(V_a + V_R)}{mS} (t - t_1) + k_1$$

$\therefore k_1$ is called velocity constant and assumed to velocity at $(t - t_1)$

$$\frac{\partial x}{\partial t} = \frac{-q(V_a + V_R)}{mS} (t - t_1) + V(t_1)$$

$$\int \partial x = \frac{-q(V_a + V_R)}{mS} \int (t - t_1) \partial t + V(t_1)$$

$$\begin{aligned} X &= \frac{-q(V_a + V_R)}{mS} t_{t1} \left[\frac{t^2}{2} \right] - t_{t1} t_{t1} [t] + \int V(t_1) \partial t \\ &= \frac{-q(V_a + V_R)}{mS} \left[\frac{t^2 - t_1^2}{2} \right] - (t_2 - t_2)t_1 + \int V(t_1) \partial t \\ &= \frac{-q(V_a + V_R)}{mS} \left[\frac{t^2 - t_1^2 - 2t_2 t_1 + 2t_2^2}{2} \right] + \int V(t_1) \partial t \\ &= \frac{-q(V_a + V_R)}{mS} \left[\frac{t^2 + t_1^2 - 2t_2 t_1}{2} \right] + \int V(t_1) \partial t \\ &= \frac{-q(V_a + V_R)}{mS} \left[\frac{(t - t_1)^2}{2} \right] + \int V(t_1) \partial t \\ &= \frac{-q(V_a + V_R)}{mS} \left[\frac{(t - t_1)^2}{2} \right] + \int V(t_1) (t - t_1) + k_2 \end{aligned}$$

Where, k_2 is displacement constant at $t = t_2, x = 0$. In practice k_2 = the cavity width which is negligible with respect to cavity space s . Here we can neglect k_2 in the expression of x

At $t = t_2, x = 0$

$$\begin{aligned} 0 &= \frac{-q(V_a + V_R)}{mS} \left[\frac{(t - t_1)^2}{2} \right] + V(t_1) (t - t_1) \\ &\Rightarrow \frac{-q(V_a + V_R)}{mS} \left[\frac{(t_2 - t_1)^2}{2} \right] = V(t_1) (t_2 - t_1) \\ &\Rightarrow (t_2 - t_1) = \frac{2V(t_1)mS}{q(V_a + V_R)} \end{aligned}$$

TRANSIT ANGLE:

$$\omega t_r = \omega(t_2 - t_1) = \omega \frac{2V(t_1)ms}{q(V_a + V_R)}$$

$$\omega t_r = \omega \frac{2V(t_1)ms}{q(V_a + V_R)}$$

We know,

$$n = +3/4, \text{ for } n=0,1,2,3,4,\dots$$

$$n = 1/4, \text{ for } n=0,1,2,3,4,\dots$$

$$2 \times \pi(t_2 - t_1) = \left(n - \frac{1}{4}\right) 2\pi \times T$$

$$2 \times \pi \times f(t_2 - t_1) = \left(2n\pi - \frac{\pi}{2}\right)$$

$$\omega t_r = \left(2n\pi - \frac{\pi}{2}\right) = \omega \frac{2V(t_1)ms}{q(V_a + V_R)}$$

OUTPUT POWER:

The beam current of Reflex klystron is given as

$$I_b = I_0 + \sum_{n=1}^{\infty} (2I_0 I_n(x')) \cos(n\omega t - \phi)$$

I_0 is dc current due to cavity voltage is given by

$$P_{dc} = V_a I_0 \text{ -----(1)}$$

The ac component of the current is given by

$$I_{ac} = \sum_{n=1}^{\infty} (2I_0 I_n(x')) \cos(n\omega t - \phi)$$

For $(n=1)$, we have fundamental current component ie,

$$\text{If } (2I_0 I_1(x')) \cos(n\omega t - \phi)$$

$$\text{For } n=2; \cos(n\omega t - \phi) = 1$$

$$I_2 = 2I_0 K_1 J_1(X_1)$$

$$\text{Power} = \frac{v_i I_2}{2} = \frac{v_i}{2} (2I_0 k_i J_i(X_i))$$

$$\omega(t_2 - t_1) = \omega \frac{2V(t_1)ms}{q(V_a + V_R)}$$

$$\omega t_2 = \omega t_1 + \frac{2V_0 ms \omega}{q(V_a + V_R)} \left[1 + \frac{k_i v_i}{2v_a} \sin \omega t_1 \right]$$

$$\therefore \left[\frac{2V_0 ms \omega}{q(V_a + V_R)} \right] = \alpha$$

$$\rightarrow \omega t_2 = \omega t_1 + \alpha \left[1 + \frac{k_i v_i}{2v_a} \sin \omega t_1 \right]$$

$$\rightarrow \omega t_2 = \omega t_1 + \alpha + \frac{k_i v_i \alpha}{2v_a} \sin \omega t_1$$

Where, $\frac{k_i v_i \alpha}{2v_a} = x$, bunching parameter

$$V_i = \frac{2v_a x}{\alpha k_i} = \frac{2v_a x}{(2n\pi - \frac{\pi}{2}) k_i}$$

$$P_{o/p} = \frac{2v_a x I_0 \times J_i(x)}{(2n\pi - \frac{\pi}{2})}$$

$$\text{Efficiency:} \quad \eta\% = \frac{P_{\text{output}}}{P_{\text{input}}} \times 100\%$$

$$= \frac{2v_a x I_0 \times J_i(x)}{(2n\pi - \frac{\pi}{2}) v_a x I_0}$$

$$= \frac{2 \times J_i(x)}{(2n\pi - \frac{\pi}{2})} \times 100\%$$

Electronics admittance of reflex klystron

It is defined as the ratio of current induced in the cavity by the modulation of electron beam to the voltage across the cavity gap.

TRAVELING-WAVE TUBE :

A **traveling-wave tube (TWT)** is an electronic device used to amplify radio frequency signals to high power, usually in an electronic assembly known as a **traveling-wave tube amplifier (TWTA)**.

The TWT was invented by Rudolf Kompfner in a British radar lab during World War II, and refined by Kompfner and John Pierce at Bell Labs. Both of them have written books on the device.[1][2] In 1994, A.S. Gilmour wrote a modern TWT book[3] which is widely used by U.S. TWT engineers today, and research publications about TWTs are frequently published by the IEEE.

Cutaway view of a TWT. (1) Electron gun; (2) RF input; (3) Magnets; (4) Attenuator; (5) Helix coil; (6) RF output; (7) Vacuum tube; (8) Collector.

The device is an elongated vacuum tube with an electron gun (a heated cathode that emits electrons) at one end. A magnetic containment field around the tube focuses the electrons into a beam, which then passes down the middle of a wire helix that stretches from the RF input to the RF output, the electron beam finally striking a collector at the other end.

A directional coupler, which can be either a waveguide or an electromagnetic coil, fed with the low-powered radio signal that is to be amplified, is positioned near the emitter, and induces a current into the helix.

The helix acts as a delay line, in which the RF signal travels at near the same speed along the tube as the electron beam. The electromagnetic field due to the current in the helix interacts with the electron beam, causing bunching of the electrons (an effect called *velocity modulation*), and the electromagnetic field due to the beam current then induces more current back into the helix (i.e. the current builds up and thus is amplified as it passes down). A second directional coupler, positioned near the collector, receives an amplified version of the input signal from the far end of the helix. An attenuator placed on the helix, usually between the input and output helices, prevents reflected wave from travelling back to the cathode.

The bandwidth of a broadband TWT can be as high as three octaves, although tuned (narrowband) versions exist, and operating frequencies range from 300 MHz to 50 GHz. The voltage gain of the tube can be of the order of 70 decibels.

A TWT has sometimes been referred to as a traveling-wave amplifier tube (TWAT),[4][5] although this term was never really adopted. "TWT" is sometimes pronounced by engineers as "TWIT".[6]

MULTICAVIDITY KLYSTRON

In all modern klystrons, the number of cavities exceeds two. A larger number of cavities may be used to increase the gain of the klystron, or to increase the bandwidth.

MAGNETRON :

The magnetron encompasses a class of devices finding a wide variety of applications. Pulsed magnetrons have been developed that cover frequency ranges from the low UHF band to 100 GHz. Peak power from a few kilowatts to several megawatts has been obtained. Typical overall efficiencies of 30 to 40 percent may be realized, depending on the power level and operating frequency. CW magnetrons also have been developed, with power levels of a few hundred watts in a tunable tube, and up to 25kW or more in a fixed-frequency device. Efficiencies range from 30 percent to as much as 70 percent. The magnetron operates electrically as a simple diode. Pulsed modulation is obtained by applying a negative rectangular voltage waveform to the cathode with the anode at ground potential. Operating voltages are less critical than for beam tubes line-type modulators often are used to supply pulsed electric power. The physical structure of a conventional magnetron is shown in Figure.

High-power pulsed magnetrons are used primarily in radar systems. Low-power pulsed devices find applications as beacons. Tunable CW magnetrons are used in ECM (electronic countermeasures) applications. Fixed-frequency devices are used as microwave heating sources

Figure :Reentrant emitting-sole crossed-field amplifier tube.

Tuning of conventional magnetrons is accomplished by moving capacitive tuners or by inserting symmetrical arrays of plungers into the inductive portions of the device. Tuner motion is produced by a mechanical connection through flexible bellows in the vacuum wall. Tuning ranges of 10 to 12 percent of bandwidth are possible for pulsed tubes, and as much as 20 percent for CW tubes.

Operating Principles

Most magnetrons are built around a cavity structure of the type shown in Figure. The device consists of a cylindrical cathode and anode, with cavities in the anode that open into the cathode-anode space—the interaction space as shown. Power can be coupled out of the cavities by means of a loop or a tapered waveguide. Cavities, together with the spaces at the ends of the anode block, form the resonant system that determines the frequency of the generated oscillations. The actual shape of the cavity is not particularly important, and various types are used, as illustrated in Figure. The oscillations associated with the cavities are of such a nature that alternating magnetic flux lines pass through the cavities parallel to the cathode axis, while the alternating electric fields are confined largely to the region where the cavities open into the interaction space. The most important factors determining the resonant frequency of the system are the dimensions and shape of the cavities in a plane perpendicular to the axis of the cathode. Frequency also is affected by other factors such as the end space and the axial length of the anode block, but to a lesser degree.

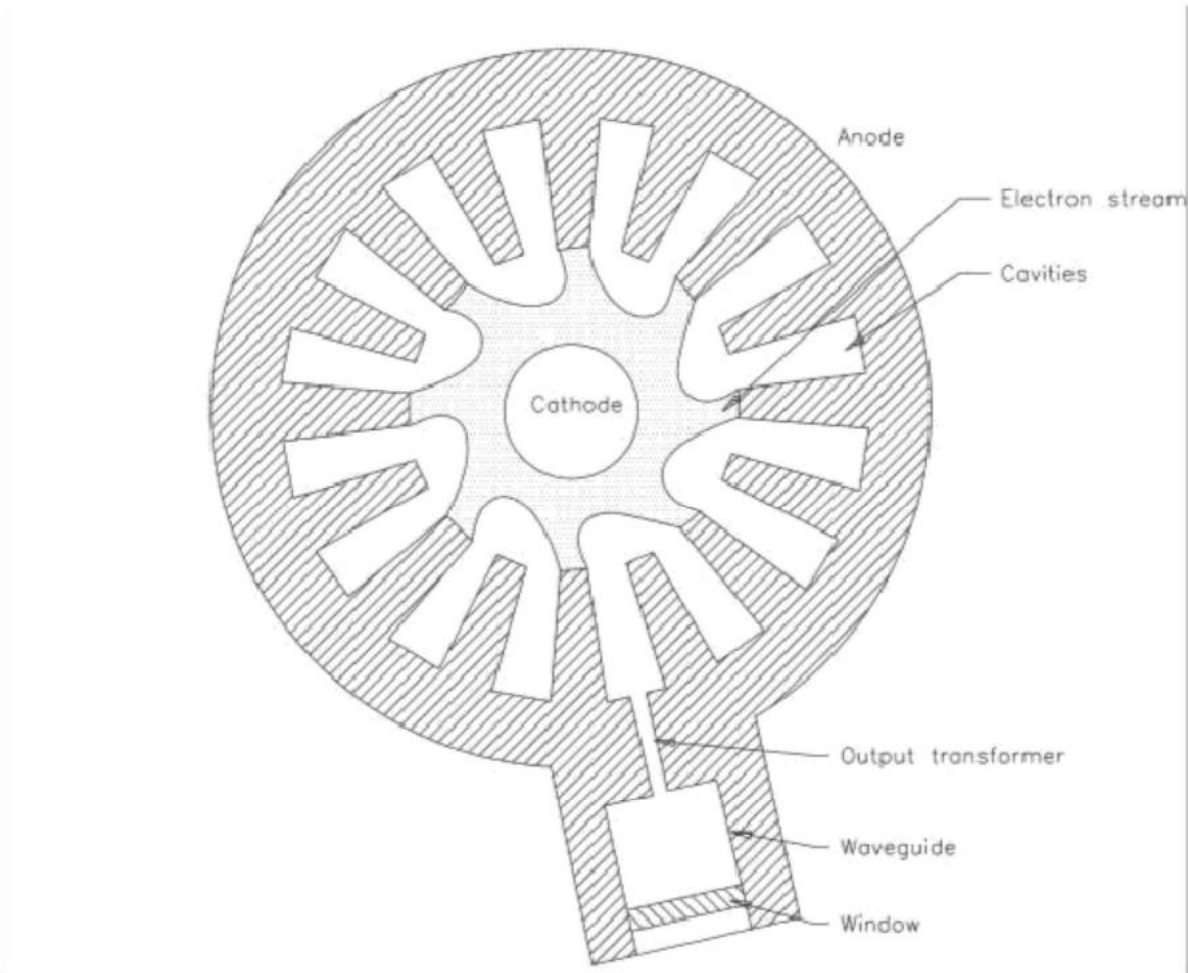


Figure : Conventional magnetron structure.

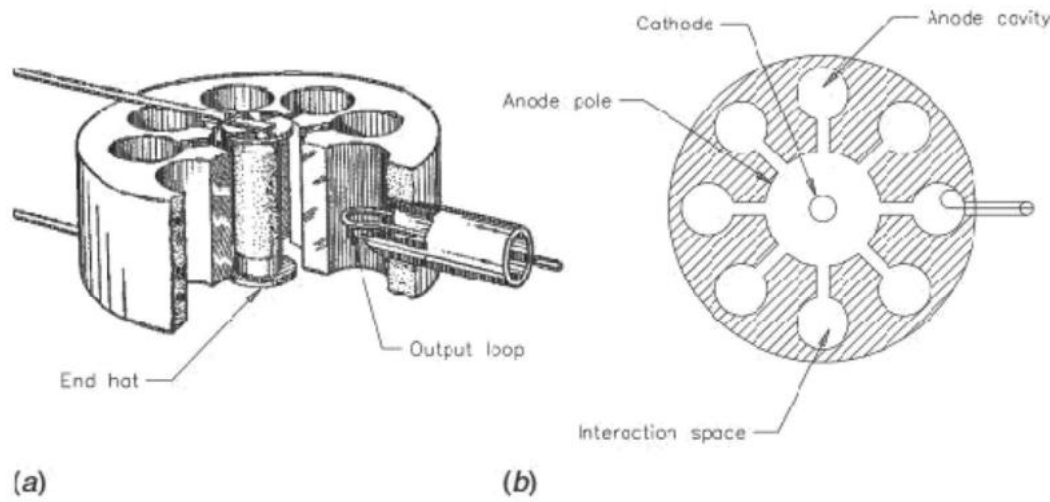


Figure: Cavity magnetron oscillator: (a) cutaway view, (b) cross section view perpendicular to the axis of the cathode.

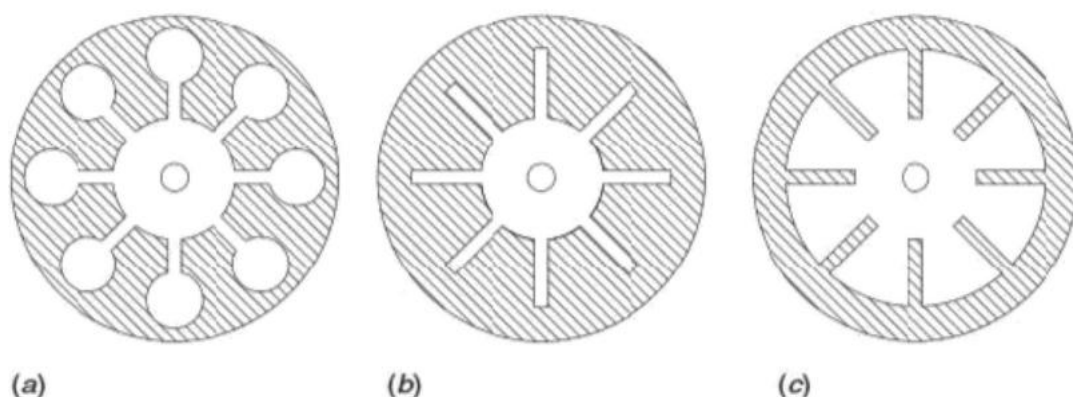


Figure : Cavity magnetron oscillator anode: (a) hole-and-slot type, (b) slot type, (c)vane type.

The magnetron requires an external magnetic field with flux lines parallel to the axis of the cathode. This field usually is provided by a permanent-magnet or electromagnet. The cathode is commonly constructed as a cylindrical disk

DIFFERENCE BETWEEN REFLEX KLYSTRON AND MAGNETRON:

REFLEX KLYSTRON	MAGNETRON
It is a linear tube in which the magnetic field is applied to focus the electron and electric field is applied to drift the electron.	In magnetron the magnetic field and electric field are perpendicular to each other hence it is called as cross field device.
In klystron the bunching takes places only inside the cavity which is very small ,hence generate low power and low frequency.	In magnetron the interacting or bunching space is extended so the efficiency can be increase.

APPLICATION:

- Used as oscillator.
- Used in radar communication.
- Used in missiles.
- Used in microwave oven (in the range of frequency of 2.5Ghz).

TYPES OF MAGNETRON:

Magnetron is of 3 types:

Negative resistance type.

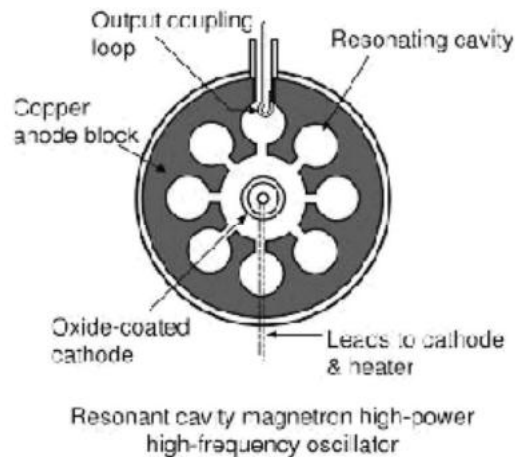
Cyclotron frequency type.

Cavity type.

Construction of cavity magnetron:

Cavity type magnetron depends upon the interaction of electron with a rotating magnetic field with constant angular velocity.

FIGURE:

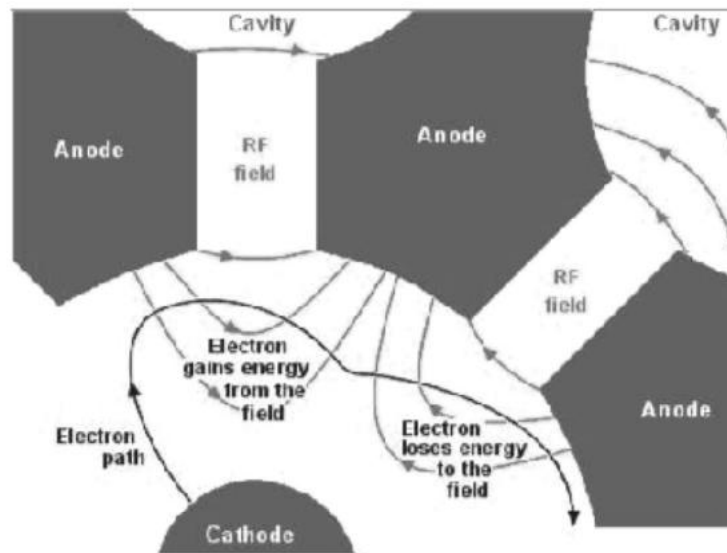


magnetron consists of a cathode which is used to emit electrons and a number of anode cavities a permanent magnet is placed on the backside of cathode. The space between anode cavity and cathode is called interacting space. The electron which are emitted from cathode moves in different path in the interacting space depending upon the strength of electron and magnetic field applied to the magnetron.

OPERATION:

EFFECT OF ELECTRIC FIELD ONLY:

- In the absence of magnetic field($B=0$) the electron travel straight from the cathode to the anode due to the radial electric field force acting on it(indicated by path A).
- If the magnetic field strength increases slightly it will exert a lateral force which bends the path of the as indicated in path B.



The radius of the path is given as

$$r = \frac{mv}{eB}$$

where v = velocity of electron

B = magnetic field strength

- If from reaching the anode current become zero (indicated by path D). the strength of magnetic field is made sufficiently high enough, so to prevent the electron
- The magnetic field required to return the electron back to the cathode just touching the surface of anode is called critical magnetic field or cut off magnetic field (B_c).
- If $B > B_c$ the electron experiences a greater rotational force and may return back to the cathode quite faster this results in heating of cathode.

Effect of magnetic field:

The force experienced by the electron because of magnetic field only.

$$F = q(v \times B)$$

$$= qvB \sin \theta$$

For maximum force $\theta = 90^\circ$

$$F_{\max} = qvB$$

And hence the electron which are emitted, moved in a right angle with respect to force.

If the magnetic field strength is sufficiently large enough, then the electrons emitted will return back to the cathode with high velocity which may destroy the cathode. This effect is called Back heating of cathode.

MODE OF OSCILLATION:

- The shape consisting of oscillation can maintain if the phase difference between anode cavity is $\pi/4$ where n is the mode of operation and the best result can be obtained for $n=4$ $\Rightarrow \pi/4$ (for $n=4$ hence it is called mode operation).
- It is assumed that each anode cavity is of $\pi/4$ length, hence a voltage antinode will exist at the opening of anode cavity and the lines of forces present due to the oscillation started by high quality factor device.
- In the above figure the electron followed by path 'b' is so emitted that it is not influenced by the electric lines of forces hence it will spend very less time inside the cavity and doesn't contribute to the oscillation so it is called unfavourable electron.
- The electron followed by path a is so emitted that it is influenced by the electric lines of forces at positions 1, 2 & 3 respectively where the velocity increases or decreases, hence more time is spent inside the cavity therefore it is called favourable electron.
- Any favourable electron which is emitted earlier or later with respect to reference electron (let a) may be bunched together due to change in velocity by the effect of electric lines of forces. This type of bunching is called phase focusing effect.
- Electrons emitted from the cathode may rotate around itself in a confined area in a shape of spoke (spiral) at an angular velocity and before delivering the energy to the anode cavity. They will rotate until they reach the anode and are completely absorbed by them. Hence the magnetron is also called travelling wave magnetron.

CUT OFF MAGNETIC FIELD (BC):

- Assume a cylindrical magnetron whose inner radius is 'a' and outer radius is 'b' and the magnetic field is as shown in the figure. Under the effect of magnetic field the electrons will rotate in a circular path at any point the force on the electron will be balanced by the centrifugal force.

$$\begin{aligned}\frac{mv^2}{r} &= qvB_z \\ \Rightarrow v &= \frac{qB_z}{m} r \\ \rightarrow \frac{v}{r} &= \frac{qB_z}{m} \\ \Rightarrow \omega &= \frac{qB_z}{m} \\ \Rightarrow f &= \frac{1}{2\pi} \left(\frac{qB_z}{m} \right)\end{aligned}$$

In cylindrical coordinate system the equation of motion is given as

$$\Rightarrow \frac{1}{\partial} \frac{d}{dt} (v^2 \frac{d\theta}{dt}) = \omega \frac{dr}{dt}$$

$$\Rightarrow \int \frac{d}{dt} (r^2 \frac{d\theta}{dt}) = \omega \int \frac{dr}{dt} \cdot r$$

$$\Rightarrow r^2 \left(\frac{d\theta}{dt} \right) = \frac{\omega r^2}{2} + k$$

Due to electric field only, the electrons move radially from cathode to anode.

$$\Rightarrow \frac{1}{2} m v^2 = q V_{dc}$$

$$\Rightarrow v = \sqrt{\frac{2qV_{dc}}{m}}$$

$$B_c = \frac{\frac{8mV_{dc}}{b(1-a^2/b^2)}}{b(1-a^2/b^2)}$$

$$V_{dc} = \frac{q}{8m} b^2 (1 - a^2/b^2) B_z^2$$

UNIT-V

Microwave Solid State Devices & Measurements

GUNN DIODE BASICS:

It has negative resistance property by which gunn diode act as oscillator. To achieve this capacitance and shunt load resistance need to be tuned but not greater than negative resistance. The figure describes **GUNN diode** equivalent circuit. Here active region is about 6-18 μm long. It has negative resistance of about 100 Ohm with parallel capacitance of about 0.6 PF. Gunn diode will have efficiency of only few percentage.

Commercial GUNN diode need supply of about 9V with operating current of 950mA and available from 4GHz to 100GHz frequency band. It is preferably placed in a resonant cavity. The GUNN diode is basically a TED i.e. Transferred Electron Device capable of oscillating based on different modes. In a unresonant transit time mode, radio frequencies of upto 1-18 GHz with power of upto 2 watt can be achieved. In a resonant limited space charge mode, radio frequencies of upto 100 Ghz with about 100watts of pulsed power can be achieved.

Gunn Effect:

Gun effect was first observed by GUNN in n_type GaAs bulk diode. According to GUNN, abovesome critical voltage corresponding to an electric field of 2000-4000v/cm, the current in

every specimen became a fluctuating function of time. The frequency of oscillation was determined mainly by the specimen and not by the external circuit.

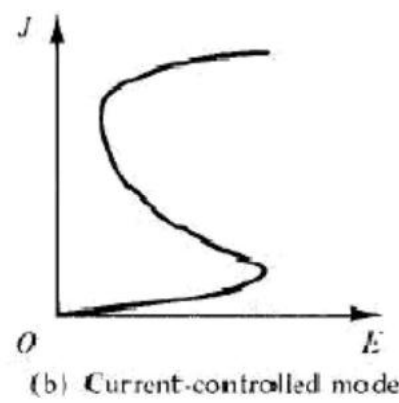
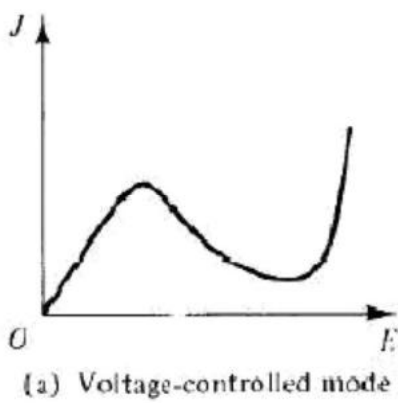
**RIDLEY-WATKINS-HILSUM (RWH)
THEORY Differential Negative Resistance**

The fundamental concept of the Ridley-Watkins-Hilsum (RWH) theory is the differential negative resistance developed in a bulk solid-state III-V compound when either a voltage (or electric field) or a current is applied to the terminals of the sample.

There are two modes of negative-resistance devices:

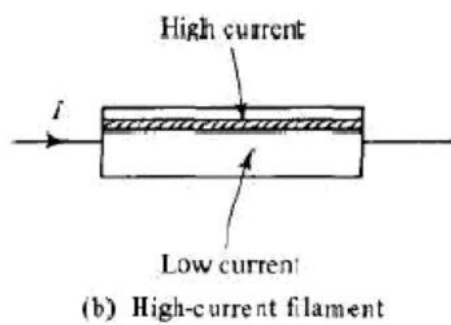
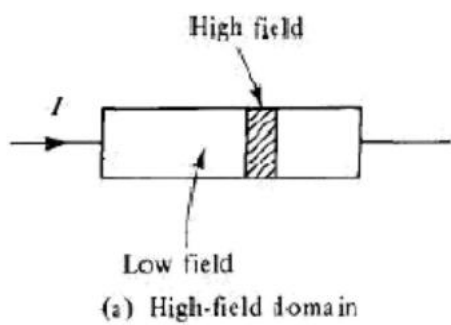
i) Voltage-controlled and

ii) current controlled modes as shown in Fig.



ii) current controlled modes as shown in Fig.

In the voltage-controlled mode the current density can be multivalued, whereas in the current controlled mode the voltage can be multivalued.



The major effect of the appearance of a differential negative-resistance region in the current density- field curve is to render the sample electrically unstable. As a result, the initially homogeneous sample becomes electrically heterogeneous in an attempt to reach stability.

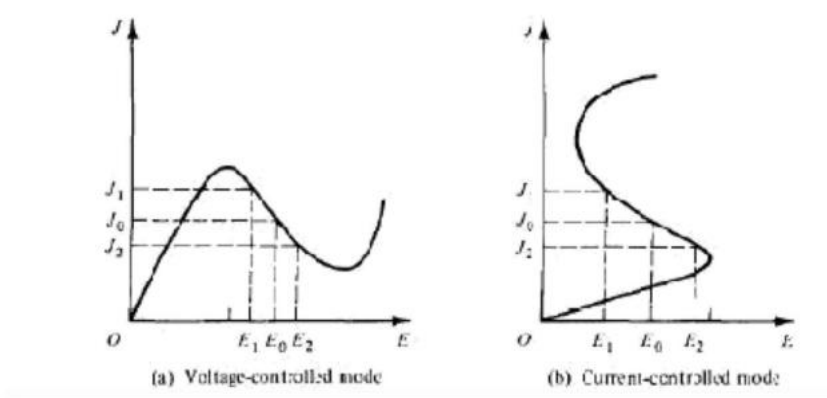
In the voltage-controlled negative-resistance mode high-field domains are formed, separating two low- field regions. The interfaces separating low and high-field domains lie along equipotentials;

thus they are in planes perpendicular to the current direction as shown in Fig. (a). In the current- controlled negative-resistance mode splitting the sample results in high-current filaments running along the field direction as shown in Fig. (b). Expressed mathematically,

$$\frac{dI}{dV} = \frac{dJ}{dE} = \text{negative resistance}$$

If an electric field E_0 (or voltage V_0) is applied to the sample, for example, the current density is generated. As the applied field (or voltage) is increased to E_2 (or V_2), the current density is decreased to J_2 .

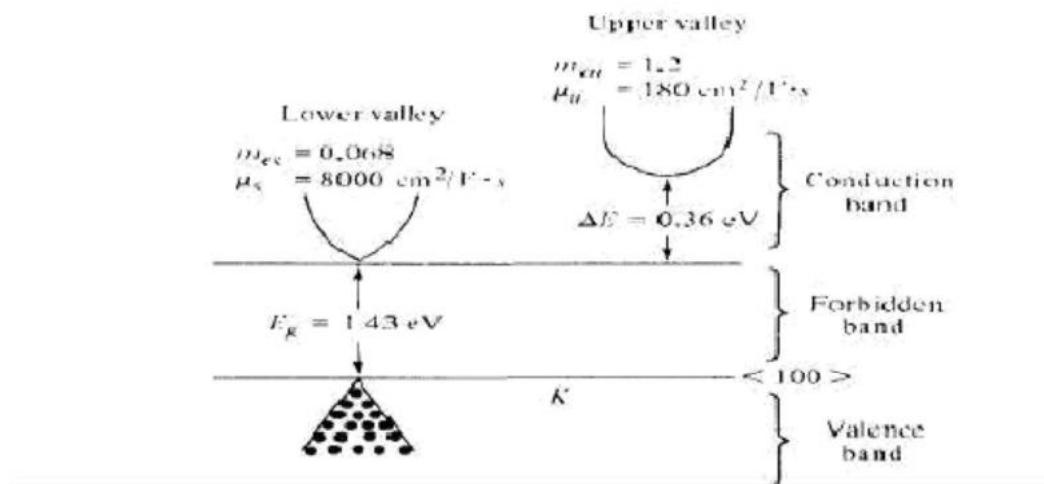
When the field (or voltage) is decreased to E_1 (or V_1), the current density is increased to J_1 . These phenomena of the voltage controlled negative resistance are shown in Fig. (a).Similarly, for the current controlled mode, the negative-resistance profile is as shown in Fig. (b).



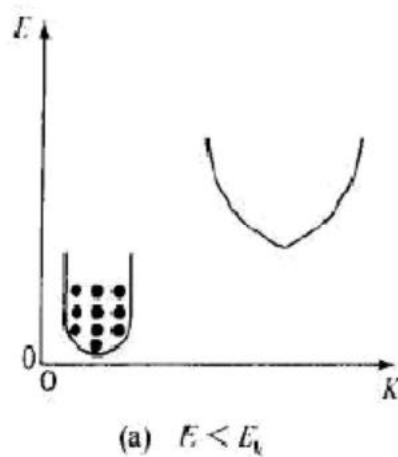
Two-Valley Model Theory

According to the energy band theory of then-type GaAs, a high-mobility lower valley is separated by an energy of 0.36 eV from a low-mobility upper valley

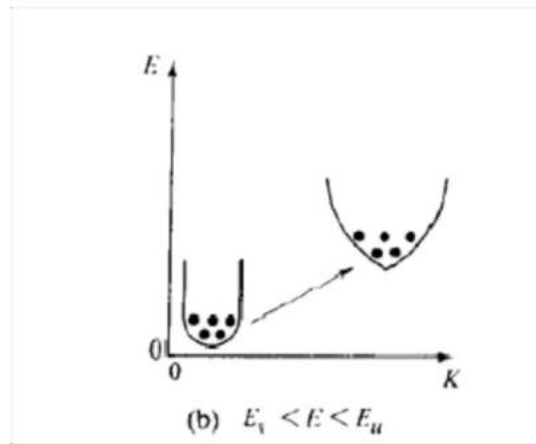
Valley	Effective Mass M_e	Mobility μ	Separation ΔE
Lower	$M_{el} = 0.068$	$\mu_l = 8000 \text{ cm}^2/\text{V-sec}$	$\Delta E = 0.36 \text{ eV}$
Upper	$M_{eu} = 1.2$	$\mu_u = 180 \text{ cm}^2/\text{V-sec}$	$\Delta E = 0.36 \text{ eV}$



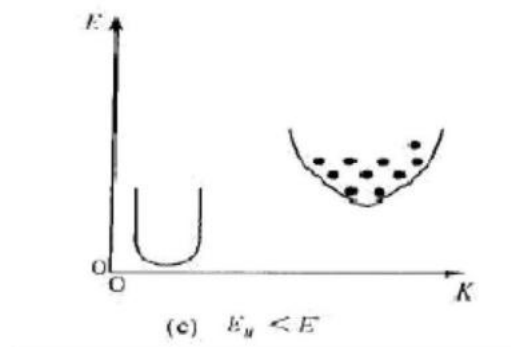
When the applied electric field is lower than the electric field of the lower valley ($E < E_c$), no electrons will transfer to the upper valley as show in Fig. (a).



When the applied electric field is higher than that of the lower valley and lower than that of the upper valley ($E_c < E < E_u$), electrons will begin to transfer to the upper valley as shown in Fig.(b).



And when the applied electric field is higher than that of the upper valley ($E_u < E$), all electrons will transfer to the upper valley as shown in Fig. (c).



If electron densities in the lower and upper valleys are n_l and n_u , the conductivity of the n -type GaAs is

$$\sigma = e(\mu_l n_l + \mu_u n_u)$$

where e = the electron charge
 μ = the electron mobility
 $n = n_l + n_u$ is the electron density

When a sufficiently high field E is applied to the specimen, electrons are accelerated and their effective temperature rises above the lattice temperature. Furthermore, the lattice temperature also increases. Thus electron density n and mobility μ are both functions of electric field E .

Gunn diode advantages

Following are major advantages of the Gunn diode.

∴ High frequency stability

-
- Higher bandwidth and reliability
 - Smaller size
 - Ruggedness in operation
 - low supply voltage
 - noise performance similar to klystron
 - low cost of manufacturing

Gunn diode disadvantages

- High turn on voltage
- low efficiency below 10GHz
- Poor bias and temperature stability
- Small tuning range
- Higher spurious FM noise
- higher device operating current and hence more power dissipation
- Lower efficiency and power at millimeter band

GUNN Diode Applications

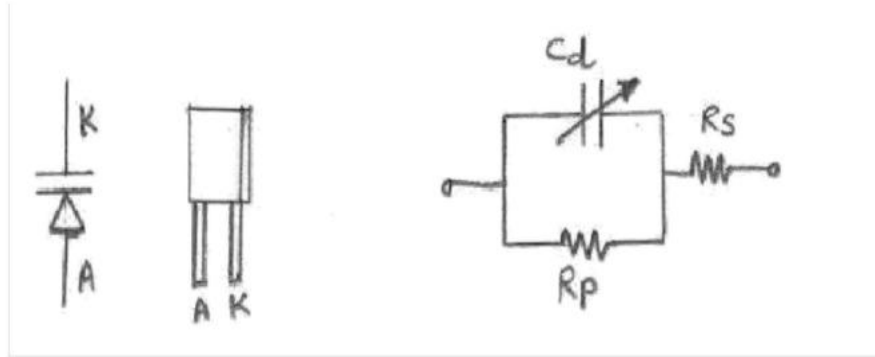
- Gunn diode is used as low and medium power oscillators in microwave instruments and receiver circuits
- As pump sources in parametric amplifiers
- Used in police radars and also in CW doppler radars
- Gunn diode oscillators are used to generate power at microwave frequencies for various applications such as automatic door openers, traffic gates, traffic signal controllers etc.

VARACTOR DIODE BASICS:

This page describes **varactor diode** basics and varactor diode applications. The link to varactor diode calculator, PIN diode, Tunnel diode and GUNN diode basics and applications are also mentioned.

Varactor diode is one of the many [microwave semiconductor devices](#) in use today.

They are manufactured with gallium arsenide. Figure depicts symbol of varactor diode and also typical manufacturing package. This diode is used as variable capacitor and as variable reactor in microwave circuits.



Varactor diode is special type of PN junction diode, in which PN junction capacitance is controlled using reverse bias voltage. When the diode is forward biased, current will flow through the diode. When the diode is reverse biased, charges in the P and N semiconductors are drawn away from the PN junction interface and hence forms the high resistance depletion zone. The equation of the varactor capacitance proportional to the reverse bias voltage is outlined below.

$$C_j = CK / (V_b - V)^m$$

Where,

C_j is the diode capacitance

C is the diode capacitance when the device is unbiased i.e. when V is zero.

V is the applied reverse voltage

V_b is barrier voltage at junction

m is the constant and depends of the material

K is also constant often taken as value of 1

The equivalent circuit of the varactor diode is mentioned in the figure along with the symbol. From the circuit maximum operating frequency of the varactor diode depends on the series resistance and diode capacitance and it is mentioned in the equation below.

$$F = 1 / 2\pi R_s C_j$$

Quality factor of the varactor diode is mentioned in the equation below.

$Q = F/f$, where F is the cutoff frequency and f is the operating frequency.

Varactor Diode Applications

Following are the varactor diode applications:

- ☐ It is used in variable resonant tank LC circuit. Here C part is varied using varactor diode.
- ☐ AFC(Automatic Frequency Control) where in varactor diode is used to set LO signal.
- ☐ Varactor is used as frequency modulator.
- ☐ It is used as frequency multiplier in microwave receiver LO.
- ☐ It is used as RF phase shifter.

TUNNEL DIODE BASICS:

Tunnel diode is one of the many microwave semiconductor devices in use today. This page covers Tunnel diode basics and its applications. The tunnel diode operates on tunneling principle which is a majority carrier event. There are few conditions which ensures high tunneling probability.

□ thin depletion zone

□ heavy doping of semiconductor material

There should be filled and empty energy state for each tunneling carrier.

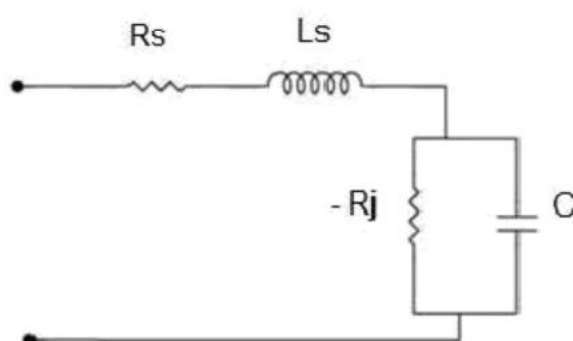


Fig:1 Tunnel Diode Equivalent Circuit

From the tunnel diode equivalent circuit it is imperative that tunnel diode oscillates when total reactance is zero and negative resistance balance out the series resistance. There are two type of frequencies equations for which are mentioned below.

$$\text{Resistive Frequency, } F_r (\text{Hz}) = (1/2 * \pi * R * C_j) * ((R/R_s) - 1)^{0.5}$$

$$2^{0.5} \text{ Self resonant Frequency } F_s (\text{Hz}) = (1/2 * \pi) ((1/L_s C_j) - (1/(R \emptyset_j)))$$

Where in,

R is the absolute value of negative resistance (Ohms)

R_s is the series resistance (Ohms)

C_j is the junction capacitance (Farads)

L_s is the series inductance (Henrys)

Tunnel Diode Applications

Tunnel diode is used as an amplifier, as one shot multivibrator and as an oscillator.

AVALANCE TRANSIT TIME DEVICES:

Avalanche transit-time diode oscillators rely on the effect of voltage breakdown across a reverse biased p-n junction to produce a supply of holes and electrons. Ever since the development of modern semiconductor device theory scientists have speculated on whether it is possible to make a two-terminal negative-resistance device.

The tunnel diode was the first such device to be realized in practice. Its operation depends on the properties of a forward-biased p - n junction in which both the p and n regions are heavily doped. The other two devices are the transferred electron devices and the avalanche transit-time devices.

The transferred electron devices or the Gunn oscillators operate simply by the application of a dc voltage to a bulk semiconductor. There are no p - n junctions in this device. Its frequency is a function of the load and of the natural frequency of the circuit. The avalanche diode oscillator uses carrier impact ionization and drift in the high-field region of a semiconductor junction to produce a negative resistance at microwave frequencies.

The device was originally proposed in a theoretical paper by Read in which he analyzed the negative-resistance properties of an idealized n - p - i - p diode. Two distinct modes of avalanche oscillator have been observed. One is the IMPATT mode, which stands for ***impact ionization avalanche transit-time*** operation. In this mode the typical dc-to-RF conversion efficiency is 5 to 10%, and frequencies are as high as 100 GHz with silicon diodes.

The other mode is the TRAPATT mode, which represents *trapped plasma avalanche triggered transit* operation. Its typical conversion efficiency is from 20 to 60%. Another type of active microwave device is the BARITT (*barrier injected transit-time*) diode. It has long drift regions similar to those of IMPATT diodes. The carriers traversing the drift regions of BARITT, however, are generated by minority carrier injection from forward-biased junctions rather than being extracted from the plasma of an avalanche region. Several different structures have been operated as BARITT diodes, such as p - n - p , p - n - v - p , p - n -metal, and metal- n -metal. BARITT diodes have low noise figures of 15 dB, but their bandwidth is relatively narrow with low output power.

IMPATT AND TRAPATT DIODE:

Physical Structures

□ theoretical Read diode made of n - p - i - p or p - n - i - n structure has been analyzed. Its basic physical mechanism is the interaction of the impact ionization avalanche and the transit time of charge carriers. Hence the Read-type diodes are called IMPATT diodes. These diodes exhibit a differential negative resistance by two effects:

- 2) The impact ionization avalanche effect, which causes the carrier current $i_o(t)$ and the ac voltage to be out of phase by 90°
- 3) The transit-time effect, which further delays the external current $i_e(t)$ relative to the ac voltage by 90°

The first IMPATT operation as reported by Johnston et al. in 1965, however, was obtained from a simple p - n junction. The first real Read-type IMPATT diode was reported by Lee et al., as described previously. From the small-signal theory developed by Gilden it has been confirmed that a negative resistance of the IMPATT diode can be obtained from a junction diode with any doping profile.

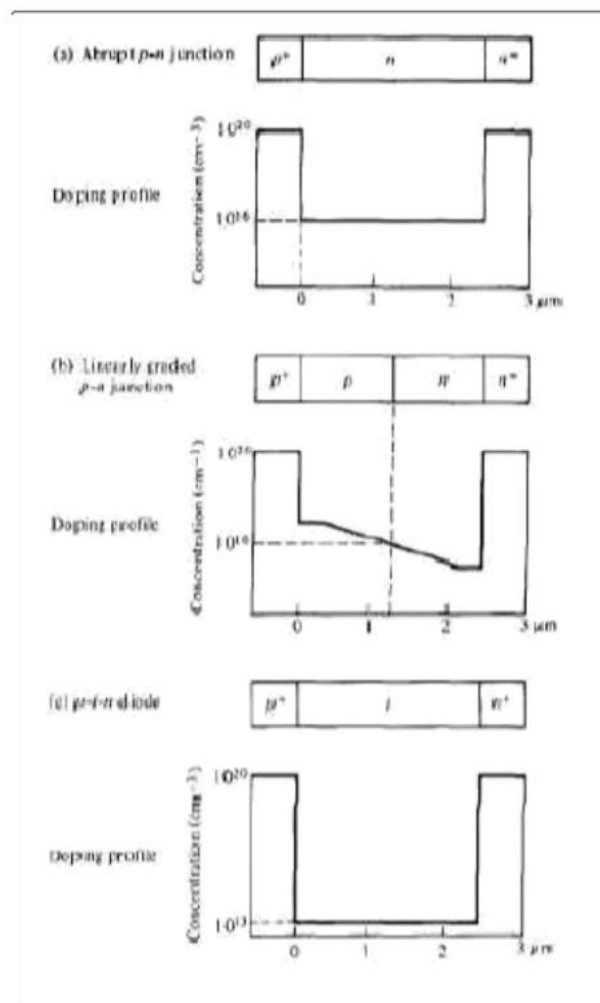
Many IMPATT diodes consist of a high doping avalanching region followed by a drift region where the field is low enough that the carriers can traverse through it without avalanching. The

Read diode is the basic type in the IMPATT diode family. The others are the one-sided abrupt p - n junction, the linearly graded p - n junction (or double-drift region), and the p - i - n diode, all of which are shown in Fig.

The principle of operation of these devices, however, is essentially similar to the mechanism described for the Read diode. Negative Resistance Small-signal analysis of a Read diode results in the following expression for the real part of the diode terminal impedance

$$R = R_i + \frac{2L^2}{v_d \epsilon_s A} \frac{1}{1 - \omega^2/\omega_c^2} \frac{1 - \cos \theta}{\theta}$$

where R_i = passive resistance of the inactive region
 v_d = carrier drift velocity
 L = length of the drift space-charge region
 A = diode cross section
 ϵ_s = semiconductor dielectric permittivity



MICROWAVE MEASUREMENTS

DESCRIPTION OF MICROWAVE BENCH

Introduction:

Electrical measurements encountered in the microwave region of the electromagnetic spectrum are discussed through microwave measurement techniques. This measurement technique is vastly different from that of the more conventional techniques. The methods are based on the wave character of high frequency currents rather than on the low frequency technique of direct determination of current or voltage.

For example, the measurement of power flow in a system specifies the product of the electric and magnetic fields. Whereas the measurement of impedance determines their ratio. Thus these two measurements indirectly describe the distribution of the electric field and magnetic fields in the system and provides its complete description. This is, in fact, the approach to most of the measurements carried out in the microwave region of the spectrum.

Microwave Bench:

The microwave test bench incorporates a range of instruments capable of allowing all types of measurements that are usually required for a microwave engineer. The bench is capable of being assembled or disassembled in a number of ways to suit individual experiments. A general block diagram of the test bench comprising its different units and ancillaries are shown below.

1. Klystron Power Supply:

Klystron Power Supply generates voltages required for driving the reflex Klystron tube like 2k25. It is stable, regulated and short circuit protected power supply. It has built on facility of square wave and saw tooth generators for amplitude and frequency modulation. The beam voltage ranges from 200V to 450V with maximum beam current 50mA. The provision is given to vary repeller voltage continuously from -270V DC to -10V.

Gunn Power Supply:

Gunn Power Supply comprises of an electronically regulated power supply and a square wave generator designed to operate the Gunn oscillator and PIN Modulator. The Supply Voltage ranges from 0 to 12V with a maximum current, 1A.

Gunn oscillator:

Gunn oscillator utilizes Gunn diode which works on the principle that when a DC voltage is applied across a sample of n-type Gallium Arsenide; the current oscillates at microwave frequencies. This does not need high voltage as it is necessary for Klystrons and therefore solid state oscillators are now finding wide applications. Normally, they are capable of delivering 0.5 watt at 10GHz, but as the frequency of operation is increased the microwave output power gets considerably reduced.

4. Isolator:

This unattenuated device permits un attenuated transmission in one direction (forward direction) but provides very high attenuation in the reverse direction {backward direction). This is generally used between the source and rest of the set up to avoid overloading of the source due to reflected power.

5. Variable Attenuator:

The device that attenuates the signal is termed as attenuator. Attenuators are categorized into two categories namely, the fixed attenuators and variable attenuators. The attenuator used in the microwave set is of variable type. The variable attenuator consists of a strip of absorbing material which is arranged in such a way that its profusion into the guide is adjustable. Hence, the signal power to be fed to the microwave set up can be set at the desired level.

6. Frequency Meter:

It is basically a cavity resonator. The method of measuring frequency is to use a cavity where the size can be varied and it will resonate at a particular frequency for given size. Cavity is attached to a guide having been excited by a certain microwave source and is tuned to its resonant frequency. It sucks up some signal from the guide to maintain its stored energy.

Thus if a power meter had been monitoring the signal power at the resonating condition of the cavity it will indicate a sharp dip. The tuning of the cavity is achieved by a micrometer screw and a curve of frequency versus screw setting is provided. The screw setting at which the power indication dip is noted and the frequency is read from the curve.

7. Slotted Section:

To sample the field with in a wave guide, a narrow longitudinal slot with ends tapered to provide smoother impedance transformation and thereby providing minimum mismatch, is milled on the top of broader dimension of wave guide. Such section is known as slotted wave guide section. The slot is generally so many wave lengths long to allow many minima of standing wave pattern to be

covered. The slot location is such that its presence does not influence the field configurations to any great degree. On this Section a probe inserted with in a holder, is mounted on a movable carriage. The output is connected to detector and indicating meter. For detector tuning a tuning plunger is provided instead of a stub.

8. Matched Load:

The microwave components which absorb all power falling on them are matched loads. These consist of wave guide sections of definite length having tapered resistive power absorbing materials. The matched loads are essentially used to test components and circuits for maximum power transfer.

9. Short Circuit Termination:

Wave guide short circuit terminations provide standard reflection at any desired, precisely measurable positions. The basic idea behind it is to provide short circuit by changing reactance of the terminations.

10. VSWR meter:

Direct-reading VSWR meter is a low-noise tuned amplifier voltmeter calibrated in db and VSWR for use with square law detectors. A typical SWR meter has a standard tuned frequency of 100-Hz, which is of course adjustable over a range of about 5 to 10 per cent, for exact matching in the source modulation frequency. Clearly the source of power to be used while using SWR meter must be giving us a 1000-Hz square wave modulated output. The band width facilitates single frequency measurements by reducing noise while the widest setting accommodates a sweep rate fast enough for oscilloscope presentation.

11. Crystal Detector:

The simplest and the most sensitive detecting element is a microwave crystal. It is a nonlinear, non reciprocal device which rectifies the received signal and produces a current proportional to the power input. Since the current flowing through the crystal is proportional to the square of voltage, the crystal is rejoined to as a square law detector. The square law detection property of a crystal is valid at a low power levels (<10 mw). However, at high and medium power level (>10 mw), the crystal gradually becomes a linear detector.

FREQUENCY MEASUREMENT.

Counters and pre-scalers for direct frequency measurement in terms of a quartz crystal reference oscillator are often used at lower frequencies, but they give up currently at frequencies above about 10GHz. An alternative is to measure the wavelength of microwaves and calculate the frequency from the relationship (frequency) times (wavelength) = wave velocity. Of course, the direct

frequency counter will give a far more accurate indication of frequency. For many purposes the 1% accuracy of a wavelength measurement suffices. A resonant cavity made from waveguide with a sliding short can be used to measure frequency to a precision and potential accuracy of $1/Q$ of the cavity, where Q is the quality factor often in the range 1000-10,000 for practical cavities.

"Precision" and "accuracy".

Precision is governed by the fineness of graduations on a scale, or the "tolerance" with which a reading can be made. For example, on an ordinary plastic ruler the graduations may be $1/2\text{mm}$ at their finest, and this represents the limiting precision. Accuracy is governed by whether the graduations on the scale have been correctly drawn with respect to the original standard. For example, our plastic ruler may have been put into boiling water and stretched by 1 part in 20. The measurements on this ruler may be precise to $1/2\text{mm}$, but in a 10 cm measurement they will be inaccurate by $10/20\text{ cm}$ or 5mm, ten times as much. In a cavity wavemeter, the precision is set by the cavity Q factor which sets the width of the resonance. The accuracy depends on the calibration, or even how the scale has been forced by previous users winding down the micrometer against the end stop...

WAVELENGTH MEASUREMENT.

Wavelength is measured by means of signal strength sampling probes which are moved in the direction of wave propagation by means of a sliding carriage and vernier distance scale. The signal strength varies because of interference between forward and backward propagating waves; this gives rise to a standing wave pattern with minima spaced $1/2$ wavelength. At a frequency of 10 GHz the wavelength in free space is 3 cm. Half a wavelength is 15mm and a vernier scale may measure this to a precision of $1/20\text{mm}$. The expected precision of measurement is therefore 1 part in 300 or about 0.33% The location of a maximum is less precise than the location of a minimum; the indicating signal strength meter can be set to have a gain such that the null is very sharply determined. In practice one would average the position of two points of equal signal strength either side of the null; and one would also average the readings taken with the carriage moving in positive and negative directions to eliminate backlash errors. Multiple readings with error averaging can reduce the random errors by a further factor of 3 for a run of 10 measurements.

Measurements of impedance and reflection coefficient.

A visit to your favourite microwave book shows that a measurement of the standing wave ratio alone is sufficient to determine the magnitude, or modulus, of the complex reflection coefficient. In turn this gives the return loss from a load directly. The standing wave ratio may be measured directly using a travelling signal strength probe in a slotted line. The slot in waveguide is cut so that it does not cut any of the current flow in the inside surface of the guide wall. It therefore does not disturb the field pattern and does not radiate and contribute to the loss. In the X band waveguide slotted lines in our lab, there is a ferrite fringing collar which additionally confines the

energy to the guide. To determine the phase of the reflection coefficient we need to find out the position of a standing wave minimum with respect to a "reference plane". The procedure is as follows:- First, measure the guide wavelength, and record it with its associated accuracy estimate. Second, find the position of a standing wave minimum for the load being measured, in terms of the arbitrary scale graduations of the vernier scale. Third, replace the load with a short to establish a reference plane at the load position, and measure the closest minimum (which will be a deep null) in terms of the arbitrary scale graduations of the vernier scale. Express the distance between the measurement for the load and the short as a fraction of a guide wavelength, and note if the short measurement has moved "towards the generator" or "towards the load". The distance will always be less than 1/4 guide wavelength towards the nearest minimum.

Fourth, locate the $r > 1$ line on the SMITH chart and set your dividers so that they are on the centre of the chart at one end, and on the measured VSWR at the other along the $r > 1$ line. (That is, if $VSWR = 1.7$, find the value $r = 1.7$). Fifth, locate the short circuit point on the SMITH chart at which $r = 0$, and $x = 0$, and count round towards the generator or load the fraction of a guide wavelength determined by the position of the minimum. Well done. If you plot the point out from the centre of the SMITH chart a distance "VSWR" and round as indicated you will be able to read off the normalised load impedance in terms of the line or guide characteristic impedance. The fraction of distance out from centre to rim of the SMITH chart represents the modulus of the reflection coefficient $[\text{mod}(\gamma)]$ and the angle round from the $r > 1$ line in degrees represents the phase angle of the reflection coefficient $[\text{arg}(\gamma)]$.

IMPEDANCE MEASUREMENT:

The impedance at any point on a transmission line can be written in the form $R + jx$.
For comparison SWR can be calculated as

$$S = \frac{1 + |R|}{1 - |R|}$$

where reflection coefficient 'R' given as

$$R = \frac{Z - Z_0}{Z + Z_0}$$

Z_0 = characteristics impedance of wave guide at operating frequency.

Z is the load impedance

The measurement is performed in the following way.

The unknown device is connected to the slotted line and the position of one minima is determined. The unknown device is replaced by movable short to the slotted line. Two successive minima positions are noted. The twice of the difference between minima position will be guide wave length. One of the minima is used as reference for impedance measurement. Find the difference of reference minima and minima position obtained from unknown load. Let it be 'd'. Take a smith chart, taking '1' as centre, draw a circle of radius equal to S. Mark a point on circumference of smith chart towards load side at a distance equal to d/λ_g . Join the center with this point. Find the point where it cut the drawn circle. The co-ordinates of this point will show the normalized impedance of load.

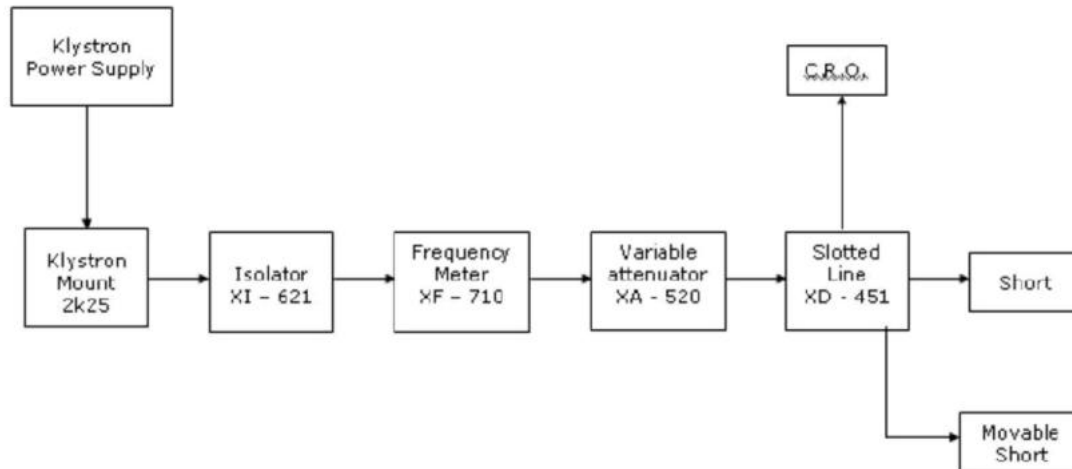


Figure : setup for impedance measurement

Steps:

- ☐ Calculate a set of V_{min} values for short or movable short as load.
- ☐ Calculate a set of V_{min} values for S-S Tuner + Matched termination as a load.

Note: Move more steps on S-S Tuner

- ☐ From the above 2 steps calculate $d = d_1 \sim d_2$
- ☐ With the same setup as in step 2 but with few numbers of turns (2 or 3). Calculate low VSWR. **Note:** High VSWR can also be calculated but it results in a complex procedure.
- ☐ Draw a VSWR circle on a smith chart.
- ☐ Draw a line from center of circle to impedance value (d/λ_g) from which calculate admittance and Reactance ($Z = R + jX$)

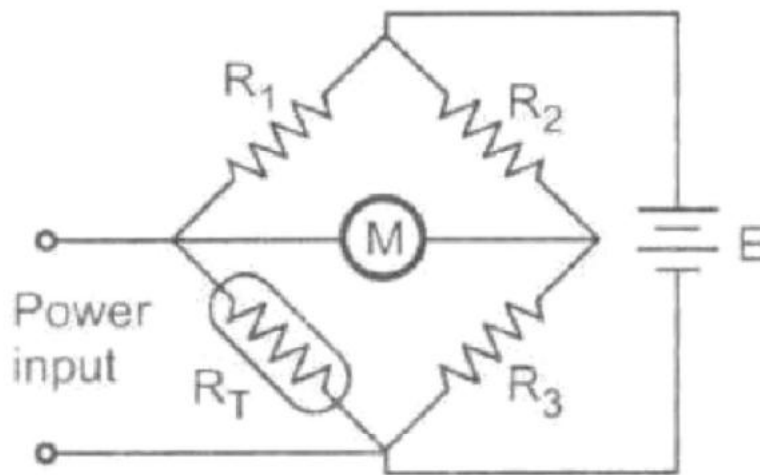
MEASUREMENT OF POWER:

To measure power at high frequencies from 500 MHz to 40 GHz two special type of absorption meters are popularly used. These meters are,

- Calorimeter power meter
- Bolometer power meter

Both these meters use the sensing of heating effects caused by the power signal to be measured.

Introduction to Bolometer power meter:



Bolometer Power meter.

The Bolometer power meter basically consists of a bridge called Bolometer Bridge. One of the arms of this bridge consists of a temperature sensitive resistor. The basic bridge used in Bolometer power meter is shown in the Fig 8.14. The high frequency power input is applied to the temperature sensitive resistor R_T . The power is absorbed by the resistor and gets heated due to the high frequency power input signal. This heat generated causes change in the resistance R_T . This change in resistance is measured with the help of bridge circuit which is proportional to the power to be measured.

The most common type of temperature sensitive resistors are the thermistor and barretter. The thermistor is a resistor that has large but negative temperature coefficient. It is made up of a semiconductor material. Thus its resistance decreases as the temperature increases. The barretter consists of short length of fine wire or thin film having positive temperature coefficient. Thus its resistance increases as the temperature increases. The barretters are very delicate while thermistors are rugged. The bolometer power meters are used to measure radio frequency power in the range 0.1 to 10 mW.

In modern bolometer power meter set up uses the differential amplifier and bridge [or] an oscillator which oscillates at a particular amplitude when bridge is unbalanced. Initially when temperature sensitive resistor is cold, bridge is almost balance. With d.c. bias, exact balance is achieved. When power input at high frequency is applied to R_T , it absorbs power and gets heated.

Due to this its resistance changes causing bridge unbalance. This unbalance is in the direction opposite to that of initial cold resistance. Due to this, output from the oscillator decreases to achieve bridge balance.

MEASUREMENT OF VSWR

High VSWR by Double Minimum Method:

The voltage standing wave ratio of

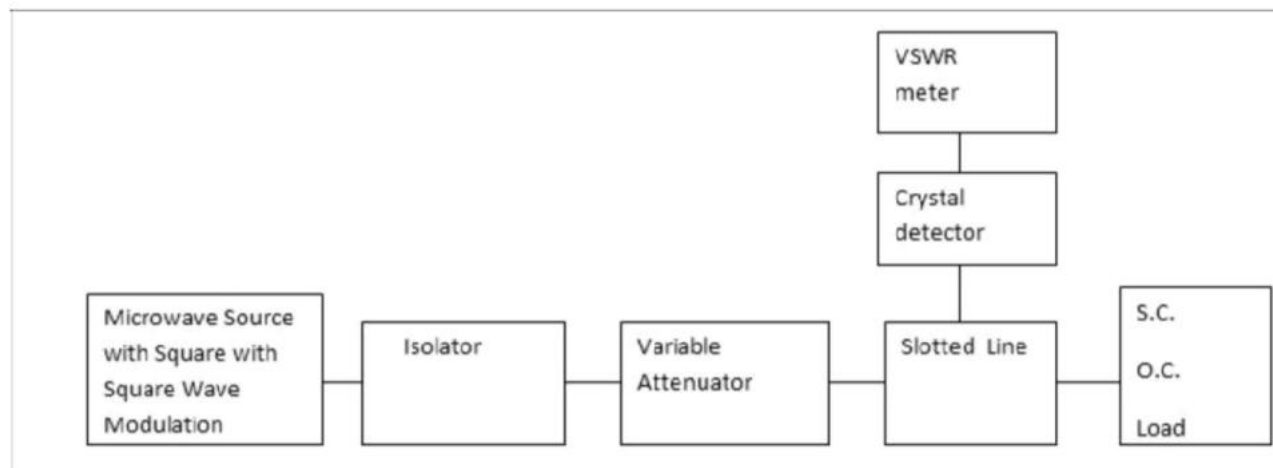
$$VSWR = \frac{V_{max}}{V_{min}}$$

where V_{max} and V_{min} are the voltage at the maxima and minima of voltage standing wave distribution. When the VSWR is high (, the standing wave pattern will have a high maxima and low minima. Since the square law characteristic of a crystal detector is limited to low power, an error is introduced if ≥ 5) V_{max} is measured directly. This difficulty can be avoided by using the 'double minimum method' in which measurements are take on the standing wave pattern near the voltage minimum. The procedure consists of first finding the value of voltage minima. Next two positions about the position of V_{max} are found at which the output voltage is twice the minimum value.

If the detector response is square

$$VSWR = [1 + \frac{1}{\sin^2(\frac{\pi d}{\lambda_g})}]^{\frac{1}{2}}$$

where λ_g is the guide wavelength and d is the distance between the two points where the voltage is 2 V_{min} .



Measurement of high VSWR:

Select “Unmatched Load” to terminate the slotted line by pressing the button.

- ☐ Use slider to fix the value of “Resistance” and “Reactance” of the load.
- ☐ Locate the position of V_{min} and take it as a reference. (If VSWR meter is used in actual experiment, set the output so that meter reads 3dB).
- ☐ Move the slider (probe of slotted line) along the slotted line on either side of V_{min} so that the reading is 3 db below the reference i.e. 0 db. Record the probe positions and obtain the distance between the two. Determine the VSWR using equation (2).
- ☐ The simulated value for VSWR can be seen by clicking the buttons “Technique used to calculate VSWR 1 & 2”.
- ☐ Then match the calculated value with the value displayed in the simulated VSWR

