



Applied Physics

Unit I-Quantum Mechanics

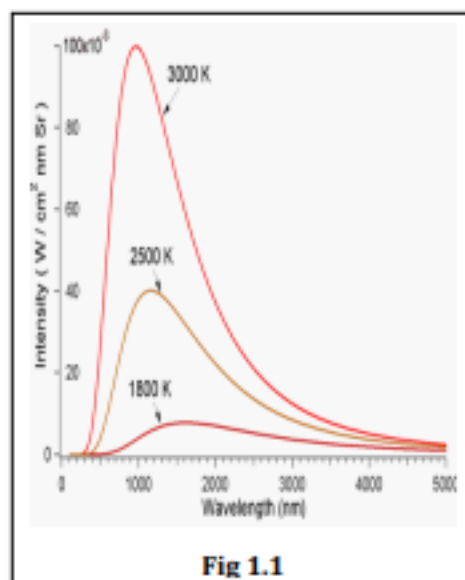
Blackbody radiation

A blackbody is a material which absorbs radiation of all possible wavelengths incident on it and emits that radiation when heated. The radiation emitted by such a body is called blackbody radiation. In practice there is no perfect blackbody. Carbon black and platinum black are materials with 99% emission. A recent addition to these is VANTA black with 99.96% emission.

Spectral Distribution of blackbody radiation

The distribution of radiation intensity or radiation energy density emitted by a blackbody in a given wavelength range is known as spectral distribution of a blackbody. Lummer and Pringsheim have carried out various experiments on the spectral distribution of blackbody radiation and came out with some important findings as mentioned below.

1. A blackbody can emit radiation at all possible wavelengths and the radiation intensity is not uniformly distributed among these wavelengths as shown in fig 1.1.
2. For a given temperature, the intensity of the radiation increases with increase in wavelength, becomes maximum at a particular wavelength. With further increase in the wavelength, the intensity of radiation decreases.
3. The radiation intensity for a given wavelength increases with increasing temperature.
4. The wavelength corresponding to the maximum intensity represented by the peak of the curve shifts towards shorter wavelengths as the temperature increases.
5. The area under the curve gives the total energy emitted at a particular temperature. It was observed that this area is directly proportional to the 4th power of the temperature of the blackbody.



Laws of blackbody radiation

Stefan's Fourth power law

Stefan and Boltzmann in 1884 showed that the radiation energy per unit volume due to all wavelengths is proportional to the fourth power of absolute temperature of the blackbody known as Stefan's fourth power law. This is given as $E = \sigma T^4$, where σ is called as Stefan-Boltzmann constant which is equal to $\sigma = 5.67 \times 10^{-8} \text{ W/m}^2 \cdot \text{K}^4$

Wien's Displacement Law

The product of wavelength corresponding to the maximum intensity and the absolute temperature is constant,

$$\text{i.e., } \lambda_{\text{max}} \cdot T = \text{constant} = 2.898 \times 10^{-3} \text{ m-K}$$

Wien's Radiation Law

In order to explain the observed spectral distribution, Wein in 1893 showed that the energy density in the wavelength range λ and $\lambda + d\lambda$ is given by $E_\lambda d\lambda = C_1 \lambda^{-5} e^{\left(\frac{-C_2}{\lambda T}\right)} d\lambda$, where C_1 & C_2 are constants, T is the absolute temperature of the blackbody.

Wien's formula works well for shorter wavelengths but has considerable deviations for long wavelengths and high temperatures.

Rayleigh-Jeans Formula

In 1900, Rayleigh and Jeans used a more vigorous method and obtained the following formula for the energy distribution as, $E_\lambda d\lambda = 8\pi kT \lambda^{-4} d\lambda$.

It was found that Rayleigh-Jeans formula is in coincidence with the experimental results only in the longer wavelength region.

Planck's Quantum theory

In 1900, Max Planck suggested that the correct results on the spectral distribution of blackbody radiation can be obtained by considering the energy of the oscillating particles to be discrete rather than continuous. He derived the radiation law using the following assumptions.

1. A blackbody radiator contains simple harmonic oscillators of molecular dimensions which can vibrate with all possible frequencies.
2. The frequency of radiation emitted by an oscillator is the same as the frequency of its vibration.
3. The emission or absorption of radiation energy by an oscillator is not continuous but in multiples of $h\nu$ called quantum. Here ' ν ' is the frequency of the emitted or absorbed radiation and h is Planck's constant.
4. This implies that the exchange of energy between radiation and matter takes place in discrete proportions of $0, h\nu, 2h\nu, 3h\nu, \dots, nh\nu$.

Planck's Radiation Law

Let there be $N_0, N_1, N_2, \dots, N_r, \dots$ etc. oscillators having energy $0, \epsilon, 2\epsilon, 3\epsilon, \dots, r\epsilon, \dots$ etc. respectively. If N is the total number of Planck's oscillators and E be their total energy, then

$$\begin{aligned} \text{Total no. of oscillators} &= N = N_0 + N_1 + N_2 + \dots + N_r + \dots \\ \& \quad \text{Total energy} &= E = 0 + \epsilon N_1 + 2\epsilon N_2 + 3\epsilon N_3 + \dots + r\epsilon N_r + \dots \end{aligned}$$

The average energy of a Planck's oscillator can be written as $\bar{\epsilon} = \frac{E}{N} \rightarrow [1]$

According to Maxwell's distribution formula, the no. of oscillators having energy $r\epsilon$ is given by,

$$N_r = N_0 e^{\left(\frac{-r\epsilon}{kT}\right)}, \text{ where } k \text{ is Boltzmann's constant.}$$

$$\text{Total number of oscillators} = N = \frac{N_0}{\left\{1 - e^{\left(\frac{-\epsilon}{kT}\right)}\right\}}, \& \quad \text{Total energy} = E = \left[\frac{\epsilon N_0 e^{\left(\frac{-\epsilon}{kT}\right)}}{\left\{1 - e^{\left(\frac{-\epsilon}{kT}\right)}\right\}^2} \right] \rightarrow [2]$$

$$\text{Now the average energy of the oscillator is given by, } \bar{\epsilon} = \frac{E}{N} = \frac{\epsilon}{\left\{ \exp\left(\frac{\epsilon}{kT}\right) - 1 \right\}} = \frac{h\nu}{\exp\left(\frac{h\nu}{kT}\right) - 1} \rightarrow [3]$$

$$\text{The no. of oscillators / unit volume, in the range between } \nu \text{ and } \nu + d\nu \text{ is given as, } f = \frac{8\pi\nu^2}{c^3} d\nu \rightarrow [4]$$

The radiation energy density, i.e. the energy per unit volume within the interval $d\nu$ is given as the **product of no. of oscillators per unit volume** in the same interval and **the average energy of oscillator**.

$$\text{i.e., } E_\nu d\nu = \frac{8\pi\nu^2}{c^3} d\nu \times \frac{h\nu}{\exp\left(\frac{h\nu}{kT}\right) - 1} \Rightarrow E_\nu d\nu = \frac{8\pi h\nu^3}{c^3} \frac{1}{\exp\left(\frac{h\nu}{kT}\right) - 1} d\nu \rightarrow [5]$$

Equation [5] is known as Planck's radiation law.

$$\text{It is also expressed in terms of wavelength as } E_\lambda d\lambda = \frac{8\pi hC}{\lambda^5} \frac{1}{\exp\left(\frac{hC}{\lambda kT}\right) - 1} d\lambda \rightarrow [6]$$

This formula agrees well with the experimental curves throughout the whole range of wavelengths.

Photoelectric Effect

The emission of electrons from a metal plate when illuminated by light or any other radiation of suitable wavelength (or frequency) is called photoelectric effect. The emitted electrons are called as photo electrons, and the phenomenon is called as photoelectric effect.

Experimental study of photoelectric effect

A simple experimental arrangement to study the photoelectric effect is shown in fig 1.2. The apparatus consists of two photosensitive surfaces A and B, enclosed in an evacuated quartz bulb. The plate A is connected to the negative terminal of a potential divider while the plate B is connected to the positive terminal through a galvanometer G or micro ammeter μA . In the absence of any light there is no flow of current and hence there is no deflection in the micro-ammeter. But when monochromatic light is allowed to fall on plate A, current starts flowing in the circuit which is indicated by galvanometer. The current is known as photocurrent. This shows that when the light falls on the metallic surface, electrons are ejected. The number of electrons emitted, and their kinetic energy depend on the following factors.

- 1) The potential difference between plates A and B
- 2) The intensity of incident radiation
- 3) The frequency of incident radiation
- 4) The photo metal used.

1) The effect of potential difference

For a given photo metal, keeping the intensity and frequency of incident radiation fixed, the photo current varies with the potential difference V between the two plates as shown in fig 1.3.1. When the positive potential of the plate B is increased, photoelectric current is also increased and reaches a maximum value. **This value of the current is known as saturation current.** Further increase in the potential hardly produces any appreciable increase in current. If the potential of the plate B is made negative and increased further, photocurrent decreases and finally becomes zero at a particular value. **This retarding potential at which the photoelectric current becomes zero is called as stopping potential.**

Fig 1.2

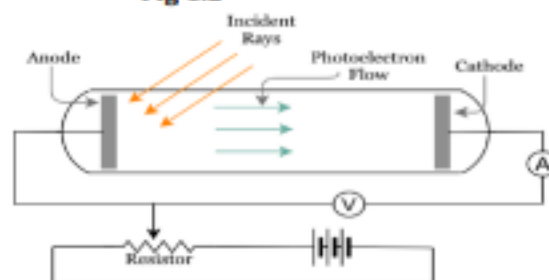


Fig 1.3.1

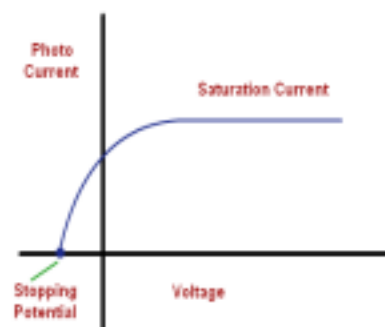


Fig 1.3.2

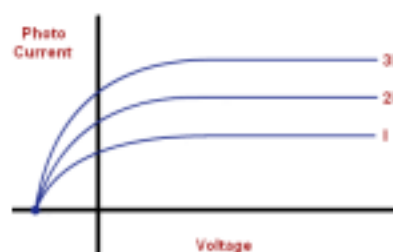
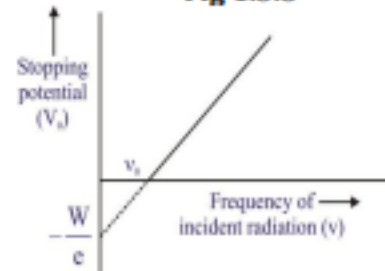


Fig 1.3.3



2) Effect of intensity of incident radiation

Fig 1.3.2 shows the variation of photocurrent with the potential difference for different intensities of incident radiation of constant frequency. It shows that,

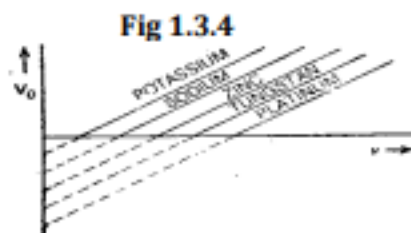
- i) the saturation current is proportional to the intensity of incident radiation
- ii) the stopping potential is independent of intensity of incident radiation

3) Effect of frequency of incident radiation

Fig.1.3.3 shows the variation of stopping potential with frequency of incident light for a constant intensity of incident radiation. Here stopping potentials are measured for different frequencies. The graph shows that at frequency ν_0 , the stopping potential is zero. This is known as threshold frequency and it is the minimum frequency of the incident radiation which can cause photoelectric emission i.e. this frequency is just able to liberate electrons without giving them additional energy.

4) Effect of photo metal

Fig. 1.3.4 shows graph between stopping potential V_0 and frequency for a few photo-metals. It shows that all the lines have the same slope but their intersections with frequency axis are different. Thus, we conclude the threshold frequency is a characteristic constant of the photo-metal.



Einstein's Photo-electric Equation

Following Planck's idea that light consists of photons Einstein proposed an explanation of photoelectric effect as early as 1905. According to this, in photoelectric effect one photon is completely absorbed by one electron, which thereby gains the quantum of energy and may be emitted from the metal. The photon's energy is utilized in the following way.

- i) A part of it is used to free the electron from the metal surface. This energy is known as **photoelectric work function** of the metal, denoted by W_0 .
- ii) The other part is used in giving kinetic energy ($\frac{1}{2}mv^2$) to the electron where v , is the velocity of the emitted electron.

Thus,
$$h\nu = W_0 + \frac{1}{2}mv^2$$

The above equation is known as **Einstein's photoelectric equation**.

de Broglie Concept

Louis de-Broglie in 1924 extended the wave-particle parallelism of radiation to all fundamental material particles (like electrons, protons, neutrons, atoms and molecules) stating that "every material particle in motion is associated with a wave, whose wavelength is inversely proportional to its momentum". The corresponding wavelength is called as **de-Broglie wavelength**, which is given

as, $\lambda = \frac{h}{mv}$, where h is Planck's constant, m and v are the mass and velocity of the particle.

Expression for de Broglie wavelength

Using Planck's law and Einstein's mass-energy relation, de Broglie has estimated the wavelength associated with an electron subjected to an accelerating potential of V as $\lambda = \frac{12.27}{\sqrt{V}} \text{ \AA}$.

Davisson-Germer Experiment (1927)

Davisson and Germer were studying scattering of electrons from a metal target measuring the intensity of scattered electron beams in different directions. The apparatus essentially consists of 3 major components.

1. Electron Gun
2. Crystal Target
3. Electron Collector

The main functionality of the **electron gun** is to produce a fine beam of fast-moving electrons. This essentially consists of a tungsten filament **F** connected to a low-tension battery **B1**, produces electrons. These electrons are accelerated to desired velocity by applying suitable potential using a high-tension source **B2**. The accelerated electrons are collimated into a fine beam by a collimator **C**.

The fast-moving beam of electrons is made to strike the **Nickel target** capable of rotating about an axis perpendicular to the plane. The electrons are scattered by the atomic planes of the crystal in all directions.

The **electron collector** is a metal cylinder that collects electron beams scattered in different directions. This can slide on a semicircular scale which measures scattering angle. The metal cylinder is connected to a sensitive galvanometer whose deflection is proportional to the intensity of the collected electron beam. The whole instrument is kept in an evacuated chamber as shown in fig 1.4 (a). An electron beam accelerated by a potential of **54 V**, scattered at angle of **50°** from the crystal target recorded **maximum intensity** as shown in fig 1.4 (b).

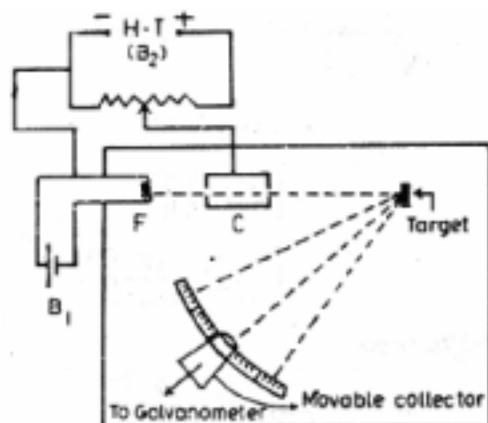


Fig 1.4 a: Experimental Arrangement

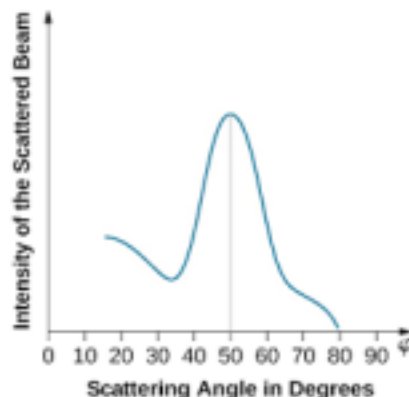


Fig 1.4 b

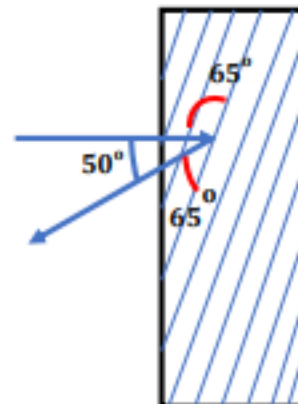


Fig 1.4 c

For a scattering angle of 50° , the Bragg angle is 65° as shown in fig 1.4(c). The interplanar spacing (d) for nickel is found out to be $d = 0.91 \text{ \AA}$.

The wavelength of scattered beams is calculated using, **Bragg's law**, $2d \sin \theta = n \lambda$

$$\text{For, } n=1, \lambda = 2 \times 0.91 \times 10^{-10} \times \sin 65^\circ = 1.65 \text{ \AA}$$

de Broglie wavelength for an electron accelerated by 54 V, is

$$\lambda = 12.25 / \sqrt{V} = (12.25 / \sqrt{54}) = 1.66 \text{ \AA}.$$

Both the values of λ are in great agreement. There are two conclusions for this experiment.

One, electrons also undergo diffraction like electromagnetic waves and **two**, the wave nature of electrons is as predicted by de Broglie.

Heisenberg Uncertainty Principle

Heisenberg's uncertainty principle can be stated as follows:

"It is impossible to determine **precisely and simultaneously** the **values** of both members of a **particular pair of physical variables** that describe the motion of a microscopic system".

If Δx is error in the measurement of position of the particle along X-axis Δp is the error in the measurement of momentum, then, $(\Delta x) \cdot (\Delta p) \geq h/2$, where $h = h/2\pi$.

A particle can be exactly located ($\Delta x \rightarrow 0$) only at the expense of an infinite momentum ($\Delta p \rightarrow \infty$). There are uncertainty relations between position and momentum, energy and time, and angular momentum and angle. If the time during which a system occupies a certain state is not greater than Δt , then the energy of the state cannot be known within ΔE , i.e. $(\Delta E) (\Delta t) \geq h/2$.

Schrodinger Wave Equation

The differential equation representing the matter waves is known as Schrödinger's Wave Equation. Schrodinger in 1926 has introduced a wave function Ψ to explain the motion of a microscopic particle. The wave function Ψ for a particle moving freely in positive x-direction is given as

$$\psi(x, t) = Ae^{-i(\omega t - kx)} \quad \text{Here } \omega = 2\pi\nu \quad \text{and } \nu = v/\lambda.$$

The above equation is valid only for a free particle. However, for a particle under certain restrictions we require a differential equation in Ψ . Differentiating Ψ with respect to x twice,

$$\frac{\partial \psi}{\partial x} = Ae^{-i(\omega t - kx)}(ik)$$

$$\text{and } \frac{\partial^2 \psi}{\partial x^2} = Ae^{-i(\omega t - kx)}(ik)^2$$

$$\text{i.e. } \frac{\partial^2 \psi}{\partial x^2} = -(k)^2 \psi$$

$$\Rightarrow \frac{\partial^2 \psi}{\partial x^2} + k^2 \psi = 0, \quad \text{where } k = \frac{2\pi}{\lambda}$$

→[1]

Since $\lambda = \frac{h}{p}$ (from de Broglie equation), we have $k^2 = \frac{4\pi^2}{\lambda^2} = \frac{4\pi^2}{(\frac{h}{p})^2} = \frac{4\pi^2}{h^2} p^2$

$$\begin{aligned} \therefore \frac{\partial^2 \psi}{\partial x^2} + \frac{4\pi^2}{h^2} p^2 \psi &= 0 \quad \rightarrow [2] \\ \Leftrightarrow \frac{\partial^2 \psi}{\partial x^2} + \frac{1}{h^2} \frac{p^2}{2m} \cdot 2m \psi &= 0 \quad (\because \hbar = \frac{h}{2\pi}) \end{aligned}$$

The total energy of the particle $E = \text{K.E.} + \text{P.E.} = E = \frac{1}{2}mv^2 + V(x) = \frac{p^2}{2m} + V$

$$\begin{aligned} \therefore \frac{p^2}{2m} &= E - V \\ \Leftrightarrow \frac{\partial^2 \psi}{\partial x^2} + \frac{2m}{h^2} (E - V) \psi &= 0 \quad \rightarrow [3] \end{aligned}$$

Equation [3] is known as **Schrodinger Time Independent Wave Equation** in one dimension.

Born's interpretation of Wave Function Ψ

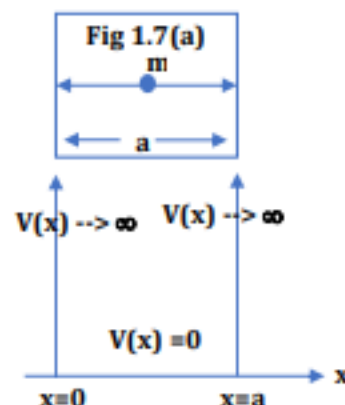
1. The solution of Schrodinger's wave equation, the wave function is assumed to give the probability of finding the particle. According to Max Born, wave function Ψ can not be an observable quantity to predict the presence of a particle in space at a given instant of time since it is both complex and negative.
2. The square of the absolute magnitude of Ψ which is $|\Psi|^2$, called as **probability density** (where $|\Psi|^2 = \Psi(x)\Psi(x)^*$) can be a quantity to find the particle's presence at any instant of time in space since it is real and positive.
3. The probability of finding a particle in a volume dx, dy, dz is $|\Psi|^2 dx dy dz$. Since the particle can be certainly found somewhere in space $\int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} \psi^*(x) \psi(x) dx dy dz = 1$.

A wave function Ψ satisfying this relation is called a normalized wave function.

Particle in one-dimensional Potential box (Application to SWE)

A small distance in which a particle is confined to move back and forth between the extremes with a constant potential energy is known as a one-dimensional potential box. As shown in fig 1.7(a), let us consider a particle of mass m moving back and forth between two infinitely hard walls at $x=0$ and at $x=a$, (in a width of a) with zero potential energy. The particle does not lose energy while colliding with the walls and has no chance of penetrating through them.

$$V(x) = 0, 0 \leq x \leq a \quad V(x) = \infty, 0 > x > a \quad \psi(x) = 0, 0 \geq x \geq a$$



Applying Schrodinger wave equation on to the particle's motion, we have

$$\frac{d^2 \psi}{dx^2} + \frac{2mE}{\hbar^2} \psi = 0, \because V(x) = 0 \text{ for } 0 \leq x \leq a \quad \rightarrow [1]$$

$$\frac{d^2 \psi}{dx^2} + k^2 \psi = 0, \text{ where, } k = \sqrt{\frac{2mE}{\hbar^2}} \quad \rightarrow [2]$$

The general solution of equation [2] is of the form, $\psi(x) = A \sin kx + B \cos kx$, where A and B are arbitrary constants. The values of these arbitrary constants can be determined by applying the following boundary conditions. As the probability of finding the particle outside the box is zero, we have,

$$\text{i) } \psi(x) = 0 \text{ at } x = 0 \Rightarrow B = 0 \Rightarrow \psi(x) = A \sin kx \quad \rightarrow [3]$$

$$\text{ii) } \psi(x) = 0 \text{ at } x = a \Rightarrow A \sin ka = 0 \Rightarrow \sin ka = 0 \quad (\because A \neq 0) \quad \rightarrow [4]$$

$$\Rightarrow k = \frac{n\pi}{a} \text{ where } n=1,2,3, \dots \quad \rightarrow [5]$$

Equation [5] is the necessary condition for the solution of the wave equation to exist. Here n can't be zero, as $n=0$, $k=0$ then both E and Ψ and thus $|\Psi|^2$ become zero everywhere inside the potential box which is not possible. So, the least value of n is 1.

Substituting the value of k from equation [5] in equation [2], we get,

$$E = \frac{k^2 \hbar^2}{2m} \Rightarrow E = \left(\frac{n^2 \pi^2}{a^2} \right) \frac{1}{2m} \left(\frac{h^2}{4\pi^2} \right) \Rightarrow E_n = \frac{n^2 h^2}{8ma^2} \rightarrow [6]$$

From equation [6], we have

1. The lowest energy of the particle (for $n=1$) is $E_1 = \frac{h^2}{8ma^2}$, called as Zero Point Energy, and n here is called as a quantum number.
2. For every n , the particle possesses a particular energy E_n called as an eigen value and a wavefunction Ψ_n , known as an eigen function. The energies exhibited by the particle are discrete corresponding to $n=1, 2, 3$ and so on.
3. The spacing between n^{th} and $(n+1)^{\text{th}}$ energy levels is $(2n+1)E_1$.

By normalization, i.e., $\int_0^a |\psi(x)|^2 dx = 1 \Rightarrow \int_0^a A^2 \sin^2 \frac{n\pi x}{a} dx = 1 \Rightarrow A = \sqrt{\frac{2}{a}}$

$$\therefore \psi(x) = \sqrt{\frac{2}{a}} \sin \frac{n\pi x}{a}, \text{ for } 0 < x < a$$

The wave functions for the first three values of n are shown in fig 1.7(b). It is clear that the wave function ψ_1 has two nodes at $x=0$ and $x=a$. The wave function ψ_2 has three nodes at $x=0$ and $x=a/2$ and at $x=a$. Similarly, wave function ψ_n will have $(n+1)$ nodes. Here each node indicates the position of maximum energy of the particle.

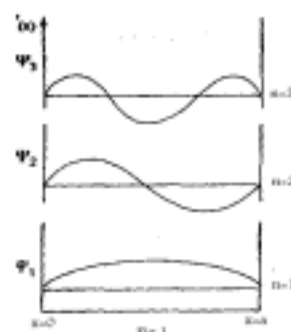


Fig 1.7(b)

Most probable positions

The probability of finding the particle in a small distance dx is $p(x)dx = |\psi_n|^2 dx = \frac{2}{a} \sin^2 \frac{n\pi x}{a} dx$

The probability density $P(x)$ is maximum if $\frac{n\pi x}{a} = \frac{\pi}{2}, \frac{3\pi}{2}, \frac{5\pi}{2}, \frac{7\pi}{2}, \dots$

$$\Rightarrow x = \frac{a}{2n}, \frac{3a}{2n}, \frac{5a}{2n}, \frac{7a}{2n}, \dots \text{ and}$$

As shown in fig 1.7(c),

The most probable position for $n=1$ is $x = \frac{a}{2}$

The most probable positions for $n=2$ are $x = \frac{a}{2}$ and $x = \frac{3a}{2}$

The most probable positions for $n=3$ are $x = \frac{a}{2}, x = \frac{3a}{2}$ and $x = \frac{5a}{2}$

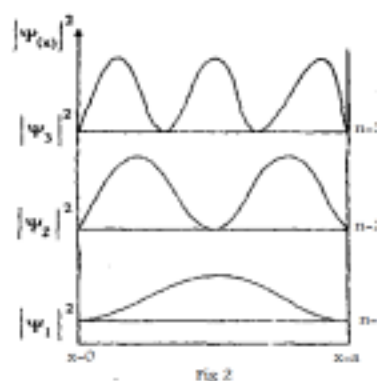


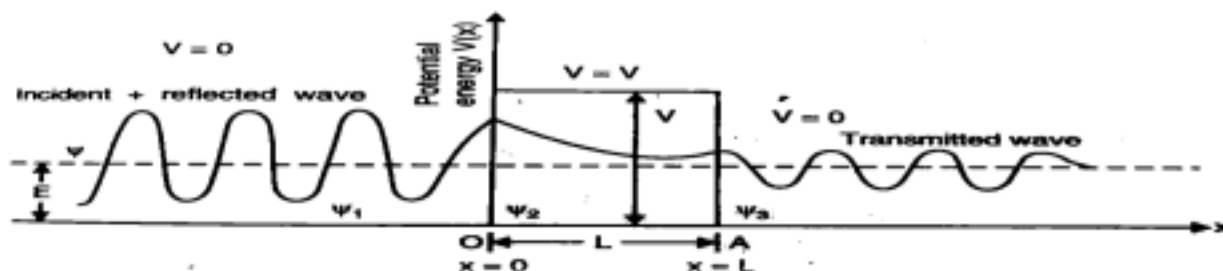
Fig 1.7(c)

Rectangular Potential Barrier

Consider a beam of particles of kinetic energy E incident from the left on a potential barrier of height V_0 and width L (OA) as shown in the figure. $V_0 > E$ and on both sides of the barrier $V=0$ which means that no forces act up on the particles there.

This potential is described by

$$\begin{aligned} V &= 0 & \text{for } x < 0 \text{ (region I)} \\ V &= V_0 & \text{for } 0 < x < L \text{ (region II)} \\ V &= 0 & \text{for } x > L \text{ (region III)} \end{aligned}$$



Writing the Schrodinger wave equations for the 3 regions

$$\text{For region I,} \quad \frac{\partial^2 \psi_1}{\partial x^2} + \frac{2m}{\hbar^2} E \psi_1 = 0 \quad \rightarrow [1]$$

$$\text{For region II,} \quad \frac{\partial^2 \psi_2}{\partial x^2} + \frac{2m}{\hbar^2} (E - V_0) \psi_2 = 0 \quad \rightarrow [2]$$

$$\text{For region III,} \quad \frac{\partial^2 \psi_3}{\partial x^2} + \frac{2m}{\hbar^2} E \psi_3 = 0 \quad \rightarrow [3]$$

Let ψ_1 , ψ_2 and ψ_3 be the respective wavefunctions in regions I, II and III as indicated in the figure.

$$\psi_1 = A e^{i a x} + B e^{-i a x} \quad \rightarrow [4]$$

$$\psi_2 = F e^{\beta x} + G e^{-\beta x} \quad \rightarrow [5]$$

$$\psi_3 = C e^{-i a x} + D e^{i a x} \quad \rightarrow [6]$$

Where the constants A, B, C, D, F and G are the amplitudes of the corresponding components which are as given below. A & B represent the amplitudes of the incident and the reflected waves in region 1, F & G represent the amplitudes of the wave penetrating the barrier and the reflected wave in region 2 and C represents the amplitude of the transmitted wave into region 3.

Since the probability density associated with a wave function is proportional to the square of the amplitude of the wave, the barrier transmission and the reflection coefficients can be interpreted as

$$T = \frac{|C|^2}{|A|^2} \text{ and as } R = \frac{|B|^2}{|A|^2}$$

If the barrier is **high** compared to the total energy of the particle and **thin** compared to the wavelength, then transmission coefficient becomes, $T = 16 \frac{E}{V_0} \left(1 - \frac{E}{V_0}\right) \exp \left[-\frac{2L}{\hbar} \sqrt{2m(V_0 - E)} \right]$, where L is the physical thickness of the barrier. Here T is also called as **penetrability** of the barrier which represents the **probability** that a particle incident on the barrier from one side will appear on the other side, which is **zero** classically but a **finite quantity** quantum mechanically. Thus, if a particle with energy E is incident on a thin energy barrier of height (V_0) greater than E, then there is a finite probability of the particle penetrating the barrier. This phenomenon is known as tunnel effect.

Free Electron Theory

We know that metals are good conductors of electricity due to the availability of free electrons in them. The explanation on the behaviour of these free electrons in metals has evolved in 3 stages. They are,

1. Classical free electron theory (CFET) proposed by **Lorentz and Drude** in 1900
2. Quantum free electron theory (QFET) proposed by **Sommerfeld** in 1927
3. Band Theory of Solids (BTS) proposed by **Bloch** in 1928

Classical Free Electron Theory (CFET)

This theory was proposed by Lorentz and Drude in 1900, which explains the behaviour of free electrons in a metal, with the following **assumptions**.

1. The free electrons move freely throughout the dimensions in all the directions like the molecules of a perfect gas in a container and obey the laws of kinetic theory of gases.
2. The free electrons moving randomly may collide with the ions of the lattice or with each other and these collisions are purely elastic in nature.
3. The velocity distribution of these electrons obeys the classical Maxwell-Boltzmann statistics.
4. The free electrons move in between the metal ions in a uniform potential field.
5. When an external electric field is applied, the electrons move in the direction opposite to the field.

Merits:

1. It verifies ohms law.
2. It could derive Wiedemann-Franz law.
3. It explains electrical and thermal conductivity of metals.
4. It also explains the optical properties of metals such as opacity and luster.

Demerits

1. It is a macroscopic theory.
2. This theory could not explain Compton & Photoelectric effects and para & ferromagnetisms.
3. The predicted values of specific heat of electrons did not match the experimental results.
4. The ratio of thermal to electrical conductivity is not constant at all temperatures as predicted by this theory.

Quantum Free Electron Theory (QFET)

This theory was proposed by Sommerfeld in 1927, according to this, the free electrons in a metal obey the laws of quantum physics. The following are important **assumptions** of this theory.

1. Only 1% of all available free electrons (those near to fermi surface only) participate in electrical conduction.
2. Electron in motion is associated with wave nature.
3. Even though electrons have freedom to move throughout the metal, they will be under some binding by the lattice and thus exhibit discrete energy values.
4. The energy distribution of the electrons is in accordance with Fermi-Dirac statistics and obeys Pauli's exclusion principle.
5. The free electrons move in between the metal ions in a uniform potential field.

Merits

1. This theory could satisfactorily explain the electronic specific heat in solids.
2. It could predict the exact value of the Lorentz number.
3. It could also explain the phenomenon of superconductivity, blackbody radiation, photoelectric effect, and Compton effect satisfactorily.
4. It could satisfactorily explain the cause for para and ferromagnetisms.

Demerits

1. This theory could not classify materials based on their electrical conductivity.
2. It could not explain the cause for higher conductivity of some conductors over the others.
3. It could not explain the negative temperature coefficient of resistivity in certain solids.
4. It could not explain the behaviour of semiconductors.

Bloch's Theorem

Bloch's theorem gives solutions to the Schrödinger wave equation for an electron moving in a periodic potential field. The Schrödinger equation for a particle of mass m , moving in a one-dimensional periodic potential is given by $\frac{\partial^2 \psi}{\partial x^2} + \frac{2m}{\hbar^2} (E - V) \psi = 0$, where E and Ψ are the total energy and the wave function of the electron.

In a periodic potential field, the particle's potential energy is periodic depending on the periodicity (a) of the crystal, so that $V(x) = V(x+a)$, then the solution of the Schrodinger Wave equation is of the form $\psi(x) = u_k(x) e^{ikx}$, where $u(x) = u(x+a)$, k is the propagation vector. Here e^{ikx} is a plane wave solution modulated by the periodic function $u_k(x)$.

Kronig-Penney Model

The Kronig-Penney model describes the motion of the electron in a periodic potential. The potential energy of an electron moving in a periodic potential depends not only on its position but also on the periodicity of the crystal. According to this model the potential energy of an electron is zero near the positive ion core and maximum when it is between adjacent ions separated by a spacing a as shown in figure 1.4(a).

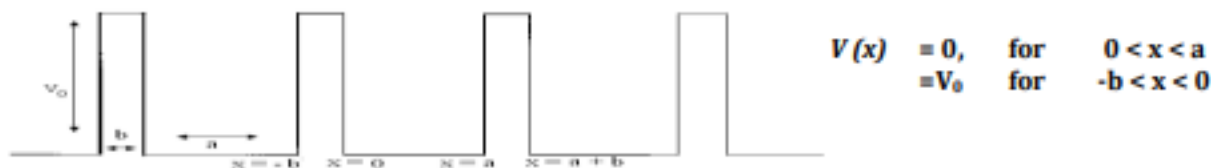


Fig 1.4(a)

Although this model is highly artificial, it illustrates many characteristic features of the behavior of an electron in a periodic lattice. The time independent Schrodinger wave equations for the two regions are

$$\frac{\partial^2 \psi}{\partial x^2} + \frac{2m}{\hbar^2} E \psi = 0 \quad \text{for} \quad 0 < x < a \quad \rightarrow (1)$$

$$\frac{\partial^2 \psi}{\partial x^2} + \frac{2m}{\hbar^2} (E - V_0) \psi = 0 \quad \text{for} \quad (-b < x < 0) \quad \rightarrow (2)$$

Assuming that the potential energy of an electron is less than the potential V_0 , we can write

$$\frac{\partial^2 \psi}{\partial x^2} + \alpha^2 \psi = 0, \text{ where } \alpha = \sqrt{\frac{2mE}{\hbar^2}} \quad \text{for} \quad 0 < x < a \quad \rightarrow (3)$$

$$\frac{\partial^2 \psi}{\partial x^2} - \beta^2 \psi = 0, \text{ where } \beta = \sqrt{\frac{2m(V_0 - E)}{\hbar^2}} \quad \text{for} \quad -b < x < 0 \quad \rightarrow (4)$$

The solution for the above equations can be expressed using Bloch's theorem, as $\psi(x) = u_k(x) e^{ikx}$

Applying the four boundary conditions to the solutions of above equations, four linear equations in arbitrary constants A, B, C and D are obtained. In order to have the non-vanishing solutions to the SWEs, the determinant of the matrix consisting of the coefficients of A, B, C, and D when made equal to zero, we get,

$$P \frac{\sin \alpha a}{\alpha a} + \cos \alpha a = \cos ka, \text{ where } P = \left(\frac{mV_0 ab}{\hbar^2} \right) \quad \rightarrow (5)$$

Here P is called as **scattering power** of the potential barrier.

Equation (5) will be satisfied only for those values of αa , for which the LHS lies between +1 and -1. Plotting the variation of LHS of equation (7) with αa , for $P = (3\pi/2)$, we get a curve as shown in fig 1.4(b), from which the following points can be accessed.

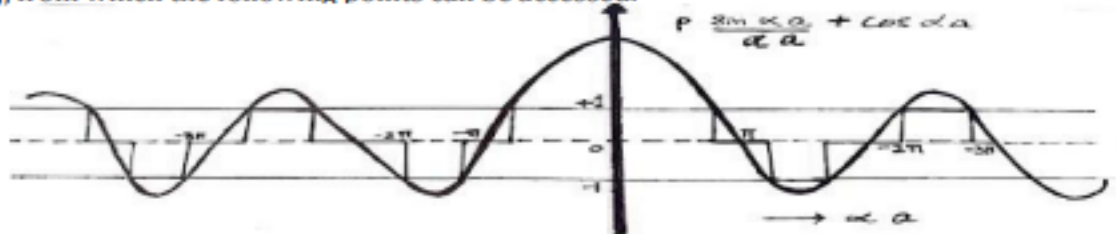


Fig 1.4(b)

Salient Features

- 1) The portions where the curve lies between +1 and -1 indicate the range of αa values for which the wave mechanical solutions exist.
- 2) The motion of an electron in a periodic potential is characterized by the bands of allowed energy separated by forbidden regions.
- 3) As αa increases, the width of the allowed energy band also increases and the width of the forbidden band decreases.
- 4) If $V_0 b$ is large, P is large then the function described by LHS of equation (5) crosses +1 and -1 regions at a steeper angle, so that the allowed energy bands get narrower and forbidden bands get wider.

The two extreme cases of P are as discussed below.

Case (i): If $P \rightarrow \infty$, then equation (6) has solution $\sin \alpha a = 0$

$$\Rightarrow \alpha a = \pm n\pi \Rightarrow \alpha = \pm n\pi/a$$

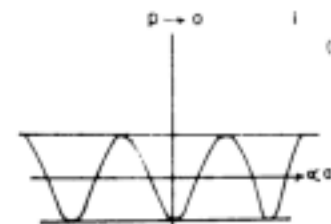
$$\alpha^2 = \frac{n^2 \pi^2}{a^2} = \frac{2mE}{\hbar^2}, \Rightarrow E = \frac{n^2 \hbar^2}{8ma^2} \quad \rightarrow (8)$$

The energy levels in this case are discrete and the result is similar to the energy levels of a particle in a constant potential box of atomic dimensions.

Case (ii): If $P \rightarrow 0$, then equation (6) has solution $\cos \alpha a = \cos ka \Rightarrow \alpha = k$

$$\Rightarrow \alpha^2 = k^2 = \frac{2mE}{\hbar^2} \Rightarrow E = \frac{p^2}{2m} = \frac{1}{2}mv^2 \quad \rightarrow (9)$$

The above expression is appropriate for completely free particles. Hence no energy levels exist, all energies are permitted to the electrons.



E-K Diagram - Brillouin Zones

The free electron moving in a periodic potential can have energies only in the allowed regions or zones. It is possible to plot the **total energy of free electrons** (V_s) their **wave vector K** as shown in the fig.1.5. We observe that the curve is not a continuous parabola, having discontinuities at $k = \pm(n\pi/a)$, $n=1,2,3$. From graph, it is evident that the electrons have allowed energy values in the region from $k = -(\pi/a)$ to $k = +(\pi/a)$. This is called the first Brillouin zone. After a break in energy values the allowed region extends from $k = -(\pi/a)$ to $(2\pi/a)$ and $k = (\pi/a)$ to $(2\pi/a)$. This zone is called the second Brillouin zone. Similarly, higher zones can also be defined. Each zone represents allowed energy values of an electron in a periodic potential field. For a completely free electron the curve is a continuous parabola.

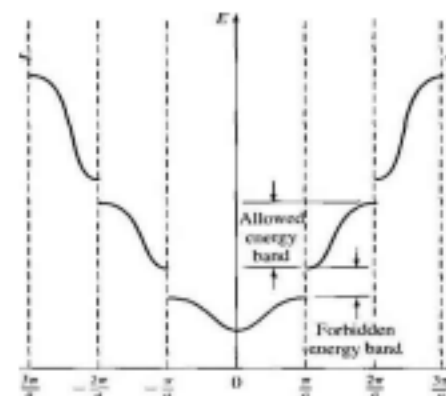


Fig.1.5

Origin of Energy Bands in Solids

In an isolated atom, the electrons are tightly bound and have discrete, sharp energy levels as shown in fig 1.6a. When two identical atoms are brought closer, the outermost orbits of these atoms overlap and interact. When the wave functions of the electrons on different atoms begin to overlap considerably, the energy levels corresponding to those wave functions split. If more atoms are brought together more levels are formed and for a solid of N atoms, each of the energy levels of an atom splits into N levels of energy.

The levels are so close together that they form an almost continuous band. The width of this band depends on the degree of overlap of electrons of adjacent atoms and is largest for the outermost atomic electrons. In a solid, many atoms are brought together so that the split energy levels form a set of bands of very closely spaced levels with forbidden energy gaps between them as shown in fig 1.6b.



Fig.1.6(a): Energy levels in an isolated atom



Fig. 1.6(b): Energy Bands in a solid

Classification of materials

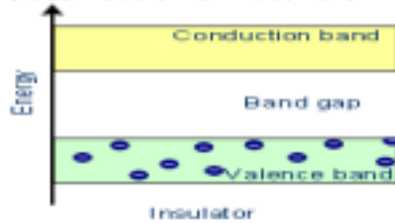


Fig 1.7(a)

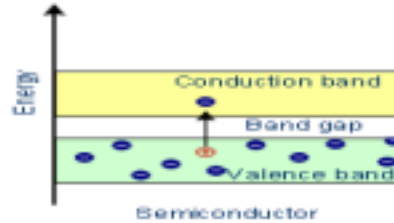


Fig 1.7(a)

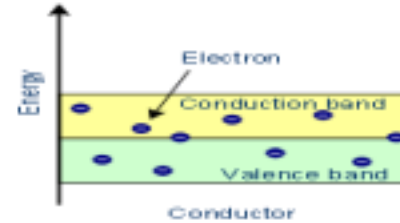


Fig 1.7(a)

According to Band theory of solids, solids are classified as follows.

(i) Insulators: In these materials, the valence band is completely filled, the conduction band is completely empty, the energy band gap is > 5 eV as shown in fig 1.7(a), and hence the valence band electrons are prevented from occupying vacant higher energy levels of the conduction band. Hence these materials have exceptionally low electrical conductivity. Eg :- Diamond, Sulphur etc.

ii) Semiconductors: They act like perfect insulators at 0K. In these materials the valence band is completely filled, the conduction band is completely vacant at 0K and the energy gap between them is about 1 eV as shown in fig1.7 (b). Above 0K, some of the valence band electrons acquire thermal energies greater than energy gap and jump into conduction band. Further the transition of electrons from VB to CB creates vacant levels in the valence band called holes. Eg :- Si, Ge, GaAs

iii) Conductors: These materials have very high values of electrical conductivity as the electrons can acquire higher energy easily. In these materials, the valence and conduction bands overlap with zero energy gap as shown in fig 1.7(c). Eg : - Cu, Ag, Au, Alkali metals

Effective Mass (m^*) of electron

The apparent mass of the electron moving in a periodic potential in the presence of an applied external electric field is called Effective mass of the electron. When an electric field ξ is applied electron attains a group velocity v_g .

Expression for m^* :

The group velocity of an electron = $v_g = \frac{d\omega}{dk} = \frac{1}{\hbar} \frac{dE}{dk}$ ($\because \omega = \frac{E}{\hbar}$ and $k = \frac{2\pi}{\lambda}$)

If ϵ is the electric field applied to the electron for a time say dt seconds, then work done on the electron by the field = Force $\times$ distance = Force $\times$ (velocity $\times$ time).

But this is equal to the change in KE of electron = dE (by work-energy theorem)

$$\text{i.e. } dE = e\epsilon \times (v_g dt) = \frac{eE}{\hbar} \cdot \frac{dE}{dk} \cdot dt \quad (\text{or}) \quad \frac{dk}{dt} = \frac{eE}{\hbar}$$

$$\text{and the acceleration of the electron } a = \frac{dv_g}{dt} = \frac{1}{\hbar} \frac{d}{dt} \left(\frac{dE}{dk} \right)$$

$$= \frac{1}{\hbar} \frac{d^2 E}{dk^2} \cdot \frac{dk}{dt} = \frac{1}{\hbar} \frac{d^2 E}{dk^2} \cdot \frac{e\epsilon}{\hbar} = \frac{e\epsilon}{\hbar^2} \left(\frac{d^2 E}{dk^2} \right) \quad (\because F = e\epsilon = m^* a)$$

$$m^* = \frac{\hbar^2}{d^2 E / dk^2}$$

Illustration of the graph

Variation of 'E' with 'k':

The variation of E with K is as shown in fig 1.8 (a), for the first allowed band. Points A and B are known as points of inflection which means that the slope of the reverses on both sides of these points.

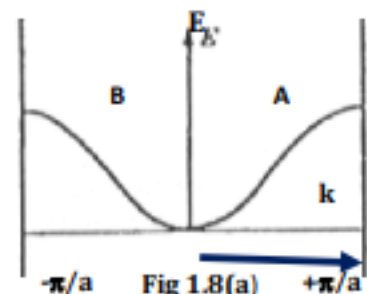


Fig 1.8(a)

Variation of ' V_g ' with ' k ':

The variation of the group velocity $V_g = \frac{1}{\hbar} \left(\frac{dE}{dk} \right)$ with k is illustrated in fig 1.8 (b). For $k=0$, $V_g=0$ as k increases reaching its maximum at A and B on both sides. Beyond the points of inflection velocity decreases and reaches zero. Note that, at the top and bottom of the energy band the velocity is zero and at the middle of the band V_g is maximum.

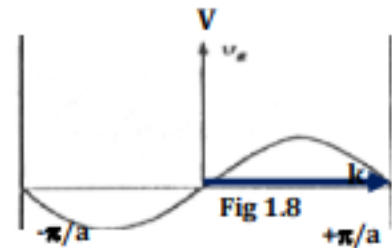


Fig 1.8

Variation of ' m^* ' with ' k ':

Fig 1.8 (c) shows the variation of m^* with k . At the bottom of the band i.e., at $k=0$, the effective mass m^* approaches the free electron mass m . As, k increases m^* also increases approaching infinity at A, as the slope $\frac{d}{dk} \left(\frac{dE}{dk} \right)$ is zero at that this point. The slope before A is positive. Beyond the point of inflection, A, m^* is negative and as K tends to (π/a) , m^* approaches a small negative value. The behaviour is similar on the other side of the graph also.

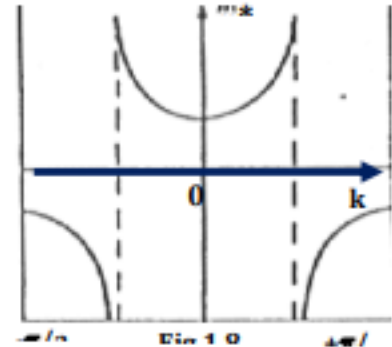


Fig 1.8

Significance of m^* :

If (d^2E/dk^2) is positive m^* is positive, if (d^2E/dk^2) is negative m^* is negative and if (d^2E/dk^2) is zero m^* is infinite. Here **positive m^*** refers to the direction of force of the electric field acting on the electron and its acceleration are in same direction. **Negative m^*** refers that the external force and acceleration are in opposite directions. As the lattice exerts a small retardation on the electron due to which effective mass becomes negative.

Unit II- Semiconductor Physics

Introduction

A Semiconductor is a solid which behaves as an insulator at absolute zero whose resistivity lies between 10^{-4} to $1 \Omega\text{-m}$. They have negative temperature coefficient of resistance, and the current conduction is carried out by two kinds of charge carriers namely electrons and holes in these materials. There are two types of semiconductors.

1. Intrinsic Semiconductor

2. Extrinsic Semiconductor

Intrinsic and extrinsic semiconductors

S. No	Intrinsic semiconductors	Extrinsic semiconductors
1.	A semiconductor in which charge carriers are created only by increasing temperature is known as an intrinsic semiconductor .	A semiconductor in which charge carriers are created not only by increasing temperature but also by adding impurities to it, known as an extrinsic semiconductor .
2.	These are in their pure form	These are doped with impurities
3.	They possess low electrical conductivity	They possess high electrical conductivity
4.	At 0 K, Fermi level exactly lies between conduction band and valence band	At 0 K, Fermi level lies closer to the bottom of the conduction band for a 'n-type' semiconductor and lies closer to the top of the valence band for a 'p-type' semiconductor
5.	Operating temperatures are low	Operating temperatures are high
6.	E.g.: Si, Ge, GaAs etc	E.g.: Si and Ge doped with Al, In, Ga, P, As, Sb etc.

Hall Effect

When a current carrying conductor or a semiconductor is placed in a transverse magnetic field, an electric field is produced inside the conductor/semiconductor in a direction normal to both the direction of current and the magnetic field. This phenomenon is known as **Hall Effect** and the generated voltage is called as **Hall Voltage**.

Consider a uniform metal strip placed with its length parallel to x-axis. Let a current I is passed through the conductor along x-direction and magnetic field (B) is established along Z-axis as shown in the figure 2.1. Due to this magnetic field the charge carriers experience a force perpendicular to X-Z plane i.e. along Y-axis.

By Fleming left hand rule this creates a transverse potential difference known as Hall voltage (V_H) which results in an electric field E_H called as Hall electric field. When equilibrium is reached, the magnetic deflecting forces are balanced by the force due to the Hall electric field E_H .

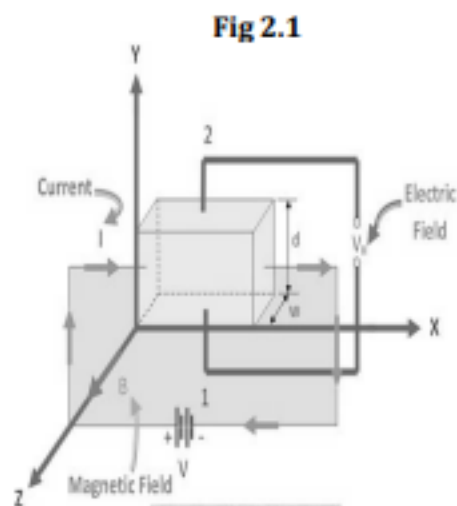


Fig 2.1

Magnetic deflecting force

$$= e(v_d \times B) = B \cdot e \cdot v_d \quad (v_d = \text{drift velocity of the charges})$$

Force due to Hall electric field

$$= e \cdot E_H$$

At equilibrium,

$$e \cdot E_H = B \cdot e \cdot v_d \Rightarrow E_H = B \cdot v_d \quad \rightarrow [1]$$

The drift velocity is related to the current density in x-direction as, $J_x = n \cdot e \cdot v_d$

$$\rightarrow [2]$$

From [1] and [2],

$$E_H = \frac{B J_x}{n e} \quad \rightarrow [3]$$

The ratio of hall electric field to the product of current density (J) and magnetic flux density (B) is known as Hall Coefficient (R_H) which numerically is the reciprocal of the charge density.

$$\text{i.e.} \quad R_H = \frac{E_H}{B \cdot J_x} = \frac{1}{ne} \quad \rightarrow [4]$$

If the thickness of the sample is t , then Hall voltage $V_H = E_H t$
 From [4], we have $V_H = R_H J_x B t$

The sign of V_H will be opposite for n and p type semiconductors.

If b is the breadth, then **area of cross section** of the sample through which current flows = $b \cdot t$

$$\therefore \text{Current density} = J_x = \frac{I_x}{bt} \Rightarrow V_H = R_H \frac{I_x}{bt} B t \Rightarrow V_H = R_H \frac{I_x}{b} B$$

$$\text{Thus} \quad R_H = \frac{V_H}{I_x} \cdot \frac{b}{B} \quad \rightarrow [5]$$

Equation [5] gives the expression for Hall coefficient.

Significance of Hall Coefficient

1. It is used to determine whether the given semiconductor material is p-type or n-type. i.e., if R_H is negative then the material is n-type and if R_H is positive then the material is p-type.
2. It is used to find carrier concentration of the given material from, $n = \frac{1}{e R_H}$
3. It is used to find the mobility of the charge carriers.
4. It is used to calculate the magnetic flux density B .

P-N Junction Diode

A P-N Junction diode is formed by combining a p-type crystal with a n-type crystal and subjecting them to high pressure so that it becomes a single piece. The assembly so obtained is called a P-N junction diode. The p-region has holes as majority charge carriers and the n- region has electrons as majority charge carriers.

Formation of P-N junction

Joining n-type material with p-type material causes electrons from the n-side to diffuse to the p-side and holes from the p-side to the n-side. Movement of electrons towards the p-side exposes positive ion cores in the n-side while movement of holes towards the n-side exposes negative ion cores in the p- side, producing an electric field at the junction. This results in a potential difference across the junction, known as junction potential or potential barrier. In the absence of an external potential the barrier controls the charge carrier movement through it.

Biasing

Connecting a p-n junction to an external dc voltage source is called biasing. There are two such connections in which a diode can be operated.

1. Forward bias
2. Reverse bias

Forward bias: When external voltage applied to the junction in such a direction that it cancels the potential barrier, allowing the current flow through it, is called forward bias. To apply forward bias, the +ve terminal of the battery is to be connected to p-side and the -ve terminal to the n-side. The applied forward potential acts against the potential barrier which results in a decrease of the width of the junction.

Reverse bias: When the external voltage applied to the junction in such a direction the potential barrier is increased, ceasing the current across it, is called reverse bias. To apply reverse bias, the -ve terminal of the battery is connected to p-side and the +ve terminal to n-side. The applied reverse voltage acts in the same direction as the potential barrier which results in the increase of the width of the junction.

I-V Characteristics of a P-N junction

The V-I characteristics of a P-N junction diode can be obtained with the help of the circuit shown in fig 2.2(a). The voltage across the diode and the current are measured with the help of a voltmeter and a milli or micro ammeter.

After connecting the diode in forward bias, by varying the supply voltage different sets of voltage and current values are obtained. By plotting them on a graph, the forward characteristics can be obtained as shown in the I quadrant of fig 2.2 (c). It can be noted from the graph that the current remains zero till the diode voltage attains the barrier potential. For a silicon diode, this is 0.7 V and for Germanium diode, it is 0.3 V. This is also known as knee voltage or cut-in voltage. Above this voltage the diode current increases exponentially with the forward voltage as the covalent bonds break releasing excess charge carriers at the junction. The diode current and applied voltage during forward bias can be expressed by $i = i_o [\exp(\frac{eV}{\eta kT}) - 1]$, where i_o is the reverse saturation current, T is the temperature and V is the forward voltage.

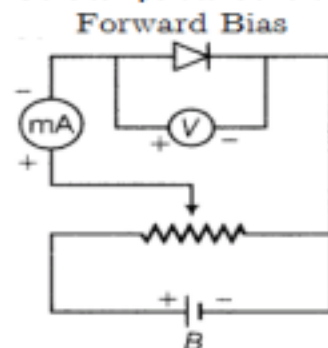


Fig 2.2(a)

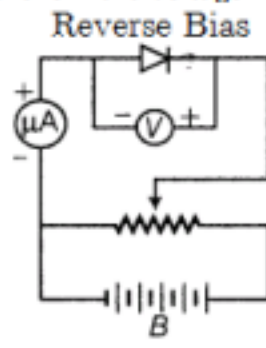


Fig 2.2(b)

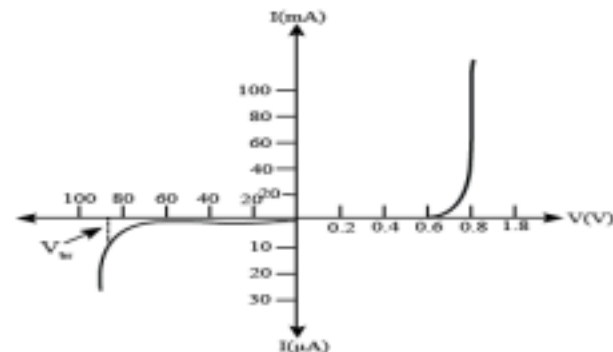


Fig 2.2(c)

The reverse characteristics can be obtained by reverse biasing the diode. During this, a very little current, almost independent of the applied reverse voltage flows across the junction due to minority carrier drift from both sides of the junction. This is known as leakage or reverse saturation current. It can be noted that at a particular reverse voltage, the reverse current increases rapidly. This voltage is called breakdown voltage as shown in figure 2.2(b).

Zener Diode

The diode which operates in the reverse breakdown region with a sharp breakdown voltage is called a Zener diode. It is an ordinary P-N junction diode except that it is properly doped to have very sharp breakdown. By adjusting the doping level, it is possible to produce Zener diodes with a breakdown voltage ranging from 2V to 800V. Zener diode undergoes two types of breakdowns depending on conditions.

Zener breakdown

In a heavily doped diode, the depletion region is narrow and when the reverse bias voltage is slightly increased the electric field across the depletion region becomes very strong of the order 10^7 V/m. A large number of electron-hole pairs are thereby produced. The reverse current rises steeply. This is called Zener breakdown.

Avalanche breakdown

The external applied voltage accelerates the minority carriers into the depletion region. These carriers gain sufficient energy to ionize atoms by collision. The electrons produced thereby accelerate to sufficiently large velocities to be able to ionize other atoms. This creates a sort of chain reaction. The cumulative effect of this chain reaction is the avalanche effect. This causes an abrupt rise in current called as avalanche breakdown. Zener breakdown is prominent at breakdown voltages less than 4V. The avalanche effect is prominent above 6V. Between 4V to 6V both effects are present.

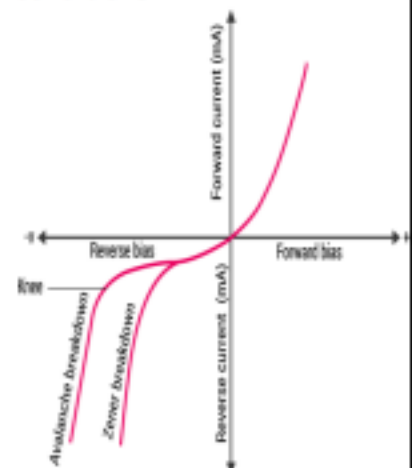


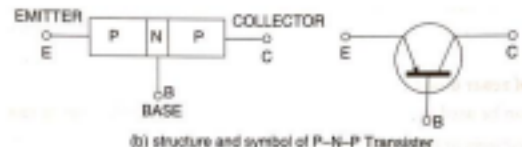
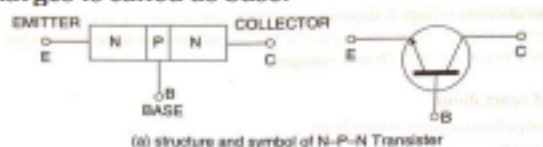
Fig 2.3

Bipolar Junction Transistor (BJT)

A bipolar junction transistor is a three terminal semiconductor device used to amplify electronic signals or switch electric power. It consists of three regions emitter, base and collector.

Emitter- This forms the left-hand section of the transistor. It is a heavily doped region which supplies charge carriers for conduction to the base.

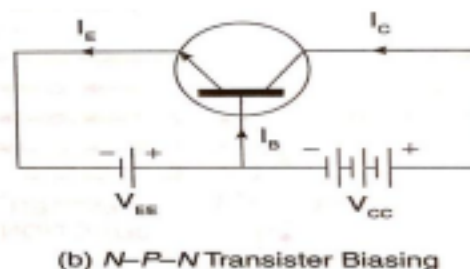
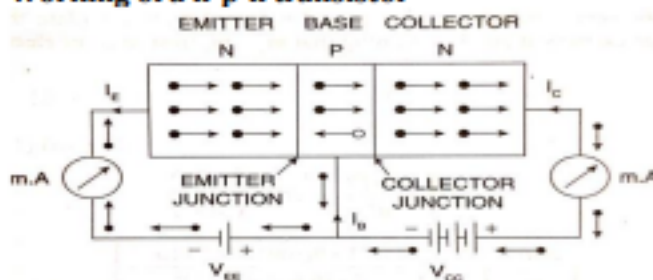
Base- A lightly doped thin region between the emitter and collector which controls the flow of charges is called as base.



Collector- A moderately doped region on the right-hand side of the transistor which collects charge carriers from the emitter through base.

The junction transistors are of 2 types. p-n-p transistor and n-p-n transistor as shown in the fig above.

Working of a n-p-n transistor



The emitter-base junction is forward biased by the battery V_{EE} . The collector-base junction is reverse biased by the battery V_{CC} . The directions of the emitter, base and collector currents are shown in the fig 2.4(a). The direction of each current is opposite to the direction of motion of electrons. The electrons being the majority carriers in the emitter are repelled due to the forward bias towards the base. The base contains holes as majority carriers and some holes and electrons combine in the base region. Since the base is lightly doped, the probability of electron-hole recombination in the base region is very small (5%). The remaining electrons cross into collector region and enter the positive terminal of the battery V_{CC} connected to the collector. At the same time an electron enters the emitter from the negative pole of the emitter-base battery V_{EE} . Thus, in a n-p-n transistors the current is carried inside the transistor as well as in the external circuit by the electrons. If I_E , I_B and I_C are the emitter, base and the collector currents respectively as shown in fig 2.4(b), then $I_E = I_B + I_C$.

Direct and indirect bandgap semiconductors

No	Direct Band gap semiconductors	Indirect Band gap semiconductors
1	The semiconductors in which the maximum of the valence band coincides with the minimum of the conduction band for the same value of the propagation constant K is known as a direct band gap semiconductor.	The semiconductors in which the maximum of the valence band does not coincide with the minimum of the conduction band for the same value of the propagation constant K is known as an indirect band gap semiconductor.
2	In these semiconductors the electron-hole recombination takes place directly from conduction band to valence band.	In these semiconductors the electron-hole recombination does not take place directly from conduction band to valence band but via a trap center in the forbidden band.
3	Lifetime of the charge carriers is less due to direct recombination	Lifetime of the charge carriers is more due to indirect recombination

4	Current amplification is less	Current amplification is more
5	Light is produced due to recombination	Heat is produced due to recombination
6	They are used for making LEDs, Laser diodes, ICs etc.	They are used in the manufacture of diodes, transistors, amplifiers etc.
7	All compound semiconductors such as GaAs, GaP, InSb are best examples for this type of semiconductors	The elemental semiconductors like Germanium and Silicon are best examples for this type of semiconductors

Light emitting diode (LED)

A forward biased P-N junction made of a direct bandgap semiconductor material. When a PN junction is forward biased, electrons diffuse from n-side to recombine with holes on the p-side emitting photons. Such junction diodes are called light emitting diodes.

Construction: A n-type layer is grown on a substrate over which a p-type layer is deposited on it by the process of diffusion. Metal contacts are made at the outer edge of the p-layer (anode) so that more upper surface is left free for the light to escape. For making cathode connections, a metal film is coated at the bottom of the substrate. This film also reflects as much light as possible to the surface of the device.

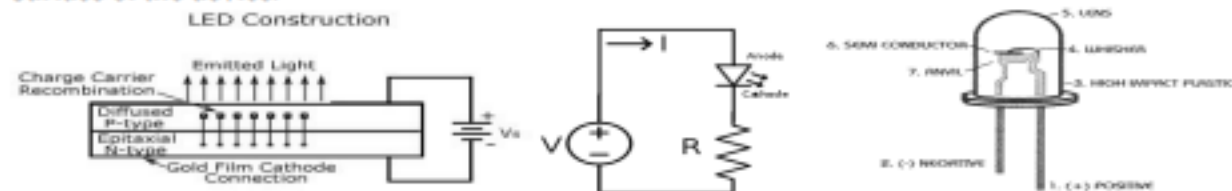


Fig 3.2 (a)

Fig 3.2 (b)

Fig 3.2 (c)

Working

When a junction diode is forward biased, electrons from n-side and holes from p-side move towards the depletion region and they recombine. During this process, energy is released because electrons make transition from conduction band to (higher energy level) to valence band (lower energy level). The energy is in the form of light if these recombinations take place in a direct band gap semiconductor. If E_g is the band gap, then the energy, $E_e = \frac{hc}{\lambda}$, and the corresponding wavelength is

given by $\lambda = \frac{hc}{E_g}$.

Materials: P-N junctions made with GaAs, InSb, GaN, InN, CdSe, ZnSe

Applications: LEDs are extensively used in numeric displays, traffic signals, emergency vehicle lighting, camera flashes, mining operations and data communication and signaling and also in flat panel displays.

Solar Cell

The solar cell is basically a p-n junction diode that converts sunlight directly to electricity with large conversion efficiency.

Construction: The cell is constructed by depositing a thin P layer on an N substrate. Contact is made to the P layer through narrow strips so that the illumination of the cell surface is not reduced by the presence of contacts. The P layer thickness should be less than one diffusion length so that the incident radiation should reach the junction region. The photo current I_L is due to the charge carriers (Electron-Hole) generated within one diffusion length of the charge carriers on either side of the depletion region. When the cell is exposed to sunlight a photon having energy less than the band gap E_g makes no contribution to the output, hence semiconducting materials with narrow band gap should be used for fabricating solar cells. Semiconductors with band gap between 1-2 eV can be used as cell materials.

Working: When a P-N junction is exposed to light, photons with energy greater than E_g are absorbed, and electron-hole pairs are generated on both P and N regions. These pairs are then separated by the

strong barrier field that exists across the depletion region. The electrons in P-region move towards the N-side and the holes in the N-region slide down to the P-side. When the P-N junction is open circuited, the accumulation of the charge carriers reaches an equilibrium value and the voltage across the junction is called as open circuit voltage (V_{oc}) and when the junction is short circuited maximum current flows in the circuit called as short circuited current (I_{sc}).

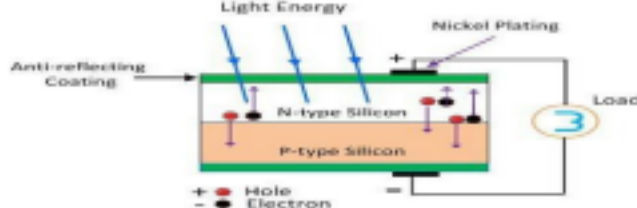


Fig.2.5(a)

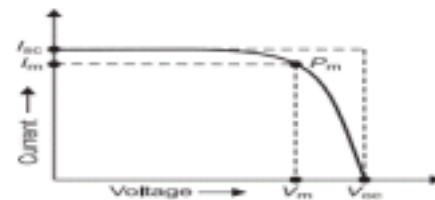


Fig.2.5(b)

The ratio of the area of the largest rectangle that can fit inside the curve area (maximum output power) to the product of the open circuit voltage and short-circuited current is called as Fill factor. The fill factor value for a solar cell will be between 0.65 and 0.8.

Materials: P-N junctions made with Silicon, cadmium Telluride (CdTe) and Copper Indium Gallium Selenide

Applications: Solar cells find extensive usage in calculators, watches, torches, garden lights, portable fans, alarm systems, in PV systems used in remote radar systems, artificial satellites, domestic electrical power generation and in telecommunications

Photo Diode

A photo diode is essentially a reverse biased P-N junction diode which is designed to respond to photon absorption.

PIN Photodiode

PIN Photodiode is a photodetector in which the depletion layer thickness can be modified for generation of large photocurrent. If the thickness of depletion layer on which light is falling is more than the surface area, the conversion efficiency of a photodiode increases, and more photo current will be generated.

Construction: The P⁺-type layer, intrinsic layer and N-type layer are sandwiched to form two junctions NI junction and PI junction. The P⁺ layer can be obtained by ion implantation and the intrinsic layer is an epitaxial layer grown on N-type substrate. The electrons from N-side and holes from P-side diffuse due to the concentration gradient. The width of the depletion region is more in intrinsic region and is narrow in the N-type layer. The P⁺-type side is connected to the negative terminal of battery and N-type side is connected to positive terminal of the battery.

Working: When the reverse bias is applied to PIN diode the width of depletion region starts increasing in the intrinsic region. And with the increase of reverse voltage, the width increases even more. A stage is reached when the depletion region's width becomes equivalent to the thickness of the intrinsic layer. At this point, the intrinsic layer is swept free of mobile charge carriers. It is assumed that there is no electron-hole recombination within the depletion region i.e., it is completely free of mobile charge carriers. And each photon absorbed generates one electron-hole pair.

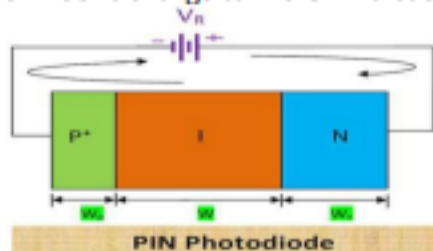


Fig 4.1

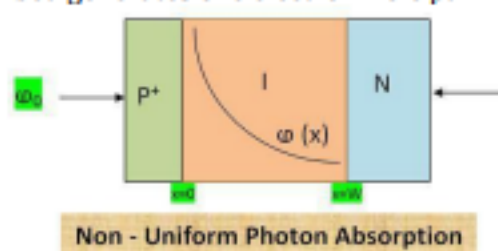


Fig 4.2

Materials: Diodes made with Silicon, Germanium, CdTe, GaP, InP and GaAs

Applications: PIN diodes are extensively used for detection of laser pulses, in ultrafast switching circuits and in logic circuits due to their ability to detect weak signals.

I B. Tech-II Semester (2023-24)
R22-Applied Physics
Important Questions for Mid I

Unit I

1. Derive Planck's radiation law. Discuss how it could successfully explain the entire spectrum of blackbody radiation.
2. What is photoelectric effect? Mention the factors affecting the photocurrent and kinetic energy of photoelectrons. Explain Einstein's photoelectric equation.
3. Describe Davisson-Germer's experiment with neat diagrams.
4. Derive time independent Schrodinger wave equation.
5. Show that the energies exhibited by the particle in a 1D potential box are discrete.
6. Discuss the salient features of Kronig-Penney model.
7. Derive an expression for the effective mass of an electron. Explain its significance.

Unit II

8. Explain Hall effect and derive an expression for the Hall coefficient of a semiconductor. Mention a few applications of Hall coefficient.
9. Explain the formation of P-N junction diode. Discuss the I-V characteristics of a P-N junction diode.
10. Describe the working of a bipolar junction transistor (BJT).
11. Discuss the construction, working and the I-V characteristics of a LED.
12. Illustrate with diagrams the construction, working and the I-V characteristics of a solar cell.

Unit III

13. Explain different types of polarization mechanisms in dielectric materials.
14. Discuss the nature of ferroelectric and piezoelectric substances.



Applied Physics

UNIT-III- Dielectric, Magnetic and Energy Materials Dielectric Materials

Electric dipole

The arrangement of two equal and opposite charges with a fixed distance of separation between them is called an electric dipole.

Dipole moment (μ)

The product of charge and distance between the two equal and opposite charges is called dipole moment. If q is the charge and d is the length of the dipole, then dipole moment $\mu = qd$.

The unit for dipole moment is coulomb-metre.

Dielectric constant (ϵ_r)

The dielectric constant is defined as the ratio between the permittivity of the medium to the permittivity of free space. i.e., $\epsilon_r = (\epsilon / \epsilon_0)$.

Electric Polarization

When an electric field is applied to a dielectric, the positive charges are displaced in the direction of the field while negative charges are displaced in the opposite direction producing local dipoles throughout the solid. This process of producing dipoles by the application of an electric field is called electric polarization.

Polarization Vector (P)

The average dipole moment per unit volume in a dielectric is known as polarization vector.

$P = N\bar{\mu}$, where N is the number of dipoles per unit volume and $\bar{\mu}$ is the average dipole moment.

Units: coulomb/metre.

Electric Polarizability (α)

The average dipole moment per unit applied electric field is termed as electric polarizability. This can be expressed as $\alpha = \frac{\bar{\mu}}{E}$. Units: Farad.metre²

Electric susceptibility (χ_e)

The electric susceptibility is defined as the ratio of polarization vector to the applied electric field vector. This can be expressed as, $\chi_e = \frac{P}{\epsilon_0 E}$.

Displacement Vector (D)

The number of lines of force received by unit surface area due to a charge 'q' of a material is known as displacement vector.

$$D = \frac{q}{4\pi r^2} \text{ and } E = \frac{q}{4\pi\epsilon_0 r^2} \text{ where } \epsilon \text{ is the permittivity of the medium.}$$

$$\therefore D = \epsilon E = \epsilon_0 \epsilon_r E$$

For a dielectric under polarization $D = \epsilon_0 E + P$ and $D = \epsilon_0 \epsilon_r E \Rightarrow \epsilon_0 E + P = \epsilon_0 \epsilon_r E$

$$\Rightarrow P = \epsilon_0 (\epsilon_r - 1)E \text{ and } \chi_e = \frac{P}{\epsilon_0 E} \Rightarrow \chi_e = \epsilon_r - 1 \Rightarrow \epsilon_r = 1 + \chi_e$$

Types of dielectrics

Polar dielectrics

The dielectrics which contain electric dipoles even in the absence of electric field are called as polar dielectrics.

Non -polar dielectrics

The dielectrics which contain induced dipoles in the presence of electric field are called non-polar dielectrics. The dipoles disappear when field is taken off.

Types of polarization

The displacement of electric charge centers results in the formation of electric dipoles in atoms, ions and molecules. There are 4 types of polarization.

1. Electronic polarization
2. Ionic polarization
3. Orientation polarization
4. Space charge polarization

Electronic polarization

When an electric field is applied to a dielectric material both electron cloud and nucleus in an atom experience Lorentz force and displace in opposite directions. The charge centers of both get separated resulting in the formation of an atomic dipole. Since the dipole formation is mainly due to the displacement of the electron cloud, this phenomenon is known as **electronic polarization**.

This type of polarization can be explained in atoms of rare gases. The polarizability for this polarization can be expressed as $\alpha_e = 4\pi\epsilon_0 R^3$ and is called as electronic polarizability, where R is the atomic radius.

Ionic Polarization

When an electric field is applied on an ionic dielectric, then the ions experience Lorentz forces in opposite directions such that the positive ions displace in the direction of the field and the negative ions in opposite direction to the field, forming ionic dipoles. This phenomenon is known as **ionic polarization**.

If the masses of the positive and negative ions are m & M and the angular frequency of their vibration is ω_0 , then the expression for ionic polarizability is given as $\alpha_i = \frac{e^2}{\omega_0^2} \left(\frac{1}{M} + \frac{1}{m} \right)$.

Orientation Polarization

When the electric field is applied on a polar dielectric then all the dipoles tend to rotate in the field direction, increasing the total dipole moment. This phenomenon is known as **orientation polarization**. This is also called as **dipolar polarization**.

This polarization occurs only in polar dielectrics. The expression for orientation polarizability is $\alpha_o = \frac{\mu^2}{3kT}$, where μ is the average dipole moment, T is the absolute temperature. This is only polarization that changes with temperature, where the polarizability varies inversely with temperature.

Space charge polarization

The space charge polarization occurs due to the accumulation of charges at the electrodes or at the interfaces in a multi-phase material. When field is applied anions accumulate at one place and cations at another place creating charge redistribution. This is called as space charge polarization. The net polarization of the material is the summation of all the 4 polarizations.

Ferroelectric Materials

1. Certain dielectrics **exhibit electric dipole moment even in the absence of an external electric field**, i.e., they exhibit spontaneous polarization, these materials are called as ferroelectric materials.
2. Valasek discovered this phenomenon in Rochelle Salt (Potassium sodium tartrate, $\text{NaKC}_4\text{H}_4\text{O}_6$) in 1920.
3. The Polarization vs Electric field (**P-E**) curves of these materials similar to the B-H curves of ferromagnetic materials, hence the nomenclature ferroelectrics.

4. Examples of such materials are BaTiO_3 , LiNbO_3 , KNbO_3 , KDP (KH_2PO_4) and ADP ($\text{NH}_4\text{H}_2\text{PO}_4$).

Properties

1. They exhibit spontaneous polarization, which vanishes above a particular temperature (T_c) called as Curie temperature.
2. In the ferroelectric region i.e., below T_c , the dielectric constant depends on the field strength and above T_c , it varies with temperature according to the relation $\epsilon_r = (C/T - T_c)$ where C and T are constants.
3. Ferroelectric materials also exhibit both piezo and pyroelectric natures.
4. Ferroelectric materials exhibit hysteresis.

Applications: Ferroelectric materials are used,

1. In the manufacture of small sized capacitors.
2. As memory elements by making use of their hysteresis property.
3. As piezoelectric transducers to detect and produce ultrasonics.

Structure of BaTiO_3

There exists an intimate relationship between ferroelectric properties and atomic arrangement, which can be illustrated with BaTiO_3 example. BaTiO_3 is a cubic crystal above 120°C . When cooled below its Curie temperature (120°C) it undergoes an asymmetric shift and the crystal changes from cubic to tetragonal shape, resulting in some net dipole moment.

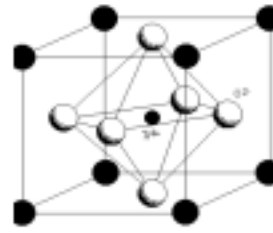


Fig 3.1(a)

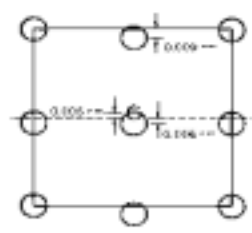


Fig 3.1(b)

Piezoelectric Materials

1. Development of electric charge (electric polarization) when subjected to stress is called piezoelectric effect. This phenomenon was discovered by Curie brothers in 1880.
2. Crystals like quartz, tourmaline (Boron Aluminum Silicate) and Rochelle salt when subjected to compression or tension develop opposite charges on their surfaces. For example, a 1 cm^3 cube of quartz with 2 kN of correctly applied force can produce a voltage of 12500 V.
3. These materials also exhibit the inverse piezoelectric effect i.e., these substances get strained when subjected to electric fields. The strains produced are quite small even for large fields.
4. All ferroelectrics are piezoelectric, but all piezoelectric materials are not ferroelectric in nature. BaTiO_3 exhibits both ferroelectricity and piezoelectricity while quartz (SiO_2) is only piezoelectric.

Applications: They are used

1. As transducers.
2. To produce and detect ultrasonic waves in a medium.
3. For the stabilization of RF oscillations for broadcasting purpose.
4. In (quartz) watches to maintain accurate time.

Pyroelectric materials

1. The dielectrics that produce polarization when subjected to uniform heating or cooling are known as Pyroelectric materials. This is due to the change in spontaneous polarization with temperature in certain anisotropic solids.
2. The pyroelectric effect can be described by the equation $\Delta P_s = P_g \Delta T$, where ΔP_s = change in polarization, ΔT = change in temperature and P_g = pyroelectric coefficient.
3. The pyroelectric coefficient is defined as the change in polarization per unit temperature change of the material. Thus, $P_g = \frac{dP}{dT}$. The higher the pyroelectric coefficient higher electric field that can be produced.

- The largest pyroelectric effects are observed in ferroelectrics. The most important pyroelectric materials are ceramics, synthetic polymers, ceramic polymer composites, lead zirconate-based ceramics.

Applications: Pyroelectric materials are used

- As thermal detectors and calorimeters
- In intruder alarms, fire alarms and radiometers.

Magnetic Materials

Hysteresis

The lagging of magnetization vector behind the magnetic field vector during one cycle of magnetization and demagnetization of a ferromagnetic specimen is known as "Hysteresis".

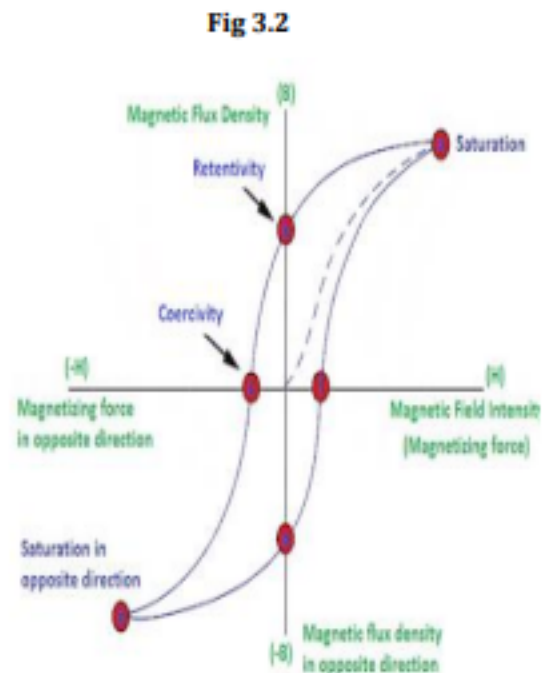
When a ferromagnetic virgin specimen is subjected to an external magnetic field the magnetization increases as shown in the figure.

For a higher magnetic field saturation is achieved. Further increase in the applied field from here, the magnetization remains constant.

If the field is decreased from here, reduces the magnetization, but some magnetization still resides in the specimen even when the field is zero known as "Retentivity".

By increasing the magnetic field in opposite direction, the remaining magnetization can be erased from the material. The demagnetizing field required to destroy the residual magnetization is known as "Coercivity".

With further increase in the field from this point the material undergoes a similar variation in opposite direction as shown in the fig 3.2.



Soft and Hard magnetic materials

S. No	Soft magnetic materials	Hard magnetic materials
1.	These materials are easily magnetized or demagnetized.	These materials are not easily magnetized or demagnetized.
2.	They possess smaller H-loop area and thus have low hysteresis loss	They possess larger H-loop area and thus have high hysteresis loss
3.	These materials have large values of μ_r & χ_m	These materials have small values of μ_r & χ_m
4.	Possess small values of Coercivity and Retentivity.	Possess large values Coercivity and Retentivity
5.	These are produced by annealing process	These are produced by quenching process.
6.	Domain wall movement is easy as the material is free from imperfections	Domain wall movement is difficult due to the presence of imperfections
7.	e.g.: Fe-Si & Fe-Ni alloys, ferrites and garnets	e.g.: Alnico, Cunife, Cunico, silmanal
8.	Used in switching circuits, microwave isolators, shift registers and matrix storage.	Used in magnetic detectors, flux meters, voltage regulators, damping devices

Magnetostriction and applications

1. Change in dimensions in a ferromagnetic material when subjected to external magnetic field is called magnetostriction. The material reverts to the original dimensions on removal of magnetic field.
2. This phenomenon was first discovered by JP Joule 1842. This process is also reversible i.e. magnetization can be produced when the dimensions of the material are changed due to an external force.
3. Magnetostriction is caused by the rotation of domains of a ferromagnetic material under the action of a magnetic field.
4. Magnetostriction can be positive or negative. Material showing positive magnetostriction expands when the magnetic field increases and contracts when it decreases. The converse is true for materials showing negative magnetostriction.
5. Iron shows positive magnetostriction in weak magnetic fields and negative magnetostriction in strong magnetic fields. Nickel and cobalt show negative magnetostriction.
6. There are 3 different ways in which magnetostriction can occur in a material.
 - i) Longitudinal: when the change in dimension occurs in the direction of the applied field
 - ii) Transverse: The change in dimension occurs perpendicular to the applied field
 - iii) Volume: The change dimension occurs both perpendicular & parallel to applied field.

Applications

Magnetostrictive materials are used,

- a) To make sensors that measure a magnetic field or detect a stress/force.
- b) In medical devices, industrial vibrators, ultrasonic cleaning devices, vibration or noise control systems.
- c) In the transfer ultrasonic energy in underwater sonar devices & in high frequency oscillators.

Arakshapur (V), Ghatkesar(M), Medchal Dist. - 501 301, T.S., INDIA

Magnetoresistance

1. The nature of the material to change its resistance with the application of magnetic field is called magnetoresistance. This effect occurs in semiconductors, non-magnetic metals and in magnetic metals. This effect was discovered in 1856 by Lord Kelvin.
2. He observed that resistance of an iron piece increased when the electric current is flowing in the same direction as the magnetic field and decreased when the electric current is flowing at right angles to the magnetic field.
3. The origin of magnetoresistance is the spin orbit coupling. As the magnetic field is applied the magnetic moment vector rotates and the electron cloud in each atom deforms slightly which changes the amount of scattering of conduction electrons.
4. Magnetoresistance depends on the strength and direction of the applied magnetic field.
5. Different types of magnetoresistance effects are given below.
 - a) Anisotropy magnetoresistance (AMR)
 - b) Giant magnetoresistance (GMR)
 - c) Extraordinary magnetoresistance (EMR)
 - d) Tunnel magnetoresistance (TMR)

Applications

Magneto resistors are used,

1. To measure magnetic field strength and direction.
2. For data storage in the hard disk drives.
3. In automotive sensors, biosensors non-volatile memories, contactless switches.
4. In electronic compass to measure Earth's magnetic field.

Magnetic bubble memories

1. Magnetic bubble memory is a type of computer memory that uses a thin film of magnetic material to hold small, magnetized areas known as bubbles. Here each bubble is a cylindrical magnetic domain capable of storing one bit of data.
2. As these memories are non-volatile the bubbles do not disappear when power is turned off.
3. Andrew Bobeck in 1967 is the first to create a sequential memory device to store data using the propagation of these magnetic bubbles.
4. Magnetic bubbles can only be formed by certain materials such as orthoferrites, hexagonal ferrites, synthetic garnets and amorphous metal films.
5. A magnetic bubble memory basically contains a chip, magnetic coils and permanent magnets. The chip is composed of a non-magnetic crystalline substrate upon which a thin crystalline magnetic epitaxial film is grown.
6. The bubbles are generated in accordance with input data by a magnetic field produced by a pulse of current. The bubbles move along a predetermined path in a soft magnetic material on the chip surface. To read or write, the bubbles are rotated past a read- write head.

Applications

Magnetic bubble memories are used,

1. for large sequential magnetic data storage up to 10 to 20 million bits in HDDs.
2. In digital signal analyzers and in electronic switching.

Magnetic field sensors

1. A magnetic sensor is a transducer that converts magnetic field into an electrical signal. They are used in detecting distance, speed, rotation, angle and position by converting magnetic information into electrical signals that are processed by electrical circuits.
2. These sensors use an internal magnet to detect a permanent or an electromagnetic field and produce either an analogue or digital output.
3. Magnetic sensors measure magnetic fields in terms of flux intensity and direction. The magnetic sensors are divided into various types based on technology or elements used.
4. A few types of magnetic sensors are coiled type of magnetic sensors, Hall element sensors, Magnetic induction sensors, Reed switches and Squids.

Applications

Magnetic sensors are used in,

1. Automobile and automotives, robotics and factory automation
2. Military, biomedical applications, geophysics applications
3. Space equipment, agriculture-power supply units, consumer electronics
4. Water & Energy appliances and green energy power plants

Multiferroic materials

1. Multiferroics are materials that simultaneously exhibit more than one type of ferro-ordering like ferro-magnetic, ferro-electric, ferro-elastic and ferro-toroidic etc. Hans Smith introduced the term multi-ferroics in 1994.
2. These materials exhibit long range order in at least one microscopic property and develop dipoles with a conjugate field.
3. Spintronics uses a lot of multiferroic materials. Spintronics is the study of intrinsic spin of an electron and its associated magnetic moment with the fundamental electronic charge in solid state devices.
4. Multiferroic materials have significant potential for producing multifunctional, miniature, high speed devices. But very few materials like BiFeO_3 , $\text{PbFe}_{12}\text{Nb}_{12}\text{O}_{30}$, YbMnO_3 exhibit multiferroic nature as the simultaneous existence two ferro-orderings has many limitations.
5. These are primarily of 2 types.

Type 1: In this the ferromagnetic and ferroelectric orderings occur independently. Magnetic and ferroelectric couplings are usually weak.

Type 2: In this, the electric and magnetic orderings occur simultaneously. In this, the coupling between magnetism and ferroelectricity is stronger than in type one.

6. In multiferroic memory devices information can be written electrically and read magnetically.

Applications

1. In spintronics multiferroic materials like BiFeO_3 are useful as spin valves, magnetic tunnel junctions, spin filters, AC/DC multiferroic magnetic field sensors, gyrators.

Energy Materials

Conductivity of liquid and solid electrolytes

An electrolyte is a substance that dissociates into ions in a solution and acquires the capacity to conduct electricity. Electrolytes serve two functions; **one to conduct electricity** and the other **to take part in oxidation reaction** that drives the electrons through the external circuit.

The conductivity of an electrolyte is a measure of its ability to conduct electricity. The conductivity can be determined by connecting the electrolyte in a Winston bridge circuit and measuring its resistance under the bridge balance condition. There are 2 types of electrolytes: liquid and solid electrolytes.

Liquid electrolytes

1. A liquid electrolyte has a better & higher conductivity. It also possesses high dielectric constant, low viscosity, and high ionic conductivity.
2. Liquid electrolytes have excellent contact area with high-capacity electrodes and can accommodate size changes of electrodes during charging and discharging cycles.
3. These electrolytes are made from a lithium salt, an organic solvent and additives.
4. Liquid electrolytes have poor physical and chemical stability.

Solid electrolytes

1. A solid-state electrolyte exhibits high conductivity for cations and anions but is an insulator electrically. The ions migrate through vacancies or interstices within the lattice which lead to ionic conductivity.
2. These electrolytes should possess very high ionic conductivity, negligible electronic conductivity, low activation energy.
3. Solid electrolytes are also known as fast ionic conductors or superionic conductors. They are widely used in batteries, supercapacitors and in sensors.
4. Solid electrolytes have good charging, increased cycle life, increased energy density and low leakage currents. The contact area being less with the high-capacity electrodes.
5. Solid electrolytes are classified into 3 types based on conducting ion type, host material and composition of materials.

Superionic conductors

Superionic conductors are solid conductors usually with a very high ionic conductivity. This high conductivity is due to the rapid movement of the ions through the lattice containing vacancies and interstices. The resistivity of these materials is of the order of 10^{-4} to $10^{-1} \Omega\cdot\text{m}$. Superionic conductors are classified as

- i) Cationic conductors like Ag^+ , Na^+ , Li^+ and Anionic conductors like F^- and O^{2-}
- ii) Amorphous conductors like glassy electrolytes and Polymer based conductors
- iii) Other materials like RbAg_4I_5 , AgI , CuI , Zirconia based materials, perovskite ceramics

Applications

Superionic conductors are used in batteries, solid oxide fuel cells & oxygen sensors

Supercapacitors

1. Supercapacitors are capacitors designed to store a large amount of energy typically 10-100 times more than electrolytic capacitors.
2. Supercapacitor combines the properties of a capacitor and a battery.
3. A supercapacitor consists of 2 electrodes, an electrolyte, and a separator.
4. The supercapacitor uses the same electrical charge storage mechanism as a capacitor, but charge accumulates at the interface of a conductor and electrolyte. Thus, a double layer is formed due to accumulated charge. The first layer is created by the charged electrode and the second layer is created by electrolytic ions.
5. Super capacitors are classified into 3 types: electrostatic double layer capacitors, pseudo capacitors and hybrid capacitors.

Materials used in supercapacitors:

For electrodes: Most commonly used materials are activated carbon, graphene, CNTs, metal oxides, metal hydroxides, conducting polymers, carbon aerogels

For electrolytes: The electrolyte materials used in supercapacitors can be classified into 3 types;

Aqueous electrolytes, Organic liquids, Ionic liquids

As separators: Cellulose, PET, polypropylene, polyethylene

Applications: Super capacitors are used in

- 1) wind power generation photovoltaic generation
- 2) rails, electric cars and in electric grid
- 3) Power supplies to portable devices like notebook computers, phones, cameras etc.

Rechargeable ion batteries

A battery is an electrochemical device that stores chemical energy and transforms it to electrical energy when connected in a circuit. It consists of 3 elements: the negative side, the positive side and an electrolyte. There are two types of batteries.

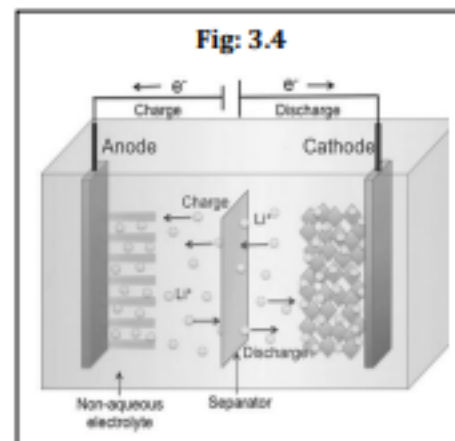
Primary batteries: The batteries made for one time use only and are non-rechargeable are known as primary batteries. **E.g.:** Alkaline battery and button cell battery.

Secondary batteries: A battery which is made for reusable purpose by recharging is called a secondary battery. These batteries have the same electrochemical reaction as alkaline batteries which can be reversed between charging and discharging. **E.g.:** Lead-acid batteries, Ni-Cd batteries, and Li-ion batteries.

Lithium-ion batteries: Li-ion battery is the best example of a rechargeable battery. Stanley Brittingham a British-American chemist is known as the founding father of lithium-ion batteries. These batteries generate DC power by utilizing chemical reactions. The battery mainly consists of a positive electrode made of a lithium metal oxide is coated on aluminium foil. A negative electrode made of graphite wrapped with a copper foil. Both foils are rolled in cylindrical shape with a separator between them soaked in an electrolyte. The separator is generally made of polypropylene and the electrolyte is made of a lithium salt mixed with an organic solvent.

During the working of the battery, Li ions move between cathode and anode internally. Electrons move in the opposite direction in the external circuit. During the process of charging, the charger passes current to the battery and lithium ions move from cathode to anode through the electrolyte. During the process of discharging, the lithium ions stored at the anode now move to cathode reversing the electrochemical process.

Applications: Lithium-ion batteries are used: **a)** In portable devices such as cameras, power banks, mobile phones, laptops and in medical devices **b)** In Military, aerospace and for marine applications **c)** In electric vehicles, UPS and grid storage **d)** For solar and wind power storage.



Applied Physics

Unit IV- Nanotechnology

Nanoscale

Nano means 10^{-9} . A nanometer is one thousand millionth of a meter (i.e., 10^{-9}m). Nanomaterials could be defined as those materials which have structured components with size less than 100nm at least in one dimension.

Nanoscience is the study of phenomena and manipulation of materials at atomic, molecular and macromolecular scales, where properties differ significantly from those at a larger scale.

Nanotechnology is the design, characterization, production and application of structures, devices and systems by controlling shape and size at the nanometer scale.

Classification of nanomaterials

Thin films or surface coatings: materials having 1 dimension in the nanoscale and are extended in the other 2 dimensions.

Nanowires and nanotubes: materials that are nanoscale in 2 dimensions and extended in 1 dimension.

precipitates and colloids: materials that are nanoscale in all 3 dimensions.

Properties of nanoparticles

The properties of nanoscale materials are very much different from those at a larger scale. Two principal factors that cause the properties to differ significantly are **increased relative surface area** and **quantum confinement** effects. These can affect or change properties such as reactivity, strength and electrical characteristics.

1. Increase in surface area to volume ratio

Nanomaterials have relatively larger surface area when compared to the volume of the bulk material. Consider a sphere of radius r , then its surface area $= 4\pi r^2$ and its volume $= 4\pi r^3 / 3$. The surface area to the volume ratio is $3/r$. i.e., as the radius of sphere decreases its surface area to volume ratio increases. Hence as the particle size decreases, greater proportions of atoms are found at the surface compared to those inside. Thus, nanoparticles have much greater surface area to volume ratio. As growth and catalytic chemical reaction occur at surfaces, then a given mass of the material in nano-particulate form will be much more reactive than the material of same mass in bulk form. This affects their strength and other properties.

2. Quantum confinement effect

When very large number of atoms are closely packed to form a bulk solid, the energy levels split and form bands. In case of isolated particles, atoms or molecules the energy levels are discrete. Nano materials represent an intermediate stage between isolated particles and the bulk form of the materials. When the dimensions of the material particles are of the order of their de-Broglie wavelength, the energy levels get affected. This effect is called Quantum confinement. This quantum confinement causes a subsequent change in the physical and chemical properties of nanomaterials like optical, electrical, magnetic, thermal, chemical behaviour etc. Thus, properties of nanomaterials differ from those of bulk materials.

Methods of preparation

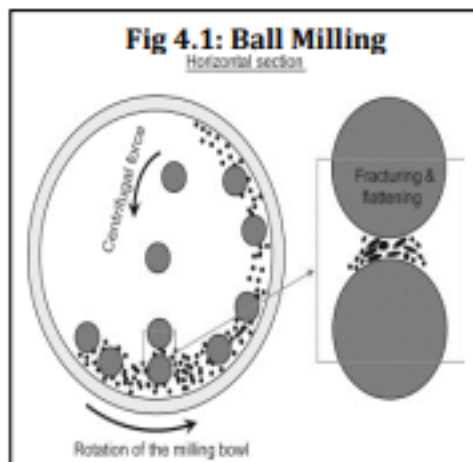
The fabrication of the nanomaterials can be done in 2 different ways.

- 1. Top-down approach:** In this process, the starting material is in bulk form. Nanostructures are obtained by careful and controlled division of the bulk material into nano size. Ex: Ball milling, PVD and CVD.
- 2. Bottom-up approach:** In this process, the atoms, molecules and even nanoparticles are used as building blocks for the creation of complex structures. Ex: Sol-gel method, Precipitation, and combustion methods.

Top-down Synthesis Methods

High energy ball Milling

Ball milling is a simple top-down method to produce nanoparticles. A ball mill contains a cylindrical stainless-steel container with many small steel, silicon or tungsten carbide balls that are made to rotate inside a drum. The material is taken inside the steel container in bulk powdered form crushed into nano sized particles using this technique. A magnet is placed outside the container to provide the pulling force to the material and this magnetic force increases the milling energy when the chamber rotates the metal balls. The significant advantage of this method is that it can be readily implemented commercially. Ball milling can be used to make carbon nanotubes and boron nitride nanotubes and is a preferred method in preparing metal-oxide nanocrystals like Cerium oxide (CeO_2) and Zinc Oxide (ZnO).



Physical Vapour Deposition (PVD)

Physical vapour deposition (PVD) is a thin-film coating process which produces coatings of pure metals, metallic alloys and ceramics with a thickness usually in the range of nanometers to a few microns. PVD as the name implies, involves physically depositing atoms, ions or molecules of a coating on to a substrate. PVD is used in a variety of applications, including fabrication of microelectronic devices, interconnects, battery and fuel cell electrodes, diffusion barriers, optical and conductive coatings, and surface modifications.

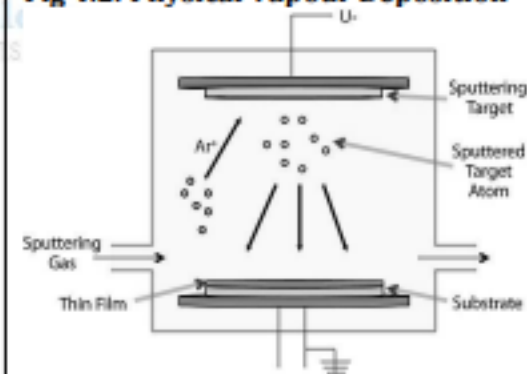
There are three main types of PVD, all of which are undertaken in a chamber as shown in fig 4.2 containing a controlled atmosphere at reduced pressure (0.1 to 1 N/m^2).

Thermal evaporation uses the heating of a material to form a vapour which condenses on a substrate to form the coating. Heating is achieved by various methods including hot filament, electrical resistance, electron or laser beam and electric arc.

Sputtering involves the electrical generation of a plasma between the coating species and the substrate.

Ion plating is essentially a combination of thermal evaporation and sputtering.

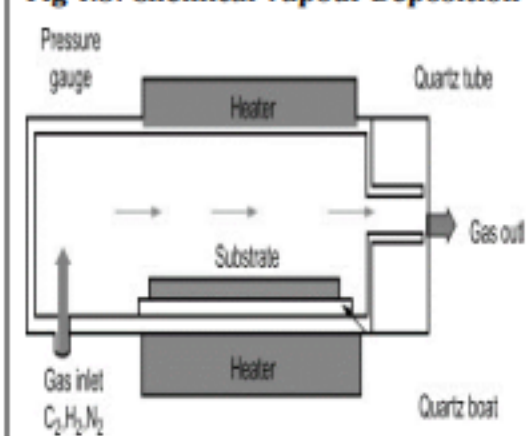
Fig 4.2: Physical Vapour Deposition



Chemical Vapour Deposition (CVD)

In this method, the material is heated to gaseous phase and allowed to condense on a solid surface in vacuum. Nanomaterials of metallic oxides or metallic carbides can be formed by heating metal and carbon or metal and oxygen in a vacuum chamber to gaseous phase and allowed to deposit on the surface of a solid as shown in fig 4.3. Metal nanoparticles are formed by exposing the metal to tuned metal-exciting microwaves so that the metal is melted, evaporated and formed into plasma at 1500°C . By cooling this plasma with water in a reaction column, nanoparticles are produced. The grain size of the nanoparticles depends on the concentration of metal vapour, its rate of flow in the reaction column and the temperature.

Fig 4.3: Chemical Vapour Deposition



Bottom-up Synthesis Methods

Sol-Gel method

A colloid that is suspended in a liquid is called a sol. The gelation of the sol in the liquid to form a network is called gel. The formation of sol-gel involves hydrolysis, condensation, growth of particles and the formation of networks as shown in fig 4.4. Using the sol-gel method, silica gels, zirconia and yttrium gels and aluminosilicate gels are formed. Nanostructured surfaces are formed using the sol-gel method.

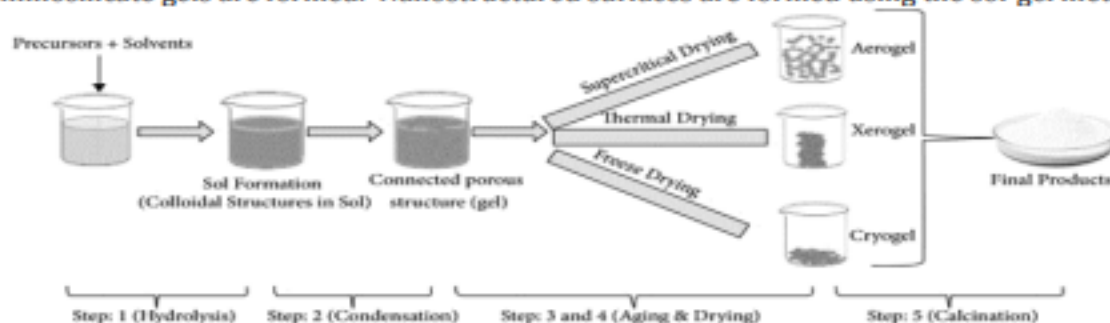


Fig 4.4: Sol-gel method

Steps involved:

Step 1: Preparation of stable solution: A stable solution of alkoxide or solvated metal precursor is formed.

Step 2: Formation of gel: An oxide (M-OH-M) or alcohol bridged gel is formed by polycondensation reaction that increases the viscosity of the solution.

Step 3: Aging of the gel: The polycondensation reactions will continue until the gel becomes a solid mass along with contraction of the gel network and expulsion of the solvent from gel pores.

Step 4: Drying of the gel: Water and other volatile liquids are removed from the gel network by drying. If the solvent such as water is extracted under supercritical condition the product is an aerogel or if isolated by thermal evaporation the resulting product is a xerogel.

Step 5: Calcination: By drying M-OH groups are removed and gel stabilizes. Heating the xerogel to temperatures around 800°C, the gel gets stabilised.

Precipitation Method

The precipitation method is a common technique to make nanoparticles by forming solid particles from a solution. The process involves mixing two solutions that contain metal ions and a precipitating agent, such as a base or a salt, respectively. The precipitating agent reacts with the metal ions and causes them to form insoluble compounds that precipitate out of the solution. The size and shape of the nanoparticles can be controlled by adjusting the concentration, temperature, pH, stirring speed, and other factors of the reaction. The precipitated nanoparticles can then be separated, washed, and dried for further use.

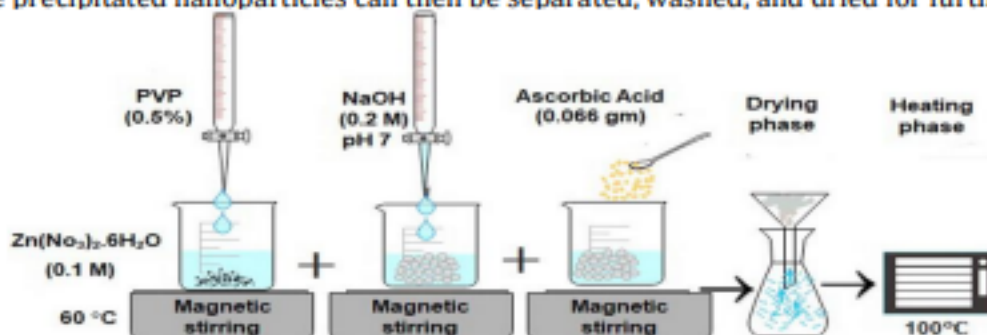


Fig 4.5: Preparation of ZnO nanoparticles using Precipitation Method

One of the examples of nanoparticle synthesis by precipitation method is Zinc oxide (ZnO) nanoparticles, which can be prepared by mixing zinc nitrate and sodium hydroxide solutions. The sodium hydroxide acts as a base and neutralizes the zinc nitrate, forming zinc hydroxide, which then decomposes into zinc oxide and water. The zinc oxide nanoparticles have applications in photocatalysis, sensors, and optoelectronics.

Combustion Method

The combustion method of preparing nanoparticles is a technique that involves the rapid heating of a solution containing metal precursors and fuel, resulting in the formation of nanosized metal oxides. This method is based on the utilisation of heat energy produced during the **exothermic spontaneous redox reaction** between an oxidiser and a fuel.

The oxidizer can be a metal nitrate and the fuel be any organic solvent like glycine, oxalic acid, urea, hexamine, sugar and EDTA etc. The advantages of this method are its simplicity, low cost, high yield, and scalability. Some of the factors that affect the properties of the nanoparticles are the type and amount of fuel, the metal-to-fuel ratio, the pH of the solution, the heating rate, and the cooling method. The nanoparticles prepared by this method find applications in catalysis, energy storage, biotechnology, and environmental remediation.

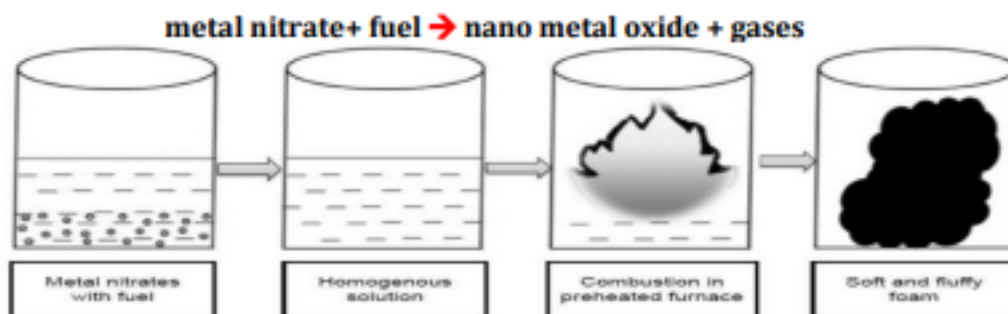


Fig 4.6: The 4 stages of the Combustion process

Procedure:

- The uniformly mixed solution of all the reactants is kept in a furnace maintained at 500°C.
- As the solution undergoes evaporation a concentrated and uniformly mixed viscous gel is obtained.
- After some time, the viscous gel catches fire and propagates spontaneously in the redox mixture in the form a flame or a smoulder.
- This combustion lasts for about 1-2 minutes after which a soft foam like substance is obtained.
- During the propagation, a large quantity of gases and high temperatures are produced which results in the formation of a nano-metal oxide.

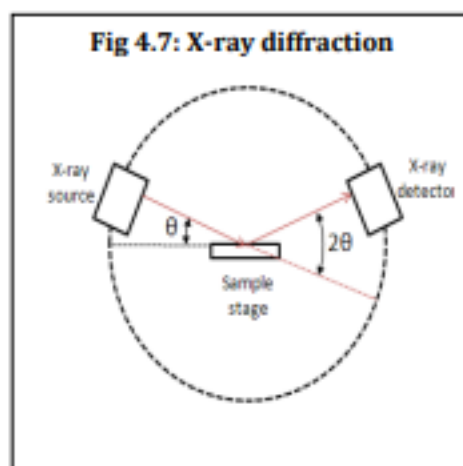
Characterization of Nanomaterials

Study and determination of various physical properties is called characterization process. The following are some important methods used in the characterization of nanomaterials.

- X-ray diffraction (XRD)
- Scanning Electron Microscopy (SEM)
- Transmission Electron Microscopy (TEM)

XRD (X-Ray Diffraction)

X-ray diffraction (XRD) is a widely used characterization process. In this, X-rays of wavelength comparable to the size of nanoparticles are allowed to get diffracted from the atomic planes of the specimen. A typical XRD instrumentation consists of 4 main components: X-ray source, specimen (sample), X-ray detector and a photographic film. When a beam of x-rays incident on the sample they are scattered by each atom of the sample and the scattered beams which are in phase interfere constructively to give maximum intensity as per Bragg's law, $n\lambda = 2d \sin\theta$. The intensities of these diffracted x-rays are recorded on a photographic film known as a **diffraction pattern**. These diffraction patterns can be analyzed to determine the size and orientation of the unit cells and also crystalline structure of the nanomaterials.

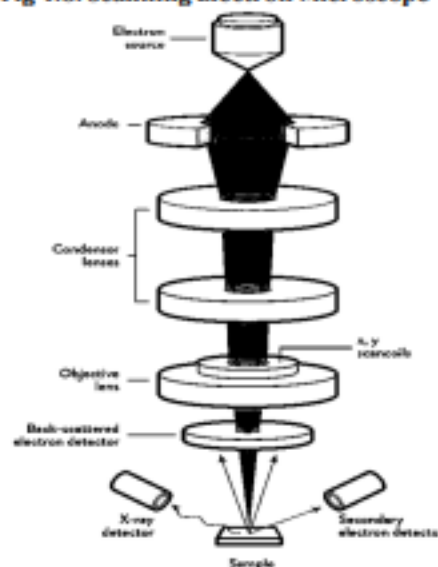


Scanning Electron Microscope (SEM)

In electron microscopes, beams of fast-moving electrons illuminate the material to be examined instead of ordinary light as in an optical microscope. As the de-Broglie wavelength of electrons is very small and comparable to dimensions of nanoparticles, images of very high magnification of the order of 10^6 can be obtained. The microscope consists of an electron source, magnetic lenses, sample stage, detectors and a display device.

In a typical SEM, the specimen is exposed to a narrow beam of electrons of energy between 0.2-40 keV from an electron gun. This beam is focused by the condenser lenses to a spot about 0.4 -5 nm in diameter. After falling on the surface of the specimen the primary electrons interact with the sample and an energy exchange takes place. This results in the reflection of high energy electrons, secondary electrons and EM radiation which give information of the size and grain structure of nanomaterials. SEM is useful because of its higher magnification, larger depth of the field and a higher resolution of the scanned image.

Fig 4.8: Scanning Electron Microscope



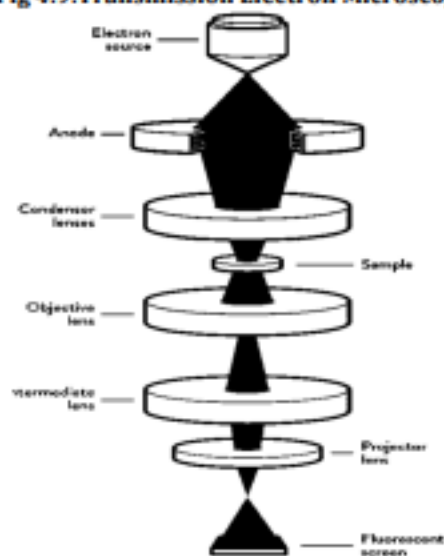
Transmission Electron Microscopy (TEM)

A Transmission Electron Microscope (TEM) operates on the same basic principles as the light microscope but uses electron beams instead of light. As the wavelength of electrons is much smaller than that of light the resolution is much better than from that of a light microscope. The resolution of TEM is about 0.2 nm which is about 100 times better than ordinary light microscope. The first TEM was developed by Max Knoll and Ernst Ruska in 1939.

A beam of high velocity electrons accelerated by the potentials around 120-200 keV under vacuum are focused by a condenser lens on to the specimen and the emergent electron beam is focused by an objective lens. The final image formed on a photo fluorescent screen or on a camera for image viewing.

TEM technique is used to provide topographical, morphological & compositional information of samples of molecular level. This information is useful to analyse the structure and texture of the metallic crystals. TEM can also be used to analyse the quality, shape, size and density of quantum wells, wires and dots.

Fig 4.9: Transmission Electron Microscope



Applications of Nanomaterials

- 1. Engineering:** i) Wear protection for tools and machines (anti blocking coatings, scratch resistant coatings on plastic parts) ii) In Lubricants – free bearings.
- 2. Electronic industry:** Data memory (MRAM, GMR-HD), flat panel displays (OLED, FED), Laser diodes, Glass Fibers, next-generation computer chips, high power magnets.
- 3. Automotive industry:** Light weight construction, Painting (fillers, base coat, clear coat), Sensors, Coating for wind screen and car bodies, self-cleaning glasses
- 4. Construction:** Construction materials, Thermal insulation, Flame retardants.
- 5. Chemical industry:** Fillers for painting systems, Coating systems based on nanocomposites, Magnetic Fluids.
- 6. Medicine:** Drug delivery systems, Agents in cancer therapy, Anti-microbial agents and coatings, Medical rapid tests, Active agents, nanomachines and nano devices
- 7. Energy:** Fuel cells, Solar cells, high energy density batteries, Capacitors.
- 8. Cosmetics:** Sun protection, Skin creams, Toothpastes, Lipsticks, shampoos etc.

Applied Physics

UNIT V: LASERS & FIBER OPTICS

LASER is an acronym for **Light Amplification by Stimulated Emission of Radiation**.

Interaction between Matter and Radiation

The interaction between matter and light is of 3 different types. They are

1. Stimulated Absorption
2. Spontaneous Emission
3. Stimulated Emission

Stimulated absorption: As shown in fig. 5.1(a), when an atom in the ground state absorbs a photon of energy $h\nu = E_2 - E_1$ and excites to a higher energy state, then it is known as stimulated absorption.

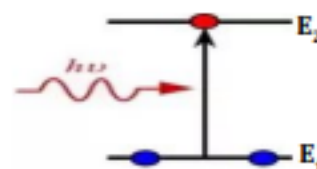


Fig: 5.1 (a)

Spontaneous emission: An atom initially present in the excited state makes transition to the ground state after completing its lifetime without any aid of external stimulus, it is known as spontaneous emission as in fig. 5.1 (b). In this process no phase relationship among the emitted photons.

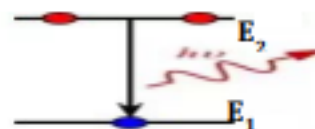


Fig 5.1 (b)

Stimulated emission: When an atom in the excited state interacts with an external photon during its lifetime it gets de-excited to the ground state by releasing another photon which is in phase with the incident photon as shown in fig. 5.1(c). This is called as stimulated emission. These two photons are coherent.

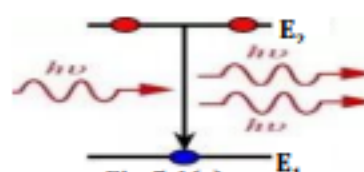


Fig:5.1(c)

Einstein's Relations

In a collection of atoms, all the three transitions, stimulated absorption, spontaneous and stimulated emissions occur simultaneously. Let N_1 be number of atoms/unit volume, with energy E_1 and N_2 be the number of atoms/unit volume, with energy E_2 .

Let 'n' be the number of photons/unit volume at frequency (ν), such that $h\nu = E_2 - E_1$, then the energy density of the interacting photons is given by $\rho(\nu) = n h \nu$. When these photons interact with atoms both upward and downward transitions are possible.

Upward transitions

The rate of stimulated absorptions depends on the number of atoms available in the lower energy state and the energy density of the interacting radiation.

$$\begin{aligned}
 \text{i.e., rate of stimulated absorptions} &\propto N_1 \\
 &\propto \rho(\nu) \\
 &= B_{12} N_1 \rho(\nu) \qquad \qquad \qquad \rightarrow [1]
 \end{aligned}$$

where B_{12} is called Einstein's coefficient for stimulated absorption.

Downward transitions

Once the atoms are excited by stimulated absorption, they stay for a short duration of time called 'lifetime'. After their lifetime atoms de-excite to lower energy level by emitting photons. This spontaneous emission rate depends on the atomic density of the excited state.

$$\begin{aligned}
 \text{i.e., rate of spontaneous emissions} &\propto N_2 \\
 &= A_{21} N_2 \qquad \qquad \qquad \rightarrow [2]
 \end{aligned}$$

where A_{21} is called Einstein's coefficient for spontaneous emission.

If the lifetime of the atoms is high in the order of a few milliseconds, they may interact with photons resulting in the stimulated emission of photons.

$$\begin{aligned}
 \text{Here, i.e., rate of stimulated emissions} &\propto N_2 \\
 &\propto \rho(\nu) \\
 &= B_{21} N_2 \rho(\nu) \quad \rightarrow [3]
 \end{aligned}$$

where B_{21} is called Einstein's coefficient for stimulated absorption.

For a system in equilibrium, Rate of upward transitions = Rate of downward transitions.

$$B_{12} N_1 \rho(\nu) = A_{21} N_2 + B_{21} N_2 \rho(\nu) \quad \rightarrow [4]$$

$$\Rightarrow \rho(\nu) = \frac{N_2 A_{21}}{N_1 B_{12} - N_2 B_{21}} = \frac{A_{21} / B_{21}}{\left[\frac{N_1}{N_2} \frac{B_{12}}{B_{21}} - 1 \right]} \quad \rightarrow [5]$$

By Maxwell-Boltzmann distribution, the population of i^{th} energy level is $N_i = g_i N_0 \exp(-E_i/kT)$, where N_0 is the population of the ground state.

Hence $N_1 = g_1 N_0 \exp(-E_1/kT)$ and $N_2 = g_2 N_0 \exp(-E_2/kT)$.

$$\Rightarrow \frac{N_1}{N_2} = \frac{g_1}{g_2} \exp\left(\frac{E_2 - E_1}{kT}\right) \Rightarrow \frac{N_1}{N_2} = \frac{g_1}{g_2} \exp\left(\frac{h\nu}{kT}\right) \quad \therefore \rho(\nu) = \frac{A_{21} / B_{21}}{\left[\frac{g_1}{g_2} \frac{B_{12}}{B_{21}} \exp\left(\frac{h\nu}{kT}\right) - 1 \right]} \quad \rightarrow [6]$$

$$\text{The radiation energy density from Planck's radiation law } \rho(\nu) = \frac{8\pi h\nu^3}{C^3} \frac{1}{\left[\exp\left(\frac{h\nu}{kT}\right) - 1 \right]} \quad \rightarrow [7]$$

$$\text{Comparing [6] with [7], we have, } \frac{A_{21}}{B_{21}} = \frac{8\pi h\nu^3}{C^3} \quad \rightarrow [8a]$$

$$\text{and } g_1 B_{12} = g_2 B_{21} \quad \rightarrow [8b]$$

Equations [8a] and [8b] are known as Einstein's relations.

Characteristics of Lasers

High Monochromaticity: A laser beam is more or less in single wavelength, i.e., the line width of laser beams is extremely narrow. The wavelength or frequency spread of light from conventional sources is usually 1 in 10^6 , whereas for lasers it is 1 in 10^{15} , i.e., if the frequency of radiation is 10^{15} Hz, then the width of line will be 1 Hz.

High Directionality: A laser beam is highly directional i.e. it can travel very long distances without divergence. The degree of divergence for laser is around 10^{-5} m per 1 meter travel whereas for a highly intensive conventional light source it is 0.5m per meter. Lesser the divergence, higher the directionality.

High Coherence: Laser beam is both spatially and temporally coherent. Spatial coherence is measured by coherence length which is around 600 KM for a laser and for any conventional light source it is only a few cm. Temporal coherence is measured by coherence time which is around 10^{-3} seconds for a laser whereas for a conventional source it is around 10^{-10} seconds.

High intensity: The intensity of light from a source can be illustrated by the number of photons coming out from it per unit area per second. For a laser this is around 10^{22} to 10^{34} photons /m²/sec whereas from a blackbody at 1000 K with wavelength $\lambda = 6000 \text{ \AA}$ is around 10^{16} only.

Lifetime: The time spent by an atom in the excited energy level is called the lifetime of that atom in that particular energy state. This is usually would be around 10^{-8} seconds.

Metastable State: An excited energy state which can accommodate atoms for a longer duration of time of the order of 10^{-3} seconds is known as a metastable state.

Population Inversion

The condition to make the excited energy state more populated as compared to the ground state is called population inversion. The number of atoms (N_1) present in the ground state (E_1), are usually larger than the number of atoms (N_2) present in the excited state (E_2). The process of making $N_2 > N_1$, is called population inversion.

For this to happen energy has to be supplied from an external source. For population inversion to occur the following conditions should prevail.

1. The system should possess at least a pair of energy levels (E_1 & E_2) such that $E_2 - E_1 = h\nu$.
2. There should be continuous supply of energy to pump the atoms to the excited states.

Components of Laser

Any laser system contains 3 major components, called as,

1. Pumping source
2. Active medium
3. Resonating cavity.

Pumping Source: It supplies suitable forms of energy to the active medium to achieve population inversion. There are five types of pumping mechanisms.

- | | | |
|-----------------------|----------------------|------------------------|
| 1. Electric Discharge | 2. Optical pumping | 3. Atom-atom collision |
| 4. Direct conversion | 5. Chemical reaction | |

In a laser, if the active medium is a transparent dielectric, then optical pumping is used. If the active medium is conductive, then electric discharge is employed.

Active Medium: It is the material in which a metastable state is present. With Metastable state only, population inversion takes place. Population inversion forces stimulated emissions; thus, a laser beam is emitted. The active medium can be a solid, liquid, gas or a forward biased PN junction.

Resonator cavity: It is an enclosure of active medium which essentially consists of two reflecting mirrors at the two ends as shown in fig. One of the mirrors is completely reflective and other is partially reflective. The resonator cavity provides optical feedback, due to the arrangement of mirrors which make the laser beam take number of reflections until it gets sufficient energy to come out.

Ruby Laser

Construction: Ruby laser is the first ever laser to be built by T.H. Maiman in 1960, which is a solid state, pulsed and is a 3-level laser.

As shown in Fig.4.26(a), the laser consists of a cylindrical shaped ruby crystal rod of length 2-20 cm and diameter 0.1-2 cm. The end faces of the rod are highly flat and parallel. One of the faces is highly silvered and the other face is partially silvered, so that it transmits 10% - 25% of incident light and reflects the rest to make the rod a resonant cavity.

Basically, ruby crystal is Al_2O_3 doped with 0.05-0.5% of Cr^{3+} ions which serve as activators. Due to the presence of 0.05-0.5% Cr^{3+} ions the ruby crystal appears in pink colour.

The ruby rod is placed along the axis of a helical xenon flash lamp of high intensity, which is surrounded by a reflector. The ends of flash lamp are connected to pulsed high voltage source so that the lamp gives flashes of intense light.

Working: When the flash lamp glows, each flash of light lasts for several milli seconds. The ruby rod absorbs the light to excite the Cr^{3+} ions to higher energy bands, $4F_1$ and $4F_2$. During the flash an enormous amount of heat is produced. The ruby rod is protected from the heat by enclosing it in a hollow tube through which cold water is circulated.

The Cr^{3+} ions are responsible for stimulated emissions whereas Al^{3+} and O_2 ions are passive. The emission of radiation by Cr^{3+} ions can be explained with the help of the energy level diagram as shown in fig.4.26(b). As the Cr^{3+} ions absorb photons of wavelengths 5500 Å and 4000 Å and get

excited to $4F_1$ and $4F_2$ levels from ground state. These are not metastable states and the de-excitation from these levels takes place in 10^{-9} S to 10^{-8} S to a bi-metastable state $2E$ which are non-radiative in nature. Population inversion takes place in the bi-metastable state due to which, stimulated emissions take place.

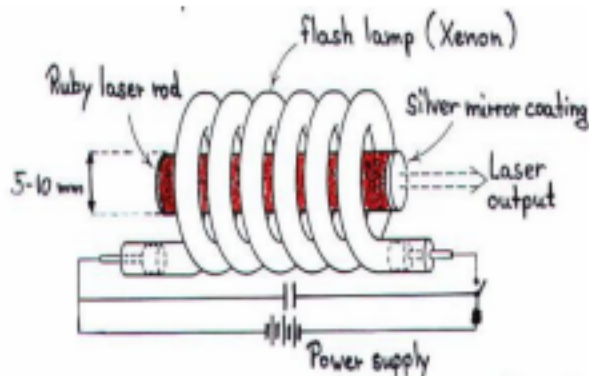


Fig. 4.26(a)

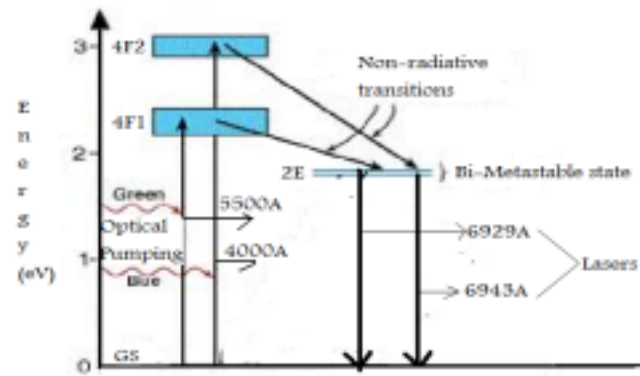


Fig. 4.26 (b)

The transition of Cr^{3+} ions from the bi-metastable state gives rise to emissions of wavelengths $6929\text{\AA}$ and $6943\text{\AA}$ respectively. The laser radiation is mostly due to $6943\text{\AA}$. Optical feedback is achieved by placing cylindrical elliptical reflector as photons travelling in all the directions are reflected to pass through the active medium. Ruby lasers are used for holography, industrial cutting and welding.

Nd: YAG Laser

Construction: A Nd: YAG laser is a solid state, 4 level laser invented by J.E. Geusic in 1965. This laser consists of yttrium aluminium garnet $Y_3Al_5O_{12}$ known as YAG which is an optically isotropic crystal. About, 0.725% of Y^{3+} ions in the crystal are being replaced by Nd^{3+} ions which serve as active centres. As shown in figure 4. 27 (a) the laser system consists of an electrically cylindrical reflector with the Nd: YAG rod is placed along one of its foci and the flash lamp at the other focus so that all the photons from the flash lamp will finally end up falling on the crystal rod. The crystal rod is of length 10 cm and 12 mm diameter. Here krypton flash lamp is used for optically pumping the Nd^{3+} ions to the excited states. Confocal mirrors are placed at the two ends of the rod to ensure proper optical feedback.

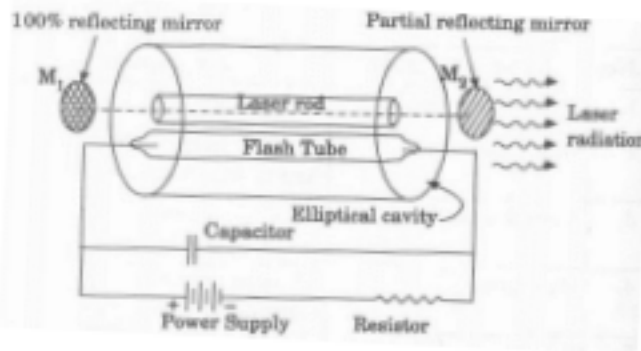


Fig 4.27 (a)

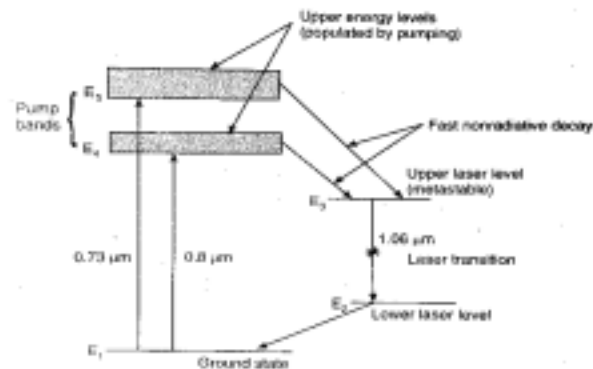


Fig 4.27 (b)

Working: when the flash lamp is on a ground state Nd^{3+} ion is excited to E_5 level if it absorbs a photon of wavelength $0.73 \mu\text{m}$ and to E_4 level if it absorbs a photon of wavelength $0.8 \mu\text{m}$ as shown in fig 4.27(b). Both E_4 and E_5 are normal excited states and the transitions from these states to the energy state E_3 are non-radiative in nature. As E_3 is a metastable state and the subsequent lower energy E_2 happens to be a normal excited state, population inversion is easily achieved and

maintained. The deexcitations from level E_3 to level E_2 are stimulated emissions resulting in an infrared laser of wavelength $1.06 \mu\text{m}$. The Nd^{3+} ions return to ground state E_1 from the lower laser level E_2 due to collisional deexcitation which are again pumped to E_4 & E_5 levels. The output can be either CW or pulsed depending on the requirement. These lasers found extensive applications in the fields of medicine, industry, and military.

He-Ne Laser

Construction: The first continuous wave gas laser was operated by Ali Javan and his co-workers in 1960. He-Ne laser consists of a long and narrow discharge tube of diameter 2-8 mm and length 10-100 cm. It is filled with helium and neon gases at pressures of 1 torr and 0.1 torr in 10:1 ratio. Neon atoms are the actual lasing atoms and helium atoms are used for selective pumping of neon atoms to the upper levels. The 2 mirrors are fixed at the two ends of the tube, such that one is partially silvered, and the other is fully silvered so that 100% reflection takes place. Two electrodes are fixed near the ends of the tube to pass electric discharge through the gas as shown in fig 4.28(a).

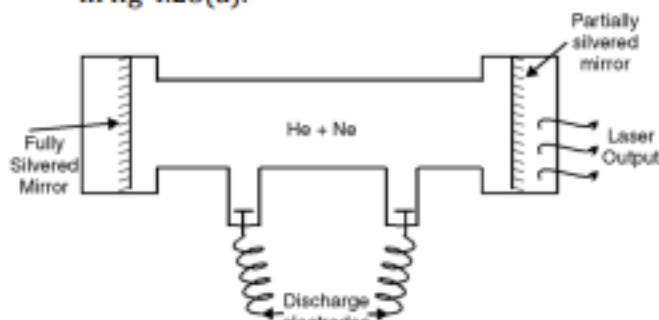


Fig. 4.28(a)

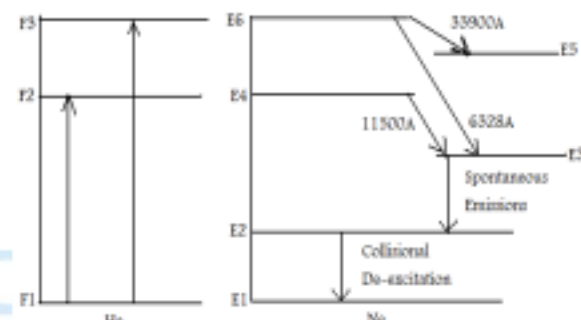


Fig. 4.28(b)

Working: When electric field is switched on, high voltage of about 10kV is applied across the gas mixture. This electric discharge accelerates electrons which excite helium atoms as they are lighter. The role of helium atoms is to excite neon atoms to their metastable states. Helium has two excited states F_1 and F_2 , whose energies match with the metastable states E_6 and E_4 of neon. When helium atoms get excited due to electric discharge, they transfer energy to neon atoms and return to the ground state. Such energy transfer is called resonant energy transfer. The probability of this energy transfer from helium atoms to neon atoms is more as there are 10 helium atoms for every single neon atom in the gas mixture. As in fig. 4.28(b), the excited Neon atoms transit to the ground state in 3 different ways leading to 3 Lasing transitions of different wavelengths. They are:

- Transition from **E_6 to E_5** results in a radiation of wavelength **$3.39 \mu\text{m}$**
- Transition from **E_4 to E_3** results in a radiation of wavelength **$1.15 \mu\text{m}$**
- Transition from **E_6 to E_3** results in a radiation of wavelength **$6328 \text{ \AA}$**

The first two transitions will lie in the infrared region and the third transition produces an orange red coloured laser. The atoms in E_5 and E_3 levels undergo spontaneous transitions to E_2 . The photons emitted due to these transitions are likely to excite atoms back to E_4 & E_6 levels. The atoms from E_2 level return to ground level, mainly due to collisions with the walls of the discharge tube.

Semiconductor Laser

Construction: A highly doped forward biased p-n junction diode made of a direct band gap semiconductor material is known as semiconductor laser. The laser is emitted from the junction due to recombination of conduction band electrons of N-side with valance band holes of P-side. Two conditions must be satisfied to achieve the laser emission.

- 1. Population Inversion**
- 2. Optical feedback**

Population Inversion is achieved by high doping concentration of the order (10^{19} atoms/ m^3). Optical feedback is obtained by polishing the ends of P-N junction at the right angles to the junction layer. As the forward bias is slowly increased through the junction stimulated emissions take place above a particular value of current density, known as Threshold Current Density.

There are 2 basic semiconductor lasers with their threshold current density values given below at room temperature.

1. Homo structure diode laser having a threshold current density around 5×10^4 A/ cm^2
2. Hetero Structure diode laser having a threshold current density around 2×10^3 A/ cm^2

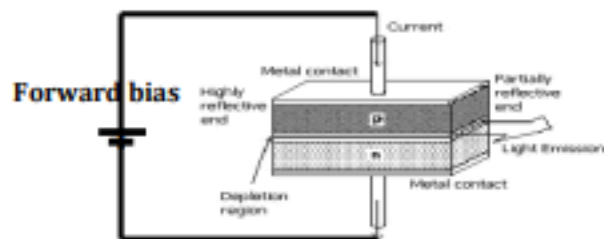


Fig 4.3(a)



Fig 4.3 (b)

Working: As shown in fig 4.3(a), when a forward bias with a dc source is applied to the semiconductor electrons from n-region and holes from p-region move to cross the and recombine in the junction region. The transition energy appears as photon energy.

Population inversion does not take place for low currents. Hence only spontaneous emissions take place and photons so released not coherent. When forward current exceeds the threshold value, population inversion take place and coherent photons are released as shown in graph of fig 4.3(b). The energy gap of **GaAs** is **1.47 eV** and the wavelength of the emitted radiation is $\lambda = \frac{hc}{E_g}$, where 'h' is Planck's constant, 'C' is velocity of light.

Applications of Lasers

In communication

- a) Lasers are used in optical fiber communications. Lasers are used as light sources to transmit audio, video signals and data over long distances without attenuation and distortion.
- b) Due to the narrow angular spread the laser beam can be used for the communication from earth to moon or to other satellites.
- c) As laser is not absorbed by water, it can be used in underwater communication networks.

Industrial Applications

- d) Lasers are used in metal cutting, welding, surface treatment and hole drilling. Holes can be drilled in steel, ceramics, diamond, and alloys. Using laser-controlled orifices and aerosol nozzles are drilled with controlled precession.
- e) Metal sheets and cloths can be cut in desired dimension by exposing them to laser beam.
- f) Welding has been carried by using a laser beam. Dissimilar metals can be welded with great ease.
- g) Laser beam is used in selective heat treatment for tampering with the desired parts in automobile industry.
- h) Lasers are widely used in the electronics industry in trimming the components of IC's.

Medical Applications

- i) Lasers are used in painless eye surgeries especially to treat the detached retina and cataract.
- j) Lasers are used in plastic surgery to treat skin injuries, in the removal of moles and tumors developed in the skin tissue.
- k) Used in Stomatology, the study of mouth and its disease. Laser is used for selective destroying, at the part of tooth affected caries. Mouth ulcers can be cured by exposing it to laser beam.

- l) Laser radiation is sent through optical fiber to open the blocked artery region, which burns the excess growth in the blocked region and regulates the blood flow without the requirement of a by-pass surgery.
- m) Lasers are used to destroy kidney stones and gall stones. The laser pulses are sent through optical fibers to the stoned region and laser breaks them into pieces.
- n) Used in cancer diagnosis and therapy.
- o) Used in bloodless surgery. During operation, the cut blood veins are fused at their tips by exposing them to infrared laser light, hence no blood loss.
- p) Lasers are used in liver and lung treatment and to control hemorrhage.
- q) Used in endoscopies to detect hidden tumors.
- r) Laser Doppler velocimetry is used to measure blood velocity in the blood vessels.

Military Applications

- s) Death rays: by focusing highly energetic laser beam for a few seconds to an aircraft missile, they can be destroyed.
- t) Laser gun: the vital part of enemy's body can be evaporated at a short range focusing on a highly convergent laser beam.
- u) LIDAR (Light Detection and Ranging): In place of RADAR, LIDAR is used to estimate the size and shape of distant objects and war weapons.

Computers

- v) By using lasers, a large amount of information or data can be stored in or retrieved from a CD-ROM. Lasers are also used in computer printers.

Scientific research

- w) Laser beams can initiate or fasten chemical reactions. Lasers are used to study the nature of chemical bonds, in counting atoms and also in isotope separation.
- x) Using laser irradiation, the monomers are united to form polymers.
- y) Lasers are also used in air pollution monitoring to find the size of the dust particles.
- z) Used in holography for recording and reconstruction of a hologram.
- aa) Used in thermonuclear fusion for creating very high temperatures of order 10^7 °C, sufficient to initiate nuclear fusion reaction.

Fiber Optics

Optical fiber is a cylindrical wave guide made of glass or plastic which transmits data in the form of electromagnetic radiation at optical frequencies by the principle of **total internal reflection** (TIR).

Principle of propagation of light through an optical fiber: Total internal Reflection

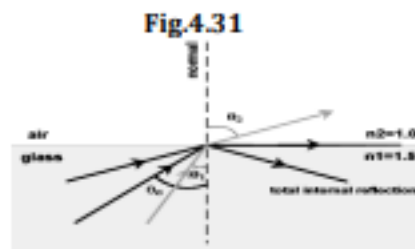
This occurs when a light beam travelling in an optically denser medium of refractive index (n_1) is incident on another medium of refractive index (n_2) with $n_1 > n_2$, at an angle greater than critical angle. The critical angle θ_c is defined as the angle of incidence for which the light ray will graze along the interface between the two media and is given as $\sin\theta_c = (n_2/n_1)$. When a light ray travels from a denser medium to a rarer medium, there exist three possibilities for the light ray to undergo depending on the angle of incidence as in fig 4.31.

Let θ_i be the angle of incidence, θ_r be the angle of refraction and θ_c be critical angle.

Case (1): when $\theta_i < \theta_c$, the angle incidence is less than the critical angle, the light ray refracts into the second medium.

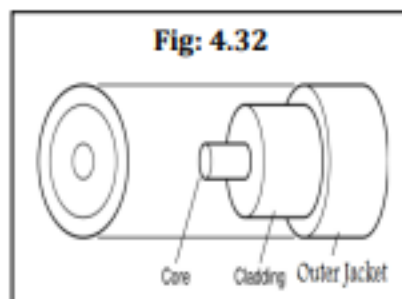
Case (2): when $\theta_i = \theta_c$, the angle incidence is equal to the critical angle, the light ray grazes along the interface between the two media.

Case (3): when $\theta_i > \theta_c$, the angle incidence is greater than the critical angle, the light ray reflects into the same medium.



The construction of a typical optical fiber essentially consists of core, cladding and the protective jacket as shown in fig 4.32. The core and cladding are made up of either glass or plastic whose values of refractive indices differ very slightly. The core is the wave guide that carries light through the fiber via total internal reflection. The main functionality of cladding is to confine light only to the core of the fiber. The protective jacket is made up of a polyurethane material which protects the fiber from mechanical stresses, abrasions, and outer changes.

A typical glass fiber consists of a central core of thickness 10-100 microns surrounded by a cladding of diameter 150-200 microns. By using the outer jacket, the fiber cable will not be damaged during hard pulling, bending, stretching, or rolling though the fiber made up of brittle glass.



Acceptance Angle and Numerical Aperture (NA)

The maximum entrance angle on the core for which the light ray travels through the core of the fiber by total internal reflection is known as **acceptance angle**.

Expression for Acceptance angle and Numerical aperture

Let us consider a cross sectional view of an optical fiber along its length as shown in figure 4.33. Let n_0 , n_1 and n_2 are the refractive indices of the outer medium (air), core and the cladding such that $n_0 < n_1 > n_2$.

Let the light ray OA strikes the core at the interface of air medium with an angle of incidence θ_0 and refracts into the core with an angle of refraction θ_1 . The refracted ray AB is again incident on the core-cladding interface at an angle $90^\circ - \theta_1$. If this $90^\circ - \theta_1$ is equal to the critical angle of the core and cladding media, then the ray grazes along the interface of core and cladding along path BC.

If a ray enters the core of the fiber at an angle greater than θ_0 , it will refract into the fiber at angle greater than θ_1 , then after falling on the core-cladding interface it would refract into cladding medium and gets lost. This is because the angle of incidence ($90^\circ - \theta_1$) at the core-clad interface is less than the critical angle. But if the angle of incidence at the surface of the core is $< \theta_0$, then $90^\circ - \theta_1$ will be greater than the critical angle. Therefore, total internal reflection takes place, and the light ray continues to propagate through fiber by multiple reflections. Hence θ_0 is the maximum entrance angle on the core surface for which the light ray is confined to travel through core of the fiber. This is known as **the Acceptance angle**.

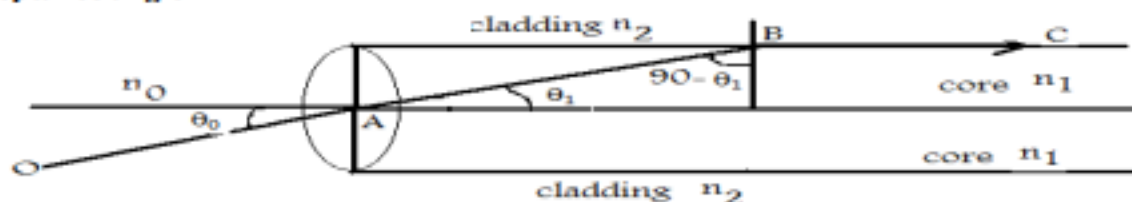


Fig. 4.33

According to Snell's law, at point A, $n_0 \sin \theta_0 = n_1 \sin \theta_1$

$$\Rightarrow \sin \theta_0 = \frac{n_1}{n_0} \sin \theta_1 \quad \rightarrow [1]$$

According to Snell's law, at point B,

$$\text{Is } n_1 \sin(90^\circ - \theta_1) = n_2 \sin 90^\circ \Rightarrow n_1 \cos \theta_1 = n_2 \Rightarrow \cos \theta_1 = \frac{n_2}{n_1} \rightarrow [2]$$

$$\sin \theta_1 = \sqrt{1 - \cos^2 \theta_1} = \sqrt{1 - \frac{n_2^2}{n_1^2}} \quad \rightarrow [3]$$

Substituting equation [3] in equation [1],

$$\sin \theta_0 = \frac{n_1}{n_0} \sin \theta_1 = \frac{n_1}{n_0} \sqrt{1 - \frac{n_2^2}{n_1^2}} = \frac{n_1}{n_0} \sqrt{\frac{n_1^2 - n_2^2}{n_1^2}} \Rightarrow \sin \theta_0 = \frac{\sqrt{n_1^2 - n_2^2}}{n_0} \rightarrow [4]$$

$$\therefore \text{Acceptance angle} = \theta_0 = \sin^{-1} \left(\frac{\sqrt{n_1^2 - n_2^2}}{n_0} \right) \rightarrow [5]$$

Acceptance Cone

If the acceptance angle is rotated around the axis of the fiber in all the directions, it forms a cone of semi-vertical angle θ_0 which is called as acceptance cone. Every light ray falling in the cone would be allowed to travel through the core of the fiber. Thus, acceptance angle is also called as acceptance cone-half angle.

Numerical Aperture (NA)

The light gathering capability of a given fiber is known as numerical aperture. It is numerically equal to the sign of acceptance angle.

$$NA = \sin \theta_0 = \sin \left[\sin^{-1} \left(\frac{\sqrt{n_1^2 - n_2^2}}{n_0} \right) \right] = \frac{\sqrt{n_1^2 - n_2^2}}{n_0}, \text{ For air medium } n_0=1 \Rightarrow NA = \sqrt{n_1^2 - n_2^2}$$

Mode of propagation

An independent path provided for the light ray to transmit through the optical fiber is known as Mode of propagation. Depending upon the number of paths available for propagation the fibers can be classified into 2 categories, a single mode fiber and a multi-mode fiber.

Types of Fibers

Based on refractive index profile, the optical fibers are classified into 2 types.

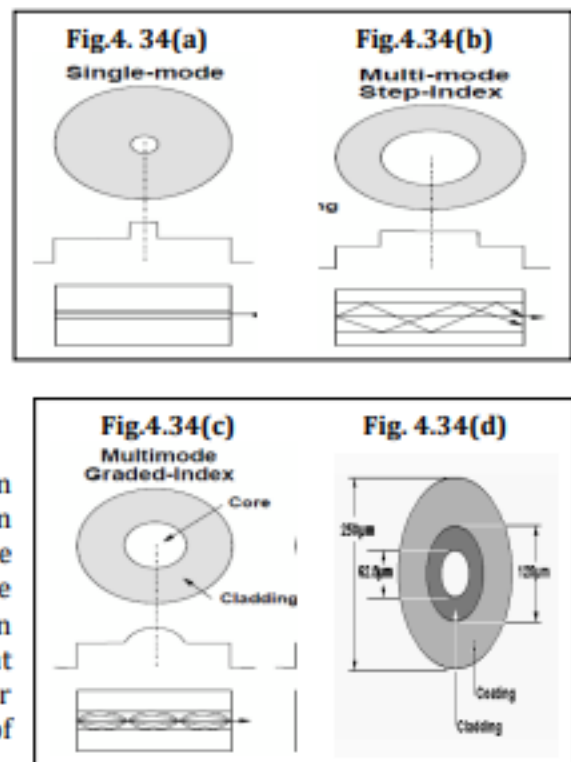
1. Step index fiber
2. Graded index fiber

Step index fibers

In step index fibers, the refractive index of the core medium is uniform throughout and undergoes an abrupt change at the interface of core and cladding as shown as fig. 4.34(a). The diameter of the core is about 50-200 micrometers in the case of multimode fiber and 10µm in case of single mode fiber. The transmitted optical signals travel through core medium in the form of **meridional rays**, which will cross the fiber axis during every reflection at the core-cladding boundary. The light rays appear to travel in a zig-zag manner as shown in fig. 4.34 (b).

Graded Index Fiber

In graded index fibers, the refractive index of core medium is made to vary in parabolic manner such that the maximum refractive index is present at the center of the core. The diameter of the core is about 50 micrometers. The transmitted optical signals travel through core medium in the form of **skew rays** which will not cross the fiber axis at any time. The shape of propagation appears in helical or spiral manner as shown in fig. 4.34(c). The typical values of the diameters of core and cladding are given in fig. 4.34(d).



Transmission of signal in optical fibers

Information can be sent from one place to another place using an optical fiber in the form of optical signals. In step-index fibers optical signals propagate in the form of meridional ray whereas in graded-index fibers in the form of skew or helical rays.

Attenuation of signal through optical fibers:

Signal attenuation is the decrease in the intensity of the signal during the propagation of the signal through the fiber. The three basic reasons for the loss of signal strength can be given as

- 1) Absorption losses
- 2) Bending losses
- 3) Scattering losses

Absorption losses

Absorption losses are mainly due to the presence of impurities. The fiber materials should be extremely pure and should have very low impurity concentrations. The OH⁻ ions of water trapped during manufacturing cause absorption at 0.95, 1.25 and 1.39 μm as shown in fig 4.36.

Bending losses

Bending losses occur whenever an optical fiber undergoes a bend of finite radius of curvature. Fibers can be subject to two types of bends, so, there are mainly two types of bending losses. Macro bending: This occurs when the radius of the bend is much larger than the radius of the fiber. Micro bending: This occurs when the radius of the bend and the radius of the fiber are of the same order.

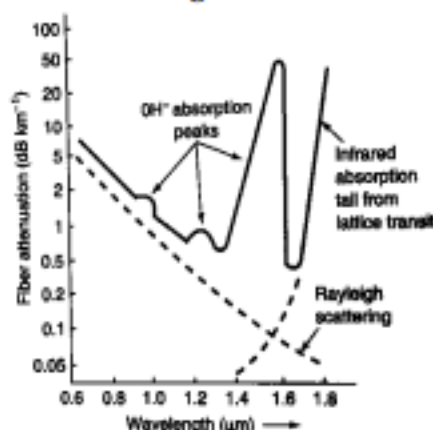
Scattering losses

Scattering losses are due to microscopic variations in the material density from compositional fluctuations and from structural defects occurring during manufacture. These are mainly due to Rayleigh scattering where the scattering losses vary inversely with the fourth power of wavelength as shown in fig 4.35 for a typical glass fiber.

The total power loss in an optical fiber due to different processes is generally measured in terms of decibels, which is a logarithmic unit. If P_{in} is the power launched into the fiber (input power) and P_{out} is the power after light propagates a distance of L kilometers through the fiber (output power), then the attenuation

coefficient can be expressed as $\alpha = -\left(\frac{10}{L}\right) \log\left(\frac{P_{out}}{P_{in}}\right)$ in dB/km.

Fig: 4.35



Advantages of optical fiber over the coaxial copper cable

- a) **Wider band width:** More information or data can be transmitted if the bandwidth of the waveguide is more. As the bandwidth is proportional to the signal frequency, which is in optical range (10^{14} - 10^{15} Hz) that is being sent through an optical fiber more data can be transmitted.
- b) **Low loss transmission:** Due to usage of ultra-low loss erbium doped fibers, the transmission losses will be very less. Due to this repeater spacing is about 100km making the transmission losses in fibers is as low as 0.2 dB/Km.
- c) **Dielectric wave guide:** As optical fiber cables are made with electrically insulating materials; they do not pick up any electromagnetic waves or high currents or any lightning. It is suitable even in explosive environments.
- d) **Freedom from electromagnetic interference:** When an electromagnetic signal is passed through a conducting medium the magnetic field around that line may interfere with the magnetic fields of any other signals which results in electromagnetic interference i.e. cross talk. As optical fibers propagate light through a closed medium there is no possibility of cross link.

- e) **Longer life:** The life span of optical fibers is expected to be 20-30 years compared to copper cables which have a life span of 12-15 years. Because they are light in weight, small, rugged, and flexible, optical fibers can be handled more easily than copper cables.
- f) **Low-cost transmission and high tolerance to temperature extremes** are some other additional features of optical fiber transmission.

Fiber optic communication system

An optical fiber communication system essentially consists of 3 parts: a transmitter, an optical fiber and a receiver. The transmitter includes a modulator, an encoder, a light source and a coupler. The light source can be a LED or laser diode. An audio signal in analog electrical form is being converted into a stream of binary electrical pulses by an encoder. The electrical pulses are made into optical pulses by a LED or a laser diode which are then to the fiber by means of a coupler.

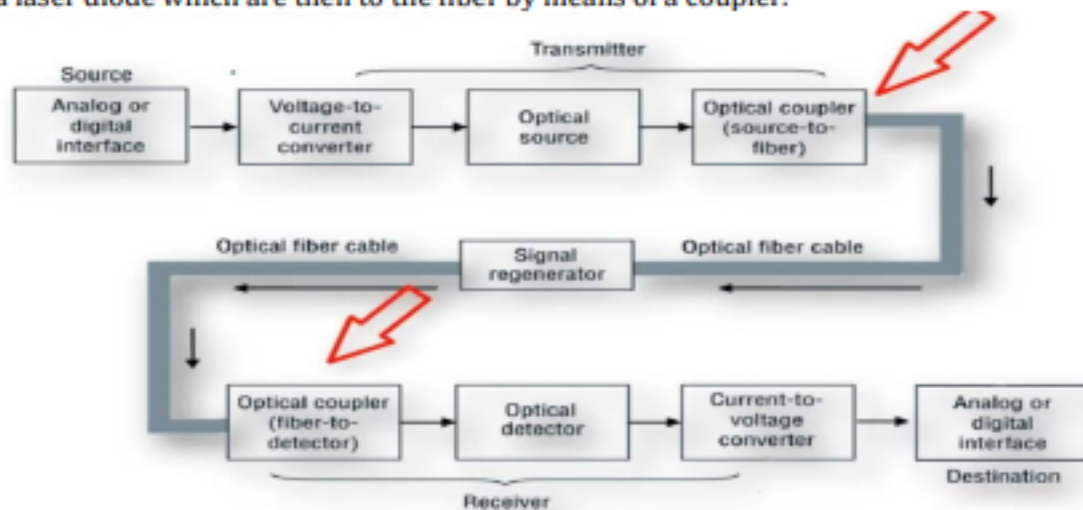


Fig 4.36: Block diagram of an optical fiber communication system

These optical pulses are carried over long distances by the optical fiber cable. While travelling through the fiber the signal may get attenuated due to various reasons like absorption, scattering and bending. By using a device known as a **repeater** the quality of the signal can be restored. A repeater improves the signal strength by amplifying the signal. It is understood that larger the repeater spacing better the quality of the fiber. With the help of a coupler the optical signals will be received by the receiver from the wave guide.

A receiver consists of 3 parts: a photo detector, a decoder, and an amplifier. The photo detector converts the optical pulses into equivalent electrical pulses. These electrical pulses in sequential form are then fed to a decoder which converts them into analog electrical signals. The amplifier amplifies the electrical signals and restores the quality of the signal back to its original.

Applications of optical fibers

- Optical fibers are used in communication to transmit data, telephone signals for secure communication without much loss due to attenuation.
- Used in computer networks to connect servers to the internal computers in LANs and WANs. Intra & internet connectors are made with optical fibers.
- The OFC cables are used in radars, digital cameras, and transmitting TV signals in digital form.
- Used electronics to produce delay in microwave signal processing.
- Used as endoscopes in medicine, used to diagnose & treat the interior of parts of human body.
- Used as sensors to monitor changes in physical parameters like pressure, temperature, stress, displacement, humidity etc. as they are highly sensitive to these changes which can be measured with high accuracy.