

بِسْمِ اللَّهِ الرَّحْمَنِ الرَّحِيمِ

(In the Name of Allah, the Most Merciful, the Most Compassionate.)

PHYSICS

12



**PUNJAB EDUCATION, CURRICULUM,
TRAINING AND ASSESSMENT AUTHORITY**

This textbook is based on Updated / Revised
National Curriculum of Pakistan-2023 and has been approved by the
Punjab Education, Curriculum, Training and Assessment Authority (PECTAA).

All rights are reserved with PECTAA.

No part of this textbook can be copied, translated, reproduced or used for
preparation of test papers, guidebooks, keynotes and helping books.

Authors

- | | |
|--|---|
|  Prof. Muhammad Ali Shahid (Izaz-i-Fazeelat)
Director Technical (Rtd.), Defunct (PTB), Lahore |  Dr. Riaz Ahmad
Prof. Emeritus, GCU, Lahore |
|  Prof. (Rtd.) Muhammad Nisar
Govt. College, Model Town, Lahore |  Prof. Muhammad Attiq-us-Salam
Govt. Graduate College of Science, Lahore |
|  Prof. Sheikh Nasir Iqbal
HOD (Physics), Govt. Graduate College, Gojra Road, Jhang |  Prof. Muhammad Arshad
HOD (Physics), RND Unique Group of Institutions, Lahore |
|  Dr. Naveed Afzal
Associate Prof., GCU, Lahore |  Dr. Mohsin Rafique
Assistant Prof., GCU, Lahore |

Editor

-  **Dr. Waris Ali**
Associate Prof. Govt. Graduate College of Science, Wahdat Road, Lahore

Reviewers

- | | |
|--|---|
|  Prof. Dr. Muhammad Jamil Alvi
COMSATS University, Lahore |  Dr. Muhammad Javaid Afzal
HOD (Physics), Govt. Islamia Graduate College, Lahore |
|  Tanvir Ahmad
Associate Prof., Forman Christian College, Lahore |  Abdul Khaliq Bhatti
M.Phil (Physics), SST, Govt. Junior Model High School
Model Town, Lahore |

 **Subject Coordinator**

 **Deputy Director (Compliance-Sciences)**

 **Director (Curriculum & Compliance)**

 **Incharge Art Cell**

 **Composing**

 **Designing & Layout**

 **Illustrations**

Abdul Rauf Zahid

Syed Saghir-ul-Hassnain

Aamir Riaz

Aisha Sadiq

Irfan Shahid

Hafiz Inam-ul-Haq

Hafiz Inam-ul-Haq , Aiyatullah

**Experimental
Edition**

CONTENTS

Chapter No.	Description	Page No.
13	Thermal Physics	1
14	Simple Harmonic Motion	15
15	Physical Optics	41
16	Electrostatics	58
17	Alternating Current	82
18	Quantum Physics	109
19	Nuclear and Particle Physics	128
20	Medical Physics	146
21	Space and Environment	156
	Pairing Scheme and Model Paper	173

Thermal Physics

Students Learning Outcomes

After studying this chapter, the students will be able to:

- ◆ explain how molecular movement causes the pressure exerted by a gas.
- ◆ derive and use the relationship $PV = \frac{1}{3}Nm \langle v^2 \rangle$.
[where $\langle v^2 \rangle$ is the mean-square speed (a simple model considering one-dimensional collisions and then extending to three dimensions using $\frac{1}{3} \langle v^2 \rangle = \langle v_x^2 \rangle$ is sufficient].
- ◆ calculate the root-mean-square speed of an ideal gas.
- ◆ derive and use the formula for the average translational kinetic energy of a gas.
- ◆ illustrate that the model of ideal gases is used as a base from which the field of statistical mechanics emerged [and has helped explain the behaviour of 'non-ideal' gases through modifications to the model e.g. the behavior of stars].

Thermal physics is an important area of study which deals with the relationship between heat, work, temperature, and other forms of energy. The laws of thermodynamics describe the mechanism of energy change within a system, allowing it to perform useful work on its surroundings. A large system containing many atoms or molecules is called a macroscopic system, and a system consisting of a single atom or molecule is called a microscopic system. Macroscopic systems have properties such as temperature and pressure, these are thermal properties of the whole system. They can be observed and studied without reference to the molecular nature of matter. Microscopic systems have properties such as kinetic energy and momentum. In thermal physics, we deal with a large number of particles of the order of Avogadro's number. These particles may be atoms or molecules in gases, liquids and solids. We can extend it to the electron motion in metals and neutrons in neutron stars.

There are a number of applications of the kinetic theory of gases in daily life such as in automobile engines, turbines, pumps, air conditioners, preparation of food, and environment etc. These diverse applications make thermal physics an important area.

Thermal physics is an area which includes the knowledge of statistical mechanics and kinetic theory of gases.

13.1 BROWNIAN MOTION

Brownian motion is the random and irregular motion of molecules in a gas. In 1827, Robert Brown, a botanist observed under a microscope that tiny plant pollen grains suspended in water moved randomly. These particles were identified as dust particles.

Later, it was proved to be one of the effects of molecular motion.

A molecule in a gas changes its path after collision with another molecule. When it keeps on colliding with other molecule, the interacting molecule follows a random or zig-zag motion. In fact, collision transfers or exchanges the momentum and energy between the molecules (Fig 13.1).

Brownian motion describes randomness and chaos, therefore, it represents one of the simple models of randomness. There are various reasons and causes of this motion which are given as under:

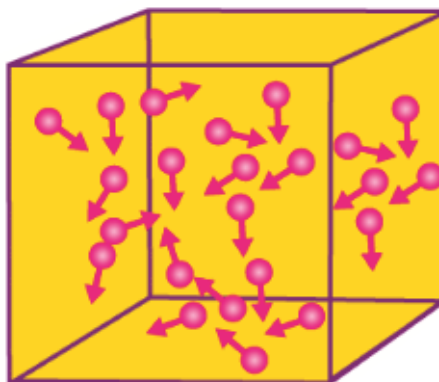


Fig. 13.1: Particles are in random motion

1. For a given impulse, lighter (smaller-mass) particles gain a large change in velocity, so they show more vigorous Brownian motion. Avoid trying it to particle size.
2. The speed of the particles is inversely proportional to the viscosity of the fluid. Low viscosity of the fluid results in faster Brownian movement.
3. Viscosity describes the magnitude of the internal friction in a fluid. It represents the resistance to flow of the fluid.
4. Brownian motion causes the particles of a fluid to be in constant motion.

For your information

Albert Einstein explained the pollen movement in a liquid assisted by the molecules in 1905. In 1908, a French physicist J Perrin experimentally verified Einstein's explanation which earned him the 1926 Nobel Prize in physics.

13.2 KINETIC THEORY OF GASES

The kinetic theory of gases describes the molecular composition of the gas and their motion. In this theory, the gas pressure arises due to particles interacting with each other and the walls of the container. It also defines properties such as temperature, volume and pressure, as well as transport properties such as viscosity, thermal conductivity and diffusivity.

The kinetic molecular theory applies to the ideal gases and it has been defined in the previous class. However, the basic assumptions are:

1. Gases consist of very large number of tiny spherical particles that are far apart from one another compared to their sizes.
2. Gas particles are in constant rapid motion in random directions.
3. Collisions between gas particles and the container walls are perfectly elastic.

Do you know?

A small 'cube' of air can have as many as 10^{20} molecules. Volume of the gas molecules is negligible compared to the volume of the empty spaces.

4. There are no forces of attraction or repulsion between gas particles until the particles perform collisions with each other.
5. The average kinetic energy of gas particles is dependent upon the temperature of the gas.

For your information

The mean free path is average distance a moving particle covered before it undergoes a collision that significantly alters its direction.

13.3 PRESSURE IN A GAS AND DERIVATION OF $PV = \frac{1}{3}Nm \langle v^2 \rangle$

We can use the kinetic theory of the gas to derive an equation which relates the macroscopic properties of a gas (pressure and volume) to the microscopic properties of its molecules (mass and speed). It is assumed that a single molecule is moving in a cube-like box of each side L (Fig.13.2). This molecule has mass m , and is moving with speed v . It moves back and forth, colliding at regular intervals with the walls and therefore, it contributes to the pressure of the gas. We are going to work out the pressure exerted by this one molecule on the wall of the box and then deduce the total pressure produced by all the molecules.

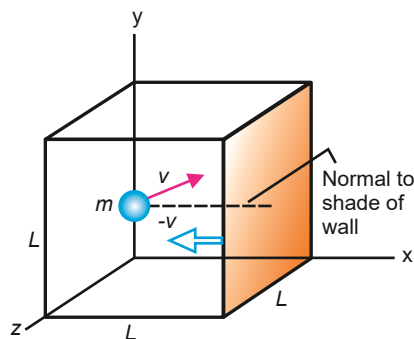


Fig. 13.2: Molecule moving with velocity v

According to the kinetic theory of gases, a gas consists of a large number of tiny particles (atoms or molecules) that are in constant random motion. These particles move freely in all directions and frequently collide with one another and with the walls of the container. When gas molecules strike the walls, they exert a force on them due to the change in momentum during each collision. The pressure of a gas is defined as the force exerted per unit area of the container walls as a result of these collisions. Since the collisions are perfectly elastic, no energy is lost, and the average kinetic energy of the molecules is directly proportional to the absolute temperature of the gas i.e., $\langle K.E. \rangle \propto T$. Thus, an increase in temperature increases the speed of the molecules, leading to more collisions, and consequently, a higher pressure.

An ideal gas obeys the ideal gas law which relate the gas temperature with other parameters and the gas law is:

$$PV = nRT \quad \dots\dots\dots (13.1)$$

where P is the pressure, n is number of moles, R is general gas constant equals to $N_A k_B$, $k_B = 1.38 \times 10^{-23} \text{ J K}^{-1}$ is the Boltzmann constant, ($N_A = 6.02 \times 10^{23} \text{ mol}^{-1}$) is the Avogadro number and T is the temperature of the gas. Therefore,

$$PV = nN_A k_B T$$

The number of moles n in terms of number of particle N is defined as; $n = \frac{N}{N_A}$, then

Brain teaser

What is the difference between mole (gram-molecule) and Avogadro's number? Does Avogadro's number depend on state of the substance?

$$PV = Nk_B T \dots\dots\dots (13.2)$$

As pressure is force per unit area, exerted on walls of the container, each time a gas particle (molecule or atom) collides with walls of the container, and exerts a force on the walls of the container. This force spreads over whole wall area of the container.

Consider a collision in which the molecule with velocity v_x strikes one face of the cube along x-axis. It rebounds elastically in the opposite direction having velocity $-v_x$, its momentum changes from mv_x to $-mv_x$.

Change in momentum along x-axis;

$$\Delta p = -mv_x - (mv_x) = -mv_x - mv_x = -2mv_x$$

The molecule travels a distance $2L$ between the consecutive collisions with the same wall of the cube. Hence, time between two consecutive collisions with one side of the cube is:

$$\Delta t = \frac{2L}{v_x} \dots\dots\dots (13.3)$$

The force on the molecule due to wall is: $F = \frac{\Delta p}{\Delta t} = \frac{-2mv_x}{\frac{2L}{v_x}} = \frac{-mv_x^2}{L}$

The average force exerted on the wall of cube by the molecule can be found using Newton's second law of motion. We know that force is equal to the rate of change of momentum. The area of the wall is L^2 . Thus,

$$F = \frac{\text{Change in momentum}}{\text{Time}} = - \left(\frac{-mv_x^2}{L} \right) = \frac{mv_x^2}{L} \dots\dots\dots (13.4)$$

This expression represents the force exerted on the container wall by one molecule in the x-direction. Thus, pressure exerted on the container wall by a molecule is:

$$\text{Pressure} = \frac{\text{Force}}{\text{Area}} = \frac{F}{L^2}$$

Using equation,
$$P = \frac{mv_x^2}{L} \bigg/ L^2$$

Brain teaser

How does the Kinetic theory of gases relate to the Boyle's law and Charles' law?

$$P = \frac{mv_x^2}{L^3} = \frac{mv_x^2}{V} \dots\dots\dots 13.5)$$

This is the pressure exerted on container wall by one molecule.

The total force, we must add up the contributions of all the molecules that strike the wall, allowing for the possibility that they all have different speeds. Dividing the magnitude of the total force F_x along x-axis by the area of the wall (L^2), then gives the pressure P on that wall:

$$P = \frac{m}{V} (v_{1x}^2 + v_{2x}^2 + v_{3x}^2 + \dots\dots\dots) \dots\dots (13.6)$$

where v_{1x}^2 is the velocity of particle 1 in the x-direction, and v_{2x}^2 the velocity of particle 2 in the x-direction and so on. If there are N number of particles in the container, then total

mass is the number of particles times mass of each particle (i.e. $M = mN$).

As we know that velocity striking the wall is v_x and after striking the wall it is $-v_x$ which would result in average velocity zero. This means simple average does not work here. In order to avoid this zero, we need to do the square of the velocities before averaging.

$$\langle v_x^2 \rangle = \frac{v_{1x}^2 + v_{2x}^2 + v_{3x}^2 + \dots + v_{Nx}^2}{N} = \frac{\sum_{i=1}^N v_{ix}^2}{N} \quad \dots\dots\dots(13.7)$$

If we take the square root of the average square of the velocities, it is known as the root mean square (rms) velocity and is written as:

$$(v_{\text{rms}})_x = \sqrt{\langle v_x^2 \rangle} \quad \dots\dots\dots(13.8)$$

Since the molecules are moving randomly in all directions, the mean-square speed is shared equally among three axes.

$$\text{So } \langle v^2 \rangle = \langle v_x^2 \rangle + \langle v_y^2 \rangle + \langle v_z^2 \rangle \quad \dots\dots\dots(13.9)$$

Since gas molecules are moving randomly and have the same average speed in all three directions. This means that;

$$\langle v^2 \rangle = \langle v_x^2 \rangle + \langle v_y^2 \rangle + \langle v_z^2 \rangle = 3 \langle v_x^2 \rangle \quad \dots\dots\dots(13.10)$$

Total mean-square velocity component along one axis is:

$$\langle v_x^2 \rangle = \frac{\langle v^2 \rangle}{3} \quad \dots\dots\dots(13.11)$$

The pressure is written as:

$$P = \frac{N}{3V} m \langle v^2 \rangle \quad \dots\dots\dots(13.12)$$

$$\text{Then } PV = \frac{1}{3} Nm \langle v^2 \rangle \quad \dots\dots\dots(13.13)$$

This is the pressure of the gas molecules on the wall. The pressure in terms of density can be written as:

$$\begin{aligned} \rho &= \frac{\text{Total mass}}{\text{Volume}} \\ &= \frac{M}{V} = \frac{Nm}{V} \end{aligned}$$

Rearranging the pressure, Eq. (13.12) for P , and substituting the density

$$P = \frac{1}{3} \rho \langle v^2 \rangle \quad \dots\dots\dots(13.14)$$

This is the pressure of the gas which depends on the density and root mean square of velocity.

Equation (13.13) can be written as:

$$P = \frac{2}{3} \frac{N}{V} < \frac{1}{2} m v^2 >$$

$$P = \frac{2}{3} N_0 < \frac{1}{2} m v^2 > \quad (\text{Here } \frac{N}{V} = N_0)$$

$$P = \frac{2}{3} N_0 < K.E._T > \dots\dots\dots (13.14-a)$$

The above expression shows that pressure is directly proportional to the average translational *K.E.* of the molecules of a gas.

Example 13.1 An ideal gas has a density of 4.5 kg m^{-3} at a pressure of $9.3 \times 10^5 \text{ Pa}$ and a temperature of 504 K . Determine the root-mean-square speed of the gas atoms.

Solution $\rho = 4.5 \text{ kg m}^{-3}$, $P = 9.3 \times 10^5 \text{ Pa}$, $v_{\text{rms}} = ?$

Pressure is defined as; $P = \frac{1}{3} \rho < v^2 >$

or $< v^2 > = \frac{3P}{\rho}$

$$< v^2 > = \frac{3 \times (9.3 \times 10^5 \text{ Pa})}{4.5 \text{ kg m}^{-3}} = 6.2 \times 10^5 \text{ m}^2 \text{ s}^{-2}$$

To find the rms value, take the square root of the mean-square-speed, thus

$$v_{\text{rms}} = \sqrt{< v^2 >} = \sqrt{6.2 \times 10^5 \text{ m}^2 \text{ s}^{-2}} = 787 \text{ m s}^{-1} \approx 7.9 \times 10^2 \text{ m s}^{-1}$$

Example 13.2 What is the rms speed of air molecules (O_2 and N_2) at room temperature e.g. 25°C ?

Solution Mass of oxygen $\text{O}_2 = 32 \text{ u}$ and Mass of nitrogen $\text{N}_2 = 28 \text{ u}$
 $1 \text{ u} = 1.66 \times 10^{-27} \text{ kg}$; $k_B = 1.38 \times 10^{-23} \text{ J K}^{-1}$; $T = 25^\circ\text{C} = 298 \text{ K}$
 Mass of oxygen $\text{O}_2 = 32 \times 1.66 \times 10^{-27} \text{ kg} = 5.3 \times 10^{-26} \text{ kg}$
 Mass of Nitrogen $\text{N}_2 = 28 \times 1.66 \times 10^{-27} \text{ kg} = 4.6 \times 10^{-26} \text{ kg}$
 The rms speed of oxygen is:

$$v_{\text{rms}} = \sqrt{\frac{3k_B T}{m}} = \sqrt{\frac{3 \times (1.38 \times 10^{-23} \text{ J K}^{-1}) \times 298 \text{ K}}{5.3 \times 10^{-26} \text{ kg}}} = 481 \text{ m s}^{-1}$$

Similarly, the rms speed of nitrogen is $v_{\text{rms}} = 515 \text{ m s}^{-1}$

The speed of air molecules is much greater than the speed of sound which is about 340 m s^{-1} .

13.4 AVERAGE TRANSLATIONAL KINETIC ENERGY OF A GAS

We have calculated the pressure of the gas molecule with $n = N/V$ molecules by using the gas law; $PV = nRT$ and kinetic theory.

We know that; $PV = \frac{N}{3} m \langle v^2 \rangle$

The gas constant R is related with Boltzmann constant and Avogadro number;

$$k_B = R/N_A$$

Equating these equations, we have

$$nRT = \frac{N}{3} m \langle v^2 \rangle \quad \dots\dots\dots(13.15)$$

or $nRT = \frac{N}{3} m v_{\text{rms}}^2$

or $m v_{\text{rms}}^2 = \frac{3nRT}{N}$

where $\frac{N}{n} = N_A$ is the number of molecules in one mole and it is Avogadro constant.

In terms of molar mass; $m v_{\text{rms}}^2 = \frac{3RT}{N_A}$

or $\frac{1}{2} m v_{\text{rms}}^2 = \frac{3RT}{2N_A}$

or $v_{\text{rms}}^2 = \frac{3RT}{mN_A}$

$$v_{\text{rms}} = \sqrt{\frac{3RT}{mN_A}} \quad \dots\dots\dots(13.16)$$

We rewrite the pressure equation as:

$$\frac{1}{2} m \langle v^2 \rangle = \frac{3}{2} k_B T \quad \dots\dots\dots(13.17)$$

In this equation, $\frac{1}{2} m \langle v^2 \rangle$ is the average translational kinetic energy and it is denoted by $\langle K.E._T \rangle$. Therefore, the above equation is written as:

$$\langle K.E._T \rangle = \frac{3}{2} k_B T \quad \dots\dots\dots(13.18)$$

The average translational kinetic energy of the molecules in a gas is given by a simple constant times the temperature. So, if this model is accurate, the temperature of a gas is a direct measure of the average translational kinetic energy of its molecules.

For an air molecule at room temperature (300 K), the quantity $k_B T$ is $(1.38 \times 10^{-23} \text{ J K}^{-1} \times 300 \text{ K}) = 4.14 \times 10^{-21} \text{ J}$. For such a small quantity, we use electron volt instead of joule (J). The electron volt is the kinetic energy of an electron that has been accelerated through a potential difference of one volt: $1 \text{ eV} = 1.6 \times 10^{-19} \text{ J}$.

Boltzmann constant is $8.62 \times 10^{-5} \text{ eV K}^{-1}$, so at room temperature, value of $k_B T = (8.62 \times 10^{-5} \text{ eV K}^{-1})(300 \text{ K}) = 0.026 \text{ eV}$.

For your information

Kinetic theory helps to understand the interaction between pollutants and the atmosphere. This knowledge is used to develop the systems to monitor the quality of air.

Daily Life Examples of Kinetic Theory of Gases

Examples from daily life are the popcorn and pressure in tyre of a car. When popcorn kernel is heated, the moisture trapped in popcorn kernel becomes steam. The pressure in the kernel is increased and it ruptures the corn cover releasing the gelatinous starch. It becomes solid after cooling.

When air is pumped into the tyres of a car, the number of air molecules increases which raises the air pressure inside the tyres. We know that more molecules result in frequent molecular collisions with tyre's walls due to limited volume. The pressure gives the tyre its hardness and bears the car's weight without deflating.



Fig. 13.3: Popcorn

13.5 KINETIC THEORY AND STATISTICAL PHYSICS

So far we know that kinetic theory of gases is used to determine the pressure of the ideal gas using the gas laws assuming the random motion of molecules. For a large number of molecule, we use statistical physics. It is the branch of physics that uses probability theory. Accordingly, atoms and molecules of a system may exist in different energy states $E_1, E_2, E_3, \dots$, etc due to different speeds. Boltzmann derived the Boltzmann kinetic equation. This equation describes the dynamic processes in gases having large number of molecules. Boltzmann constant is also related to the population of atoms in two levels.

For your information

In 1860 James Clark Maxwell (1831–1879) derived an expression that describes the distribution of molecular speeds within the system at thermal equilibrium. About 60 years later, experiments were performed to confirm Maxwell's predictions.

$$\frac{N_2}{N_1} = e^{-\Delta E/k_B T} \dots \dots \dots (13.19)$$

where $(e^{-\Delta E/k_B T} = \text{Boltzmann factor})$ N_1 and N_2 are the population of lower and higher energy states, $\Delta E = (E_2 - E_1)$ is the energy difference between these states. The number density is directly proportional to the pressure. Equation (13.19) is known as the **Boltzmann distribution law** and is important in describing the statistical mechanics of a large number of molecules. It states that the probability of finding the molecules in a particular energy state varies exponentially as the negative of the energy divided by $k_B T$; therefore, more particles reside in lower energy states than in higher ones.

In a gas container, a molecule undergoes billions of collisions every second. Each collision changes the speed of molecule thereby changing the kinetic energy, but kinetic theory concludes that average kinetic energy at temperature T is:

$$E_{av} = \frac{3}{2} k_B T \dots \dots \dots (13.20)$$

where k_B is effectively the gas constant per molecule.

13.6 STELLAR EVOLUTION

The kinetic theory of gases, when extended to astrophysical systems like stars and galaxies, helps us understand the evolution of stars in a galaxy or gas atoms in a stellar atmosphere.

Stellar evolution deals with the star changes over the time. Star can have range of lifetime from a few million years to trillions of years. The life of a star depends on the mass of the star. All stars are formed from the clouds of gas and dust which are often called nebulae or molecular clouds. This is also called a proto-star. As the cloud contracts, its density and temperature increase due to the rise in $K.E.$ of particles ($K.E. \propto T$). After millions of years, these proto-stars can become a star having achieved a state of equilibrium. The rms velocity of gases in the star increases with rise of density. This in turn increases the average kinetic energy and finally temperature increases with time. Figure 13.4 represents the competition between gravity and pressure.

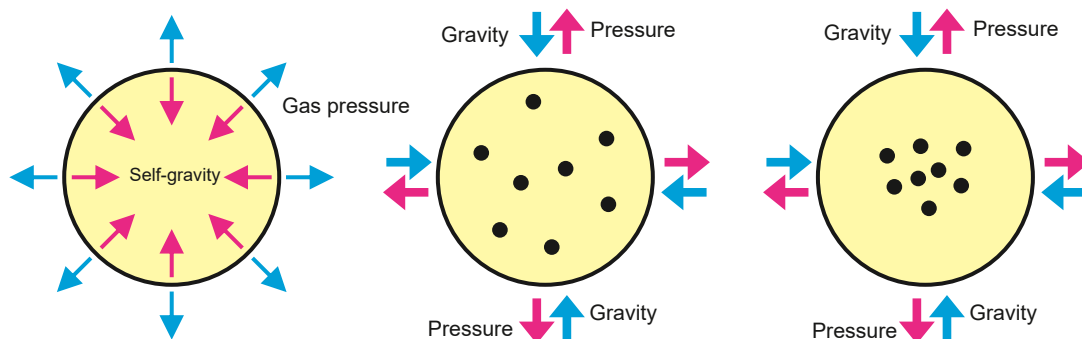


Figure 13.4: Gravity compression is balanced by pressure outward. Greater gravity compresses the gas, making it denser and hotter

As the temperature increases, the internal pressure also rises. A stable star is formed when the inward gravitational force is exactly balanced by the outward pressure. This condition is described by the hydrostatic equation; $\frac{dP}{dr} = -\frac{GM_r \rho_r}{r^2}$.

This equation ensures that the star neither collapses under gravity nor expands indefinitely.

P is the pressure inside the star, dP/dr is the rate of change of pressure with radius (how pressure changes from the centre outward), r is distance from the centre of the star, G is gravitational constant ($6.67 \times 10^{-11} \text{ N m}^2 \text{ kg}^{-2}$), M_r is mass enclosed within radius r , ρ_r is density at radius r , and negative sign indicates that pressure decreases outward, while gravity pulls inward. There are two possibilities if the balance between pressure force and gravitational force is not maintained.

- (i) Gravitational force > Internal pressure
- (ii) Gravitational force < Internal pressure

As the temperature within a star's core rises, nuclear fusion is initiated, allowing hydrogen nuclei to combine and form helium. This process releases an immense amount of energy, which sustains the star and powers it for the majority of its lifetime. The continuous energy production not only maintains equilibrium against gravitational collapse but also leads to a gradual increase in internal temperature. As a consequence, the star begins to expand.

Do you know?

A parsec is a unit of distance that is used by astronomers as an alternative to the light-year, just as kilometres are being used as an alternative to miles. In fact, one parsec is approximately 3.26 light years, or almost 19 trillion miles (31 trillion km)

With further increases in temperature and changes in core composition, the star may expand significantly to become a red giant. At this stage, stars possessing at least about half the mass of the Sun are capable of initiating helium fusion in their cores, producing heavier elements such as carbon and oxygen. In more massive stars, the process continues with the fusion of even heavier elements in successive stages, forming increasingly complex nuclei up to iron.

The final fate of a star depends primarily on its mass. Once a star exhausts its nuclear fuel, it can no longer support itself against gravitational collapse. In low and intermediate mass stars, the core contracts into a dense white dwarf, while the outer layers are expelled into space, forming a planetary nebula. In contrast, stars with masses roughly ten times greater than that of the Sun undergo a dramatic supernova explosion. During this event, the inert iron core collapses under gravity, leading to the formation of dense remnants such as neutron stars or black holes.

Red dwarf stars, due to their low mass and highly efficient fuel consumption, follow a much slower evolutionary path. The universe is not yet old enough for any red dwarf to have completed its life cycle. However, theoretical models indicate that these stars will gradually become hotter and more luminous over time before eventually exhausting their hydrogen fuel and evolving into low mass white dwarfs.

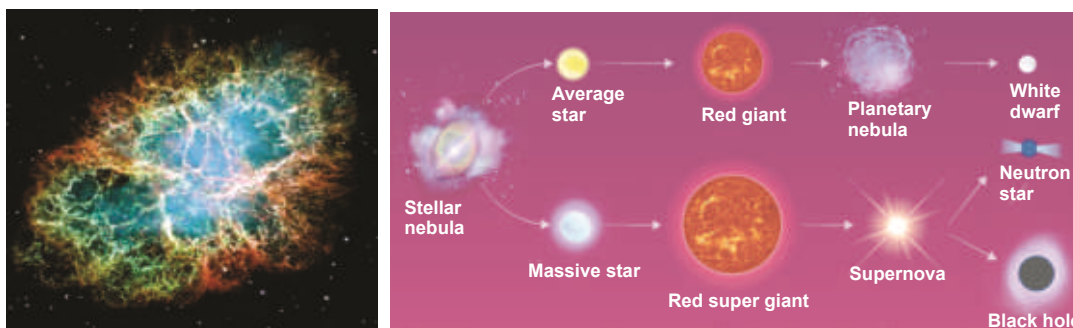


Fig 13.5: Life cycle of a star

Neutron Stars

Neutron stars are among the densest known objects in the universe, second only to

black holes. The nearest are many parsecs away, making direct study difficult. Due to their extremely high density, matter inside neutron stars behaves like a degenerate gas and their strong gravitational fields make them challenging to model. According to estimates by NASA, there may be up to a billion neutron stars in the Milky Way galaxy. Many of the neutron stars observed so far are relatively young and rotate rapidly, emitting beams of radiation. These are known as pulsars.

Scientists believe that pulsar radiation is produced when strong magnetic fields channel matter toward the magnetic poles of neutron stars. When a star collapses to form a neutron star, not only is its mass compressed, but its magnetic field is also greatly intensified. Magnetic fields represented by field lines become stronger as these lines are squeezed closer together during the collapse of the stellar core.

Our Star (SUN)

The Sun is the nearest star to Earth and the main source of heat and light for our planet. It was formed about 4.6 billion years ago from a huge cloud of gas and dust called a solar nebula. Due to gravitational contraction, the temperature at the centre increased and a protostar was formed. When the core became extremely hot, nuclear fusion started in which hydrogen changed into helium and released a large amount of energy. The life stages of the Sun are: nebula, protostar, main sequence, red giant, planetary nebula, and white dwarf. At present, the Sun is in the main sequence stage, where it is producing heat and light by converting hydrogen into helium. In thermal physics, the Sun is important because it is a natural source of enormous heat energy and transfers this energy mainly by radiation.

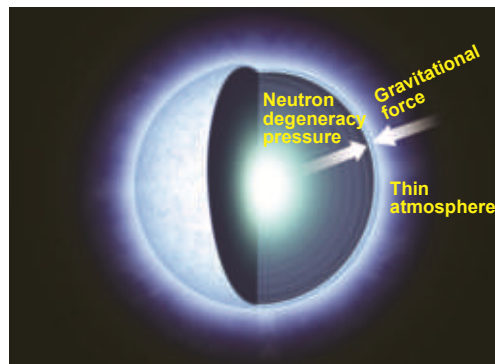


Fig.13.6: Section of neutron star

For your information

A neutron star is so dense that one teaspoon of its material would have a mass over 5.5×10^{12} kg, about 900 times the mass of the Great Pyramid of Giza. The entire mass of the Earth at neutron star density would fit into a sphere 305 m in diameter.

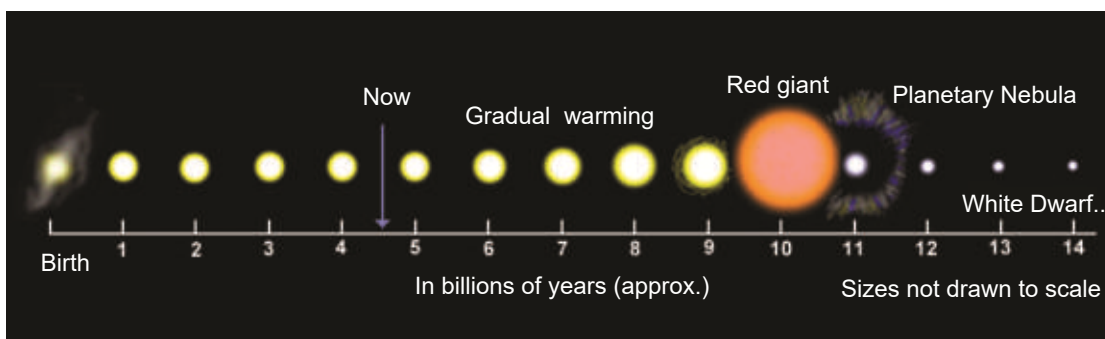


Fig 13.7: Life cycle of the Sun

Example 13.3 A star has mass 2×10^{30} kg and radius 7×10^8 m. Calculate gravitational pressure and compare it with internal pressure 1.0×10^{14} Pa.

Solution Using the formula; $P_g \approx \frac{GM^2}{R^4}$

$$P_g = \frac{(6.67 \times 10^{-11} \text{ Nm}^2 \text{ kg}^{-2})(2 \times 10^{30} \text{ kg})^2}{(7 \times 10^8 \text{ m})^4}$$

$$= \frac{6.67 \times 4 \times 10^{49}}{2.401 \times 10^{35}} \approx 1.1 \times 10^{15} \text{ Pa}$$

$$\text{i.e., } P_g = P_{\text{internal}}$$

As both pressures are nearly equal, therefore, the star is in stable equilibrium.

Example 13.4 A massive star has mass 4×10^{30} kg, radius 5×10^8 m and internal pressure 1×10^{14} Pa. Check whether the star is stable.

Solution As $P_g \approx \frac{GM^2}{R^4}$, therefore,

$$P_g = \frac{(6.67 \times 10^{-11} \text{ Nm}^2 \text{ kg}^{-2})(4 \times 10^{30} \text{ kg})^2}{(5 \times 10^8 \text{ m})^4} = \frac{6.67 \times 16 \times 10^{49}}{6.25 \times 10^{34}} \approx 1.7 \times 10^{16} \text{ Pa}$$

Gravity is stronger, therefore, the star will contract and may eventually collapse (possible neutron star or black hole formation).

Example 13.5 A star has mass 1×10^{30} kg, radius 1×10^9 m and internal pressure 5×10^{14} Pa. Determine the state of the star.

Solution As $P_g \approx \frac{GM^2}{R^4}$, therefore,

$$P_g = \frac{(6.67 \times 10^{-11} \text{ Nm}^2 \text{ kg}^{-2})(1 \times 10^{30} \text{ kg})^2}{(1 \times 10^9 \text{ m})^4} = \frac{6.67 \times 10^{49}}{10^{36}} \approx 6.67 \times 10^{13} \text{ Pa}$$

Internal pressure dominates, therefore, the star will expand and may evolve into a red giant.

QUESTIONS

Multiple Choice Questions

Choose the correct answer.

13.1 Which condition is necessary for most gases to behave nearly ideally?

- (a) Low temperature and low pressure
- (b) High temperature and low pressure
- (c) Constant temperature and low pressure
- (d) High temperature and constant pressure

- 13.2 v_{rms} of the gas molecule is 300 m s^{-1} . If its absolute temperature is reduced to half and molecular weight is doubled, then v_{rms} will become:
(a) 75 m s^{-1} (b) 150 m s^{-1} (c) 300 m s^{-1} (d) 600 m s^{-1}
- 13.3 The rms speed of the gas at 27°C is v . If temperature of the gas is raised to 327°C , the rms speed of the gas is:
(a) v (b) $\frac{v}{\sqrt{2}}$ (c) $\sqrt{2} v$ (d) $3v$
- 13.4 Ratio of the rms velocities of O_2 and H_2 at equal temperature will be:
(a) 1:1 (b) 1:4 (c) 2:1 (d) 4:1
- 13.5 In a closed room, there is a gas with pressure of $3.2 \times 10^5 \text{ N m}^{-2}$. If the gas density is 6 kg m^{-3} , what is the effective speed of each gas particle?
(a) 400 m s^{-1} (b) 300 m s^{-1} (c) 250 m s^{-1} (d) 330 m s^{-1}
- 13.6 Which statement is not the statement of kinetic theory of gases?
(a) Molecules interact during collisions
(b) Molecules are in continuous random motion
(c) Collisions are of short duration
(d) Molecules are tiny hard spheres undergoing inelastic collisions
- 13.7 Which equation ensures that a star remains stable against gravitational collapse?
(a) Mass continuity equation (b) Energy gravitational equation
(c) Hydrostatic equilibrium equation (d) Energy transport equation
- 13.8 What happens to a star when its gravitational pressure becomes greater than its internal pressure?
(a) The star expands and becomes a red giant
(b) The star contracts and its core temperature increases
(c) The star remains stable with no changes
(d) Nuclear fusion stops immediately

Short Answer Questions

- 13.1 Do all gases have the same kinetic energy at the same temperature? If yes, give example.
- 13.2 A vessel is fitted with a mixture of two different gases. Will the mean kinetic energy per molecule of gases be equal? Explain briefly.
- 13.3 If density of a gas is doubled, keeping all other factors unchanged. What will be the effect on the pressure of the gas?
- 13.4 A gas enclosed in a container is heated up. What is the effect of pressure on gas molecules?
- 13.5 The number of molecules of a gas in a container is doubled. What will be the effect on the rms speed?
- 13.6 Explain briefly, why it is not possible to increase the temperature of the gas while keeping its volume and pressure constant.
- 13.7 On reducing the volume of a gas at constant temperature, the pressure of the gas increases. Explain briefly on the basis of kinetic theory of gases.
- 13.8 How does the balance between gravitational force and internal pressure

determine whether a star expands or contracts during its evolution?

Constructed Response Questions

- 13.1 What will be the kinetic energy near absolute zero? Explain.
- 13.2 What is the average kinetic energy of the oxygen and nitrogen molecules in a room at room temperature?
- 13.3 Consider a gas enclosed in a sealed container. By what factor does the gas temperature change if: (a) the volume and pressure are both doubled? (b) the volume is halved and the pressure is tripled?
- 13.4 Is it possible to boil water at room temperature without heating? Explain.
- 13.5 Explain how the drop of ink in a beaker containing water will behave.

Comprehensive Questions

- 13.1 Derive the expression of pressure exerted by the gas on the walls of the container.
- 13.2 Write the expression for rms speed, and mean speed of a gas molecule.
- 13.3 Explain the Boltzmann distributions law on the basis of statistical physics.
- 13.4 Describe the role of pressure in a neutron star.
- 13.5 Show that temperature of a gas is a direct measure of the average translational K.E. of its molecules.
- 13.6 Explain stellar evolution in terms of the balance between gravitational pressure and internal pressure, and how stellar mass determines whether a star ends as white dwarf, neutron star or black hole.

Numerical Problems

- 13.1 Calculate the rms velocity of hydrogen molecules at standard temperature and pressure (STP). Density of hydrogen at STP; $\rho = 8.957 \times 10^{-2} \text{ kg m}^{-3}$. Density of mercury = 13600 kg m^{-3} , $g = 9.8 \text{ m s}^{-2}$.
(Ans: $1.84 \times 10^3 \text{ m s}^{-1}$)
- 13.2 Determine the pressure of oxygen gas at 0°C if the density of oxygen is 1.44 kg m^{-3} at STP and rms speed at STP is 456.4 m s^{-1} .
(Ans: $1.0 \times 10^5 \text{ N m}^{-2}$)
- 13.3 At what temperature will the rms speed of the molecules of gas be three times its value at STP?
(Ans: 2184°C)
- 13.4 If a gas is at temperature 80°C and pressure $5 \times 10^{-10} \text{ N m}^{-2}$, having 1 m^3 volume, what is the number of molecules per cubic metre, where Boltzmann's constant is $1.38 \times 10^{-23} \text{ J K}^{-1}$.
(Ans: $1.02 \times 10^{11} \text{ molecules}$)
- 13.5 The temperature and pressure in the Sun's atmosphere are $2.00 \times 10^6 \text{ K}$ and 0.0300 Pa respectively. Calculate the rms speed of free electrons (mass of the electron = $9.11 \times 10^{-31} \text{ kg}$).
(Ans: $9.5 \times 10^6 \text{ m s}^{-1}$)
- 13.6 A star has a mass enclosed within radius r equal to $2.0 \times 10^{30} \text{ kg}$, density 1500 kg m^{-3} and radius $7.0 \times 10^8 \text{ m}$. Calculate the gravitational force inside the star.
(Ans: $2.7 \times 10^2 \text{ N kg}^{-1}$)

Simple Harmonic Motion

Students Learning Outcomes

After studying this chapter, the students will be able to:

- ◆ describe simple examples of free oscillations.
- ◆ use the terms displacement, amplitude, period, frequency, angular frequency and phase difference in the context of oscillations.
- ◆ express the period of simple harmonic motion in terms of both frequency and angular frequency.
- ◆ express that simple harmonic motion occurs when acceleration is proportional to the displacement from a fixed point and in opposite direction.
- ◆ use $a = -\omega^2 x$ to solve problem.
- ◆ use the equation; $x = x_0 \cos \omega t$ and $v = \pm \omega \sqrt{x_0^2 - x^2}$.
- ◆ analyze graphical representation of the variation of displacement, velocity and acceleration for simple harmonic motion.
- ◆ analyze the interconversion of kinetic energy and potential energy during simple harmonic motion.
- ◆ apply $(1/2)m\omega^2 x_0^2$ for total energy of a system undergoing simple harmonic motion.
- ◆ describe that resistive force acting on an oscillating system causes damping.
- ◆ use the terms light, critical and heavy damping.
- ◆ sketch displacement time graph to illustrate light, critical and heavy damping.
- ◆ state that resonance involves a maximum amplitude of oscillation and this occurs when an oscillating system is forced to oscillate at its natural frequency.
- ◆ describe practical examples of free and forced oscillations.
- ◆ describe practical examples of damped oscillations. (with particular reference to the effects of the degree of damping and the importance of critical damping in cases such as a car suspension system).
- ◆ justify qualitatively the factors which determine the frequency response and sharpness of the resonance.
- ◆ identify the use of standing waves and resonance in applications. [(such as rubens tubes, Chladni plates and acoustic levitation knowledge of wave harmonic modes is not required)].
- ◆ justify the importance of critical damping in car suspension system.
- ◆ justify that there are some circumstances in which resonance is useful (such as tuning a radio, microwave oven and other circumstances in which resonance should be avoided such as airplanes, swing or a suspension bridge).

In this chapter, we will discuss concept of simple harmonic motion (S.H.M), a special type of oscillatory motion where a restoring force always pulls the object back towards its mean position. We will also explore the mathematics and physics behind S.H.M, understanding its characteristics, equations and real world applications.

Interesting Information

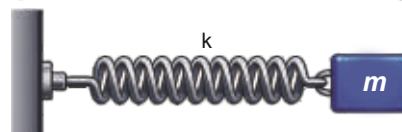
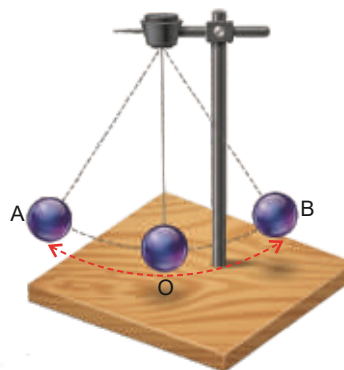
Many natural phenomena show periodic motion. Like motion of Earth around the Sun, The motion of the moon around the Earth, water waves, sound waves, a tuning fork, a rocking chair and many others.

14.1 OSCILLATORY MOTION

Oscillatory motion is a to-and-fro motion about its mean position and it is a periodic motion which repeats itself after an equal interval of time.

We come across many examples of oscillatory motion in our daily life. Some examples of oscillatory motion are given below:

- (i) A simple pendulum vibrating about its mean position when displaced from its rest position.
- (ii) A mass oscillating at the end of a spring in a horizontal or vertical plane when pulled and released.
- (iii) The atoms or molecules in a solid substance oscillate about their mean position.
- (iv) Motion of a swing.
- (v) Air molecules oscillate when sound waves travel through air.
- (vi) Most of musical instruments like guitar, and violin have strings attached to them to produce music by vibratory motion.



All the bodies that undergo vibrational or oscillation motion have an equilibrium position or mean position. When the body is displaced from this mean position then there is restoring force which brings it back to its equilibrium position and it causes oscillatory motion of the body.

Terms Related to Oscillatory Motion

We will discuss some important terms related to oscillatory motion.

Oscillation: One complete round trip (cycle) of vibrating body about its mean position

is called oscillation. Alternatively, oscillation can also be defined motion of the body from one extreme position to other extreme position and back to first extreme position crossing mean position. For example, motion of bob of pendulum from A to B and back from B to A is one oscillation as shown in Fig. 14.1.

Instantaneous Displacement (x): The distance of a vibrating body from its mean position at any instant, is called instantaneous displacement. The distance of bob of simple pendulum from O to C is called displacement and denoted by x . i.e., $OC = x$.

Amplitude (x_0): "The magnitude of the maximum displacement of the vibrating body on either side of its mean position is called amplitude denoted by x_0 , the amplitude of bob of simple pendulum is shown in Fig. 14.1 is OA or OB. i.e., $OA = OB = x_0$.

Time Period: Time period is defined as the time taken to complete one vibration or one cycle. It is represented by T and its SI unit is second (s).

Time period in terms of frequency f and angular frequency ω is expressed as:

$$T = \frac{1}{f} = \frac{2\pi}{\omega}$$

Frequency: Frequency is defined as the number of vibrations or oscillations n completed by the vibrating body per unit time. It is denoted by f .

Mathematically,

$$f = \frac{n}{t}$$

It is expressed in terms of the reciprocal of time period. i.e.; $f = \frac{1}{T}$

The unit of frequency is hertz (Hz). One hertz is defined as one oscillation per second. The dimension of frequency is $[T^{-1}]$.

Angular Frequency: Angular displacement per unit time is called angular frequency. It is represented by ω and it can be expressed as:

$$\omega = \frac{\theta}{t}$$

Now for one revolution $\theta = 2\pi$ radian and $t = T$ (time period).

So
$$\omega = \frac{2\pi}{T}$$

As
$$f = \frac{1}{T}, \quad \text{therefore,} \quad \omega = 2\pi f$$

The SI unit of angular frequency is rad s^{-1} .

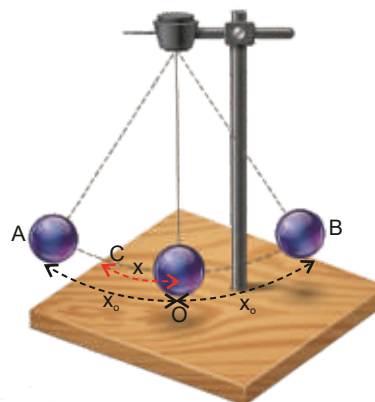


Fig. 14.1: Illustration of different terms of oscillatory motion with the help of vibrating simple pendulum

Brain teaser

If x_0 is the amplitude of simple pendulum, then what will be distance covered in one cycle (vibration) by a simple pendulum?

Brain teaser

A fish connected to a spring makes 10 vibrations in 20 seconds. What is its period and frequency?

14.2 SIMPLE HARMONIC MOTION

Simple harmonic motion is a type of oscillation or vibratory motion produced under the action of restoring force.

The necessary conditions for the execution of simple harmonic motion are:

- The restoring force shall be directly proportional to the displacement from the mean position.
- The force and displacement should follow Hooke's law ($F = -kx$), where k is spring constant depends upon the nature of material of spring.
- Acceleration of an oscillating object should be proportional to the displacement ($a \propto -x$), the negative sign indicates that acceleration is directed towards the mean position.

Interesting Information

It must be noted that not all periodic vibrations are examples of simple harmonic motion since all restoring forces are not always proportional to the displacement. Any restoring force can cause oscillatory motion. An electrocardiogram traces the periodic pattern of a beating heart, but the motion of the recorder needle is not a simple harmonic motion. As the restoring force in this case is not always proportional to the displacement from the mean position.

14.3 PRACTICAL S.H.M SYSTEMS

Mass Attached to An Elastic Spring

Consider a body of mass m attached with a horizontal spring of spring constant k lying on smooth surface (Fig. 14.2). Initially the body is at rest position O , called mean position. Now we apply some force F on the body and displace the body from O to A through displacement x towards right. According to Hooke's law, the applied force is proportional to displacement, i.e.,

$$F \propto x \text{ or } F = kx \dots\dots\dots (14.1)$$

The spring will exert the force on the body due to elastic restoring force at the same time which is equal in magnitude and opposite in direction to the applied force.

$$F = -kx \dots\dots\dots (14.2)$$

When the body is released by removing the applied force, it will move towards B and will cross the mean position O due to inertia and reaches point B , compresses the spring; it then returns and starts oscillating between A and B .

Also, according to Newton's second law of motion:

$$F = ma \dots\dots\dots (14.3)$$

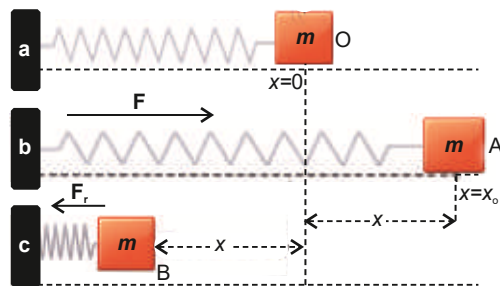


Fig. 14.2: Mass attached with spring performing S.H.M

Comparing Eqs. (14.2) and (14.3), we have

$$ma = -kx$$

or
$$a = -\frac{k}{m}x \quad \text{..... (14.4)}$$

where k is spring constant depending upon the nature (physical shape and structure of a spring). The negative sign shows that acceleration and displacement are always in opposite direction or acceleration is always directed towards the mean position.

As, k/m is a constant, so, Eq. (14.4), can be written as:

$$a = -(\text{constant})x$$

or
$$a \propto -x$$

This is mathematical form of simple harmonic motion.

It shows that acceleration of the body executing S.H.M is directly proportional to the displacement and is always directed towards the mean position.

Time period and frequency

The frequency of the spring mass system is given by the relation:

$$f = \frac{1}{2\pi} \sqrt{\frac{k}{m}}$$

and the time period is:

$$T = \frac{1}{f}$$

$$= 2\pi \sqrt{\frac{m}{k}}$$

Point to ponder!

Real world applications that utilize S.H.M principle include:

Pendulum clocks, musical instruments, and vehicle suspension system etc.

Example 14.1 A particle executing simple harmonic motion has an angular frequency of 2 rad s^{-1} . Determine its frequency and time period.

Solution

Angular frequency $\omega = 2 \text{ rad s}^{-1}$, Frequency $f = ?$

Time period $T = ?$

As $\omega = 2\pi f$, therefore,

$$f = \frac{\omega}{2\pi}$$

$$f = \frac{2 \text{ rad s}^{-1}}{2\pi \text{ rad}}$$

$$f = \frac{1}{\pi} \text{ s}^{-1}$$

$$f = 0.318 \text{ Hz}$$

$$T = ?$$

As $T = \frac{1}{f}$

$$T = \frac{1}{0.318 \text{ Hz}}$$

$$T = 3.14 \text{ s}$$

Interesting Information



Dutch Mathematician astronomer and physicist Christiaan Huygens patented the first pendulum clock in 1656. He finalized the mathematical formula that relates pendulum length to time, using simple harmonic motion.

Simple Pendulum

A simple pendulum consists of a small heavy bob of mass m suspended from a rigid support through a light inextensible string of certain length.

When the pendulum is displaced from its mean position O through an angle θ to the extreme position P and released, then a restoring force acts on the pendulum towards the mean position. Due to this restoring force, the pendulum starts oscillation to and fro under the action of gravity along a curved path about the mean position O (Fig. 14.3).

At extreme position P, the weight of the body makes an angle θ with the direction of the string. We can resolve weight into its rectangular components. As the pendulum has no motion along the direction of the string, therefore, the component $mg\cos\theta$ and tension T are along the same line but in opposite direction, so they cancel the effect of each other. The component $mg\sin\theta$ provides the necessary restoring force towards mean position and is responsible for the motion of simple pendulum. Thus,

$$F = -mg\sin\theta$$

Negative sign shows that the acceleration of the pendulum is always directed towards the mean position. According to Newton's 2nd law:

$$F = ma$$

Comparing above equations

$$ma = -mg\sin\theta$$

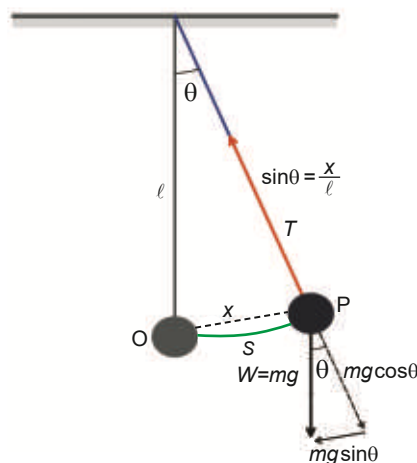


Fig. 14.3: A simple pendulum executing S.H.M

Brain teaser

All vibrating bodies produce sound. Why a vibrating simple pendulum does not produce sound?

$$a = -g \sin \theta \quad \dots\dots\dots (14.5)$$

For small values of θ , $\sin \theta$ is nearly equal to θ , thus,

$$\sin \theta \approx \theta$$

So, Eq. (14.5) is written as:

$$a = -g \theta \quad \dots\dots\dots (14.6)$$

By the definition of angular displacement,

$$\theta = \frac{S}{\ell}$$

where S is the actual path length followed by the pendulum. Thus,

$$a = -\left(\frac{g}{\ell}\right)S$$

In Fig. 14.3, x is nearly equal to the arc of length S of the circular path for a small angle θ . Hence,

$$a = -\left(\frac{g}{\ell}\right)x \quad \dots\dots\dots (14.7)$$

If length ℓ of the pendulum is fixed and g remains constant for a given place and (g/ℓ) is constant, then Eq. (14.7) can be rewritten as;

$$a = -(\text{constant})x$$

$$a \propto -x$$

This is the mathematical form of S.H.M and it is concluded that the motion of a simple pendulum is S.H.M.

$$\text{As } a = -\omega^2 x \quad \dots\dots\dots (14.8)$$

where ω is the angular frequency of oscillation.

Comparing Eqs. (14.7) and (14.8), we have

$$\omega^2 = \frac{g}{\ell}$$

$$\text{or } \omega = \sqrt{\frac{g}{\ell}} \quad \dots\dots\dots (14.9)$$

As $\omega = 2\pi f$, therefore,

$$2\pi f = \sqrt{\frac{g}{\ell}} \quad \text{or} \quad f = \frac{1}{2\pi} \sqrt{\frac{g}{\ell}}$$

$$T = \frac{1}{f} = 2\pi \sqrt{\frac{\ell}{g}} \quad \dots\dots\dots (14.10)$$

Table 14.1: Values of θ and $\sin \theta$ for small angle

θ (degree)	θ (radian)	$\sin \theta$
0.00	0.0000	0.0000
1.00	0.0175	0.0175
2.00	0.0349	0.0349
3.00	0.0524	0.0524
4.00	0.0698	0.0698
5.00	0.0873	0.0872
6.00	0.1048	0.1046
7.00	0.1222	0.1219
8.00	0.1397	0.1392
9.00	0.1571	0.1565

Brain teaser

A pendulum is shifted from Lahore to Karachi, will its time period change?

Do you know?

Pendulum is derived from Latin word Pendulus means hanging.

The above expression shows that the time period of simple pendulum is directly proportional to the square root of the length of the string and inversely proportional to the square root of acceleration due to gravity. The time period of motion of the pendulum is independent of the mass m of the bob and amplitude.

A pendulum that completes one vibration in two seconds is known as a second pendulum.

Using simple pendulum to measure value of “ g ”

The simple pendulum can be used to determine the gravitational acceleration at a particular location. We measure the length ℓ of the pendulum and then set the pendulum into motion. The time period T of the simple pendulum is measured using a stopwatch and the acceleration due to gravity is calculated by using Eq. (14.10) in the following form:

$$g = \frac{4\pi^2 \ell}{T^2}$$

Example 14.2 What is the length of a second pendulum where the value of g is 9.8 m s^{-2} .

Solution $T = 2 \text{ s}$

$$\ell = ?$$

As $T = 2\pi\sqrt{\frac{\ell}{g}}$

or $T^2 = \frac{4\pi^2 \ell}{g}$

or $\ell = \frac{gT^2}{4\pi^2} = \frac{(9.8 \text{ m s}^{-2})(2 \text{ s})^2}{4(3.14)^2}$
 $= 0.994 \text{ m}$

or $\ell = 99.4 \text{ cm}$

Brain teaser

Would the time period of a simple pendulum be same on Earth and moon, if its length is kept constant?

Brain teaser

While measuring the time period of a simple pendulum, it is recommended to take small amplitude, why?

14.4 SIMPLE HARMONIC MOTION AND UNIFORM CIRCULAR MOTION

In order to get an insight into simple S.H.M, we will co-relate it with uniform circular motion. The different parameters such as displacement, velocity, acceleration and time period of simple harmonic motion (S.H.M) can be understood by relating it with uniform circular motion. It has been observed that when a particle moves in a circular path, its projection on diameter of the circle executes S.H.M.

To study the simple harmonic motion, consider a turntable of radius x_0 with a ball attached to its rim. A beam of light casts a shadow of the ball on the screen (Fig 14.4).

When the turntable rotates with constant angular speed ω , then the ball also moves along it with uniform circular motion. Its shadow on the screen oscillates executing to and fro motion across the screen in the form of simple harmonic motion, like a body attached to a spring.

Quantitative Analysis

Consider a motion of particle P along the circumference of circle of radius x_0 with uniform angular velocity ω . Its linear velocity at point P is along the tangent ($v_p = x_0 \omega$). Let Q be the projection which oscillates along the diameter of circle about the mean position O. When the body is at point P, its projection is at distance x from the mean position.

Instantaneous Displacement

If A is the initial position of particle P, then the angle which the radius OP sweeps out in time t is $\theta = \omega t$ with OQ as shown in Fig. 14.5.

Considering triangle OQP:

$$\frac{\overline{OQ}}{\overline{OP}} = \cos \theta$$

$$\text{or} \quad \frac{x}{x_0} = \cos \omega t \quad (\because \theta = \omega t)$$

$$x = x_0 \cos \omega t \dots\dots\dots (14.11)$$

Equation (14.11) gives the instantaneous displacement of point Q which is executing simple harmonic motion.

Instantaneous Velocity

Considering Fig. 14.6, the line PR is the horizontal component of velocity v_p of the particle and it is parallel to the diameter AB of the circle. Therefore, as Q vibrates horizontally, so instantaneous velocity of the point Q is horizontal component of velocity v_p .

$$\frac{v}{v_p} = \cos (90 - \theta)$$

$$v = v_p \cos (90 - \theta)$$

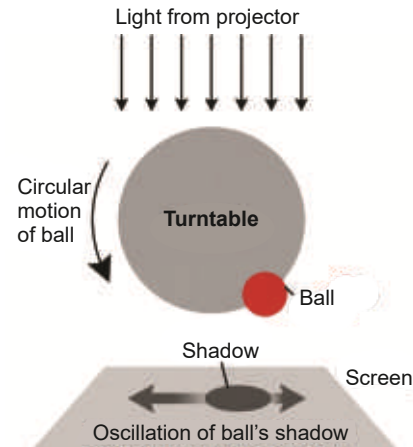


Fig. 14.4: The oscillation of the shadow of ball on screen. The ball is attached with uniformly rotating turntable.

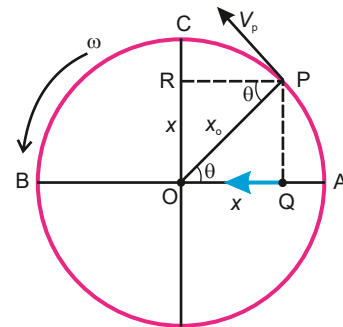


Fig. 14.5: A point P moving on circular path with constant angular velocity

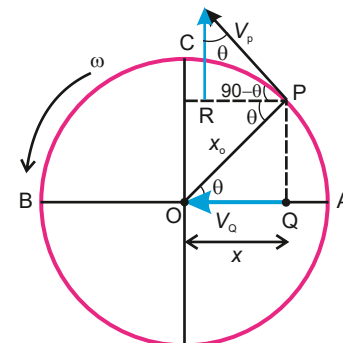


Fig. 14.6: Velocity of the point P and its horizontal component

Since $v_p = x_o \omega$ and $\cos(90 - \theta) = \sin \theta$

$$v = x_o \omega \sin \theta \quad \text{or} \quad v = x_o \omega \sin \omega t \quad \dots\dots (14.12)$$

But $\sin^2 \theta + \cos^2 \theta = 1$, therefore, Eq. (14.12) becomes:

$$\sin \theta = \sqrt{1 - \cos^2 \theta}$$

$$v = x_o \omega \sqrt{1 - \cos^2 \theta}$$

$$v = x_o \omega \sqrt{1 - \frac{x^2}{x_o^2}} \quad \left(\because \cos \theta = \frac{x}{x_o} \right)$$

$$v = \omega \sqrt{x_o^2 - x^2} = \sqrt{\frac{k}{m} (x_o^2 - x^2)} \quad \dots\dots (14.13)$$

It may be noted that at mean position $x = 0$, the velocity is maximum. Thus,

$$v_{\max} = v_o = x_o \omega = \sqrt{\frac{k}{m}} x_o \quad \dots\dots\dots (14.14)$$

Instantaneous Acceleration

Acceleration a_p of the particle at point P is directed towards the centre of the circle as shown in Fig. 14.7. The horizontal component of a_p is along the diameter. Thus, the acceleration of projection Q is equal to the horizontal component of a_p .

$$a = (a_p)_x$$

$$a = a_p \cos \theta$$

As $a = -x_o \omega^2$ (centripetal acceleration) and $\cos \theta = x/x_o$.

The negative sign indicates that the direction of acceleration is always directed towards the mean position.

$$a = -x_o \omega^2 \left(\frac{x}{x_o} \right) = -\omega^2 x$$

$$a = -x \omega^2 \quad \dots\dots\dots (14.15)$$

As particle is moving in the circle with uniform angular frequency ω , therefore, Eq. (14.15) can be rewritten as:

$$a \propto -x$$

This expression is the mathematical condition of S.H.M i.e., acceleration is directly proportional to the displacement and negative sign shows that its direction is towards mean position, therefore, it is concluded that when a particle is moving along a

Interesting Information

The expressions such as time period, displacement, velocity and acceleration which are true for projection of a particle are also applicable for particle moving in a circle.

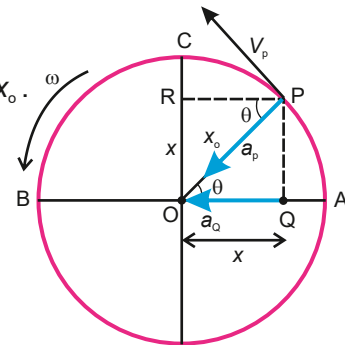


Fig. 14.7: Acceleration of the point P and its horizontal component

circumference of a circle, then its projection executes S.H.M.

Example 14.3 A 0.2 kg mass is attached to a spring with a spring constant of 10 N m^{-1} . The mass is displaced by 0.1 m from its equilibrium position and released from rest. Find: (a) acceleration at $x = 0.05 \text{ m}$ (b) velocity at $x = 0.05 \text{ m}$

Solution

$$m = 0.2 \text{ kg}$$

$$k = 10 \text{ N m}^{-1}$$

$$x_0 = 0.1 \text{ m}$$

Angular frequency $\omega = ?$

$$\begin{aligned} \text{As } \omega &= \sqrt{\frac{k}{m}} = \sqrt{\frac{10 \text{ N m}^{-1}}{0.2 \text{ kg}}} \\ &= \sqrt{50} = 5\sqrt{2} \text{ rad s}^{-1} \approx 7.07 \text{ rad s}^{-1} \end{aligned}$$

(a) For acceleration at $x = 0.05 \text{ m}$

$$a = -\omega^2 x$$

$$\begin{aligned} a &= -(7.07 \text{ rad s}^{-1})^2 \times 0.05 \text{ m} \\ &= -50 \times 0.05 \\ &= -2.5 \text{ m s}^{-2} \end{aligned}$$

(b) For velocity at $x = 0.05 \text{ m}$ $v = \omega \sqrt{(x_0^2 - x^2)}$

$$\begin{aligned} v &= 7.07 \text{ rad s}^{-1} \sqrt{(0.1 \text{ m})^2 - (0.05 \text{ m})^2} \\ &= 7.07 \sqrt{(0.01) - (0.0025)} \\ &= 7.07 \sqrt{0.0075} \\ v &= 7.07 \times 0.0866 \\ &\approx 0.61 \text{ m s}^{-1} \end{aligned}$$

14.5 PHASE

The angle $\theta = \omega t$ which specifies the displacements x as well as the direction of motion of the point oscillating S.H.M. is called phase.

Phase is the quantity which shows the state of motion of an oscillator. In circular motion, the displacement of projection of the body moving in a circle, executing S.H.M. on the diameter of the circle is given by

$$x = x_0 \cos \omega t \dots\dots (14.16)$$

where x is instantaneous displacement, x_0 is maximum displacement; ω is the angular velocity. The general way of showing this equation is:

$$x = x_0 \cos(\omega t + \phi) \dots\dots (14.17)$$

The time varying quantity $(\omega t + \phi)$ is called the phase of the motion. It describes the state of motion at a given time. The constant ϕ is called the phase constant (or phase angle). The value of ϕ depends on displacement and velocity of the particle at $t = 0$.

The quantity ϕ represents the phase difference between the states of motion of two oscillators. Let us explain, with the help of graph plotted between x and t .

If $\phi = 0$, then Eq. (14.17) becomes:

$$x = x_0 \cos \omega t$$

Putting $t = 0, T/4, T/2, 3T/4, \dots$, we obtain a graph, as shown in Fig. 14.8 (a). This graph shows that;

at $t = 0, T/2$ and T (corresponding to $\theta = 0, \pi$ and 2π), the point is at the extreme positions. At $t = T/4$ and $3T/4$ (corresponding to $\theta = \pi/2$ and $3\pi/2$), the point is at mean position.

If $\phi = \pi/2$, then Eq. (14.17) becomes:

$$x = x_0 \cos(\omega t + \pi/2)$$

Putting $t = 0, T/4, T/2, 3T/4, T \dots$, we obtain a graph, as shown in Fig. 14.8 (b). This graph shows that:

at $t = 0, T/2$ and T (corresponding to $\theta = 0, \pi$ and 2π), the point is at mean positions.

At $t = T/4$ and $3T/4$ (corresponding to $\theta = \pi/2$ and $3\pi/2$), the point is at extreme position.

If $\phi = \pi$, then Eq. (14.17) becomes:

$$x = x_0 \cos(\omega t + \pi)$$

Putting $t = 0, T/4, T/2, 3T/4, T \dots$, we obtain a graph, as shown in Fig. 14.8 (c).

It can be noted that:

When the phase difference between two oscillating systems is π , they are said to be oscillating out of phase.

When phase difference between oscillating systems is 0 or 2π , they are said to be oscillating in phase.

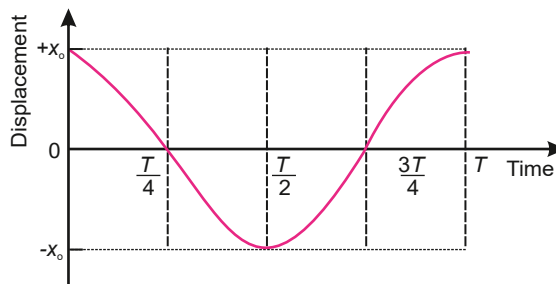


Fig. 14.8 (a): Graph of $x = x_0 \cos(\omega t + \phi)$ for $\phi = 0$

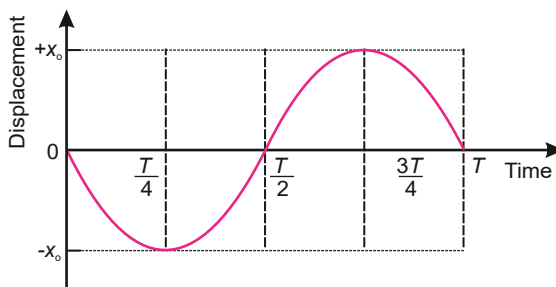


Fig. 14.8 (b): Graph of $x = x_0 \cos(\omega t + \phi)$ for $\phi = \pi/2$

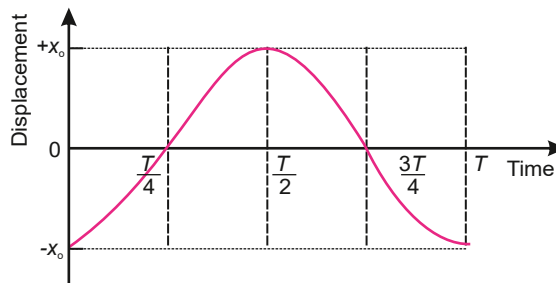


Fig. 14.8 (c): Graph of $x = x_0 \cos(\omega t + \phi)$ for $\phi = \pi$

Example 14.4 The displacement of a particle executing S.H.M is given by

$$x = 0.05 \sin \left(4t + \frac{\pi}{3} \right)$$

Find:

its velocity at the displacement of 0.03 m,
maximum velocity and amplitude.

Solution

As	Displacement	$x = 0.05 \sin \left(4t + \frac{\pi}{3} \right)$(i)
	Velocity	$v = ?$, when $x = 0.03 \text{ m}$
	Max. velocity	$v_o = ?$
	Amplitude	$x_o = ?$

We know that the equation of S.H.M is:

$$x = x_o \sin (\omega t + \phi) \quad \text{..... (ii)}$$

Comparing Eqs. (i) and (ii), we have

Amplitude	$x_o = 0.05 \text{ m}$
Angular frequency	$\omega = 4 \text{ rad s}^{-1}$
Velocity	$v = \pm \omega \sqrt{x_o^2 - x^2}$

Putting the values

$$v = \pm 4 \sqrt{(0.05 \text{ m})^2 - (0.03 \text{ m})^2}$$

$$v = 0.16 \text{ m s}^{-1}$$

$$v_o = \omega x_o = 4 \times 0.05 \text{ m s}^{-1}$$

$$v_o = 0.20 \text{ m s}^{-1}$$

Brain teaser

Two identical pendulums perform S.H.M. with the same amplitude and frequency.

- Pendulum A: $x_A = A \cos \omega t$
- Pendulum B: $x_B = A \sin \omega t$

At $t = 0$:

1. which pendulum has the greater speed?
2. what is their phase difference?
3. will they ever have the same displacement and velocity simultaneously?

14.6 GRAPHICAL REPRESENTATION OF S.H.M

The graphical representations of displacement, velocity and acceleration of a body executing SHM is given in Figs. 14.9 (a, b and c).

(i) The displacement of the particle executing SHM is given by the expression:

$$x = x_0 \cos \omega t \quad \dots\dots\dots (14.18)$$

This shows that the particle will have maximum displacement (x_0) at extreme position.

(ii) The velocity of the particle executing SHM is given by the expression:

$$v = -x_0 \omega \sin (\omega t) \quad \dots\dots (14.19)$$

The velocity of the particle is maximum (i.e., $v = x_0 \omega$) at the mean position and zero at the extreme positions.

(iii) The acceleration of the particle executing SHM is given by the expression:

$$a = -\omega^2 x$$

Putting $x = x_0 \cos \omega t$, we have

$$a = -x_0 \omega^2 \cos (\omega t) \dots\dots (14.20)$$

The acceleration will be maximum (i.e., $x_0 \omega^2$) at the extreme positions, and zero at the mean position.

From the graph of Eqs. (14.18), (14.19) and (14.20) shown in Figs. 14.9 (a, b and c) can be seen that:

- The phase difference between velocity and displacement is $\pi/2$.
- The phase difference between acceleration and displacement is π .

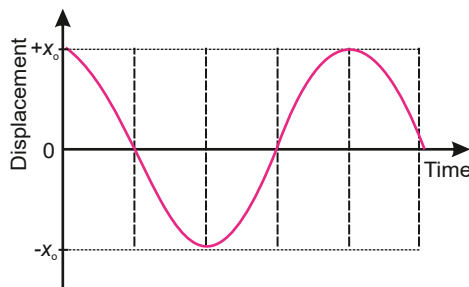


Fig. 14.9 (a): Graphical representation of displacement of oscillation performing S.H.M

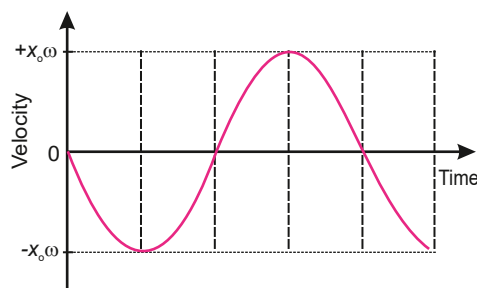


Fig. 14.9 (b): Graphical representation of velocity of oscillation performing S.H.M

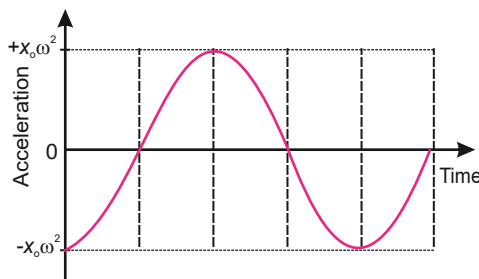


Fig. 14.9 (c): Graphical representation of acceleration of oscillation performing S.H.M

14.7 CONSERVATION OF ENERGY IN S.H.M

When a body is executing simple harmonic motion, it possesses both potential energy as well as kinetic energy. Its potential energy is on account of its displacement from mean position and the kinetic energy is due to its velocity. These energies vary during the oscillation, but the total energy at any instant remains constant in the absence of

unbalanced frictional forces. In case of mass-spring system shown in Fig. (14.10), when the mass is displaced from the mean position O, then there is a restoring force F whose value is zero at mean position, when $x = 0$ and its value is maximum at either extreme position where $x = x_0$. Thus, average value of force from the mean position to the extreme position is:

$$F_{av} = \frac{0 + F}{2} = \frac{F}{2}$$

When displacement = 0 ; Force = 0

When displacement = x_0 ; Force = $F = kx_0$

$$F_{av} = \frac{0 + kx_0}{2}$$

$$F_{av} = \frac{1}{2} kx_0$$

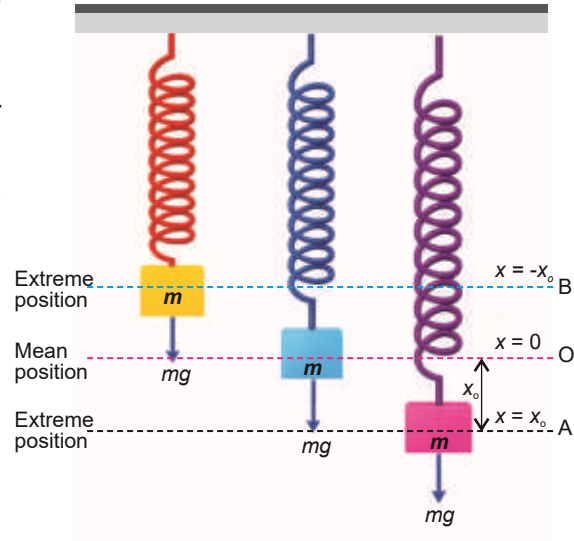


Fig. 14.10: Mass-spring system

When the spring is stretched to its maximum displacement x_0 , work is done on the spring which is given by

$$W = \frac{1}{2} kx_0 (x_0) = \frac{1}{2} kx_0^2$$

This work done on the mass attached to a spring is stored in terms of potential energy, called elastic potential energy. So, we have the expression for $P.E.$ at any instant x given by

$$(P.E.)_{inst} = \frac{1}{2} kx^2 \quad \dots\dots\dots (14.21)$$

It is clear from Eq. (14.21) that the potential energy of oscillator is zero at $x = 0$ and maximum at $x = \pm x_0$ i.e., the extreme position on either side and total energy is entirely elastic potential energy.

$$(P.E.)_{max} = \frac{1}{2} kx_0^2 \quad \dots\dots\dots (14.22)$$

After the removal of force, the mass attached to a spring starts its motion with velocity v .

As the instantaneous velocity v of the mass executing simple harmonic motion is given by

$$v = \omega \sqrt{x_0^2 - x^2}$$

As $\omega = \sqrt{\frac{k}{m}}$, so above Eq. becomes:

$$v = \sqrt{\frac{k}{m}} (x_0^2 - x^2) \quad \dots\dots\dots (14.23)$$

At mean position $x = 0$, the mass gets maximum velocity. i.e.,

$$v_{\max} = v_o = \sqrt{\frac{k}{m}} x_o \quad \dots\dots\dots (14.24)$$

Hence, at any instant where the displacement is x , the kinetic energy of mass attached to spring is given by

$$K.E. = \frac{1}{2} mv^2 = \frac{1}{2} k(x_o^2 - x^2) \quad \dots\dots\dots (14.25)$$

At mean position ($x = 0$), the velocity is maximum, hence, total energy is entirely kinetic energy.

$$T.E. = K.E._{\max} = \frac{1}{2} mv_{\max}^2 = \frac{1}{2} kx_o^2 \quad \dots\dots\dots (14.26)$$

The total energy is the sum of kinetic and potential energy of the system at displacement x , so we may write; $T.E. = K.E. + P.E.$

Substituting the values of $K.E.$ and $P.E.$ from above Eqs., we have

$$T.E. = \frac{1}{2} k(x_o^2 - x^2) + \frac{1}{2} kx^2 = \frac{1}{2} kx_o^2 \quad \dots\dots\dots (14.27)$$

As $\omega = \sqrt{\frac{k}{m}}$ or $k = m\omega^2$, so Eq. 14.27 can be written as:

$$T.E. = \frac{1}{2} m\omega^2 x_o^2 \quad \dots\dots\dots (14.28)$$

From Eqs. (14.22), (14.26) and (14.27), it is proved that total energy of mass-spring system is constant and is directly proportional to the square of the amplitude of oscillation. This statement of conservation of energy is equally valid for all bodies executing S.H.M. When $K.E.$ of mass is maximum, mass passes through the centre of oscillation, the $P.E.$ of mass spring is zero ($x = 0$), conversely when the $P.E.$ of spring is maximum, the mass is at its extreme position on either side, the $K.E.$ of mass is zero ($x = x_o$).

The energy oscillates between $K.E.$ and $P.E.$, but their sum remains constant. This can be illustrated by plotting graph of $K.E.$, $P.E.$ and $T.E.$ verses displacement as shown in Fig. (14.11).

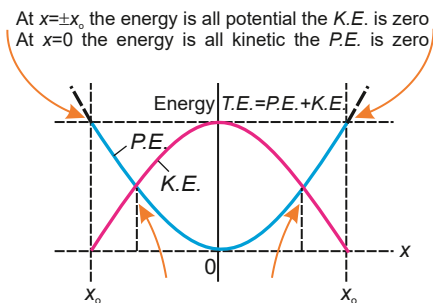


Fig. 14.11: Graphical representation of energy ($K.E.$, $P.E.$ and $T.E.$) against displacement, for S.H.M

Example 14.5 A mass of 200 g is attached to a spring. A force of 40 N is applied and the spring is stretched to 10 cm from its mean position. After release, it performs S.H.M. Find the total energy of the mass during S.H.M, its kinetic energy when displacement is 3 cm. Also find its $P.E.$ when $K.E.$ is equal to its potential energy, and its acceleration at the extreme position.

Solution

Mass $m = 200 \text{ g} = 0.2 \text{ kg}$

Applied force $F = 40 \text{ N}$

Amplitude $x_o = 10 \text{ cm} = 0.1 \text{ m}$

(i) Total Energy $E_T = ?$

(ii) Kinetic Energy $K.E. = ?$, when

$x = 3 \text{ cm} = 0.03 \text{ m}$

(iii) Potential Energy $P.E. = ?$, when

$K.E. = P.E.$

(iv) Acceleration $a = ?$

(i) Total energy $E_T = \frac{1}{2} k x_o^2$

As $F = kx$ (For spring)

or $k = \frac{F}{x}$

$$k = \frac{40 \text{ N}}{0.1 \text{ m}} = 400 \text{ N m}^{-1}$$

So,

$$E_T = \frac{1}{2} (400 \text{ N m}^{-1}) (0.1 \text{ m})^2$$

$$E_T = 2 \text{ J}$$

(ii) $K.E. = \frac{1}{2} k (x_o^2 - x^2)$, $x = 0.03 \text{ m}$

$$= \left(\frac{1}{2} \right) 400 \text{ N m}^{-1} [(0.1 \text{ m})^2 - (0.03 \text{ m})^2]$$

$$= 200(0.01 - 0.0009) \text{ J} = 1.82 \text{ J}$$

(iii) When $K.E. = P.E.$, then $P.E. = \frac{E_T}{2}$

$$P.E. = 1 \text{ J}$$

(iv) Acceleration

$$a = \frac{kx_o}{m}$$

$$a = \frac{(400 \text{ N m}^{-1}) (0.1 \text{ m})}{0.2 \text{ kg}}$$

$$a = 200 \text{ m s}^{-2}$$

14.8 FREE AND FORCED OSCILLATIONS

A body is said to be executing free vibrations if it oscillates with its natural frequency without the interference of an external force.

For example, a simple pendulum shown in Fig. 14.12 vibrates freely with its natural frequency that depends only upon its length when it is slightly displaced from its mean position.

In free oscillations, the total energy of the body remains constant i.e., energy is conserved. As we are assuming in the absence of resistance force the amplitude of the oscillation remains constant.

If a freely oscillating system is subjected to an external force, then forced vibrations will take place. Such as when swing is struck repeatedly, then forced vibrations are produced as shown in Fig. 14.13.

The vibrations of a factory floor caused by the running of heavy machinery are an example of forced vibrations. Another example of forced vibration is loud music produced by sounding wooden boards of string instruments.

14.9 DAMPED OSCILLATION

The oscillations with decreasing amplitude in the presence of various resistive forces are called damped oscillations and the resistive forces are called damping forces.

In an ideal free oscillation system, we have studied that the total mechanical energy of the oscillating body remains constant, as we discussed in mass-spring system, motion of a body in a circle and also in case of a simple pendulum. But in practical life, as we observe daily, some dissipating forces such as air resistance, friction, etc. The amplitude of the oscillation gradually decreases with time due to these forces and finally it comes to rest.

For example, a girl is swinging on a swing as shown in Fig. 14.14. Damping occurs due

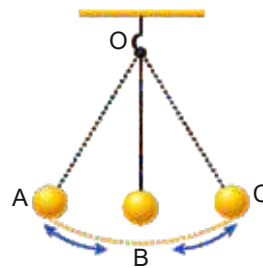


Fig 14.12: Illustration of free oscillation of simple pendulum



Fig 14.13: Illustration of forced oscillations of a swing

Interesting Information

Light damping: It reduces gradually the energy and amplitude of vibrating body. A swing in playground is the best example of light damping.

Heavy damping: A type of damping which causes the system to return to equilibrium very slowly without oscillating. For example, a pendulum submerged in thick oil experiences a heavy resistance from oil.

Critical damping: It is a type of damping in which object returns to equilibrium position in the shortest possible time, e.g., shock absorbers in a car.

to gravity and air friction. The swing will oscillate with smaller and smaller amplitude and eventually stops. Graphically, the damped oscillation of the oscillating body is shown in Fig. 14.15. Similarly, the amplitude of oscillating simple pendulum decreases gradually with time till it becomes zero and touching an oscillating tuning fork with our finger are examples of damped oscillation.



Fig 14.14: Damped oscillation of swing

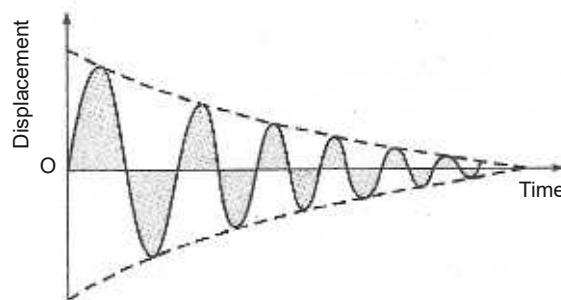


Fig 14.15: Graphical representation of damped oscillation

Application of Damped Oscillation

Shock absorbers used in the suspension system of a car shown in Fig. 14.16 are one practical application of damped oscillation. Damping system is required to ensure a comfortable ride for the passengers when the car is moving on a bumpy, rough road by producing excessive oscillations by damping using shock absorbers in such vehicles.



Fig. 14.16: Suspension system of a vehicle

14.10 RESONANCE

In damped oscillation, the oscillator cannot maintain its natural frequency for long duration due to the resistive forces and the amplitude of the oscillation decreases gradually with time. But we can maintain constant amplitude by applying a periodic external force which is called a driving force. Thus, when the oscillating body is subjected to a periodic driving force, then such oscillation is called forced oscillation and its frequency is called driving frequency. The vibration of a vehicle caused by the running of engine is an example of forced vibration. In forced oscillation, the amplitude of the oscillation depends upon the relation between the driving frequency and the natural frequency of the body.

If the frequency of the driving force is same as the natural frequency of the oscillating body, the amplitude of vibration is very much increased. This phenomenon is known as resonance. It also occurs when the applied force has frequency an integral multiple of the natural frequency of the body.

To demonstrate the resonance phenomenon, we perform a simple experiment. The experimental setup consists of two pair of pendulums A and B and C and D such that the length of A and B is ℓ_1 , and length of C and D is ℓ_2 . All the pendulums are suspended by a horizontal string (Fig. 14.17).

Now we introduce another pendulum P whose length is ℓ . Consider the case when the length of pendulum P is $\ell = \ell_1$.

If the pendulum P is set into vibration, this vibration reaches the other pendulums through the string. Then the pendulums A and B receive a driving force through the string and they also start vibrations and their amplitude increases due to the resonance phenomenon because their lengths, natural frequency and natural periods are same. At the same time, the pendulums C and D whose natural frequencies are different from natural frequency of P do not oscillate i.e., they continue to remain at rest. If the length of the pendulum P is made equal to ℓ_2 and allowed to vibrate, then the pendulums C and D start vibration due to resonance while pendulum A and B remain at rest.

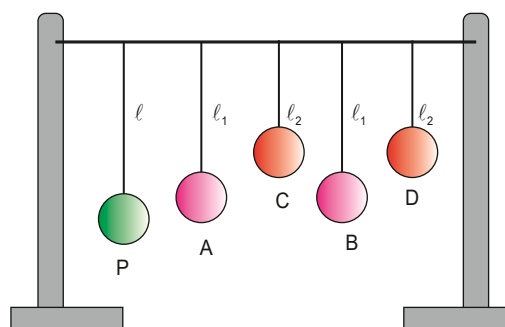


Fig 14.17: A set of five different pendulums of different lengths

Interesting Information

The resonance is the oscillation up and down back and forth motion caused by seismic waves. During an earthquake, buildings oscillate. An earthquake occurred in Pakistan magnitude 7.6 on Richter scale on 8 October, 2005 in Azad Jammu and Kashmir, a territory under Pakistan. It was centered near the city of Muzaffarabad. Over 86000 people died during this earthquake.

Some Circumstances in which Resonance is Useful

1. Heating / Cooking Food in Microwave Oven

A good example of resonance is the heating and cooking of food very efficiently and evenly by microwave oven. The waves produced in this type of oven have a wavelength of 12 cm at a frequency of 2450 MHz. At this frequency, the waves are absorbed due to resonance by water and fat molecules in the food, heating them up and so cooking the food. Food containing water molecules can only be heated by the microwave oven as shown in Fig. 14.18. The plastic or glass container do not heat up since they do not contain water molecules.



Fig. 14.18: Microwave oven

2. Radio Tuning

The tuning of an analogue radio set for a certain station is also based on resonance in which we need to turn the knob of a radio to tune a station. Here actually, we change the

natural frequency of the electrical circuit of receiver to make it equal to the transmission frequency of radio station. When two frequencies match with each other, then resonance occurs and energy absorbed by the station is maximum and this is the only station we hear (Fig. 14.19).

Following are some circumstances in which resonance should be avoided:

1. Bridge Vibration

The marching soldiers are advised to break steps while crossing a bridge. If the soldiers march in steps, then it is possible that the frequency of their footsteps becomes equal to the natural frequency of the bridge and the bridge may be set into vibrations with large amplitude due to the resonance and may collapse.

2. Aeroplane Wing

For equilibrium, the structure of an aeroplane has been designed with its two wings. These wings experience the forces such as; aerodynamics force, turbulence and engine vibration, etc. When the frequencies of these forces are matching with the natural frequency of the wings, the resonance will occur and it may be dangerous.



Fig. 14.19: Radio

Interesting Information

- The sound board of a piano, violin and guitar resonates with the vibration of string. When a musician strikes the string of a musical instrument such as a guitar or violin, the string produces sound waves in the form of vibrations. These vibrations transfer their energy to the hollow wooden box of instrument. The sound box resonates, amplifying the sound waves. The amplified sound waves are then emitted, producing a louder and richer sound.
- Magnetic resonance imaging (MRI) is an improved medical diagnostic technique in which strong radio frequency radiations are used to cause nuclei to oscillate, energy is absorbed by the nuclei. The patterns of the energy absorbed can be used to produce a computer enhanced photograph.
- The amplitude of a swing can be increased by applying a suitable periodic force on it regularly. It must be noted that natural frequency of swing matches the frequency of pushes, amplifying the motion.

14.11 SHARPNESS OF RESONANCE

We have seen that at resonance, the amplitude of the oscillator becomes very large. The amplitude as well as its sharpness, both depend upon the damping. Smaller the damping, greater will be the amplitude and more sharp will be the resonance.

A heavily damped system has a fairly flat resonance curve as is shown in an amplitude frequency graph (Fig. 14.20).

The effect of damping can be observed

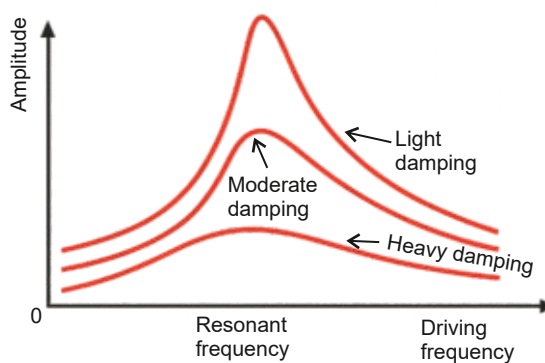


Fig. 14.20: Damping system illustration

by attaching a pendulum having light mass such as a pith ball as its bob and another of the same length carrying a heavy mass such as a lead bob of equal size, to a string (Fig. 14.17). They are set into vibrations by a third pendulum of equal length, attached to the same string. It is observed that amplitude of the lead bob is much greater than that of the pith-ball. The damping effect for the pith-ball due to air resistance is much greater than for the lead bob.

Application of Resonance and Standing Waves

There are many applications of resonance in different devices that generate and use standing waves.

There is relationship between standing wave and resonance phenomenon which can be studied by some experimental examples such as Rubens tube, chladni plates and acoustic levitations.

A Rubens tube is a long tube which is an experimental apparatus used to explain the acoustic standing wave. It consists of a metal pipe with holes drilled a long top and sealed at both ends. One sealed end is attached to a small speaker or frequency generator while the other end is connected to supply of a flammable gas as shown in Fig (14.21). The pipe is filled with gas and gas leaks from holes and produce a resonant flame.

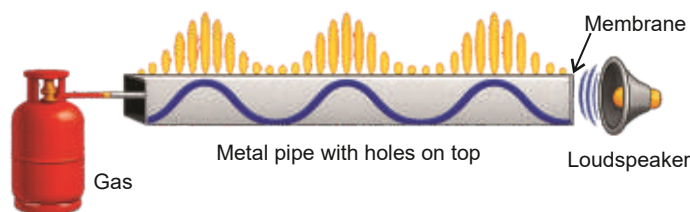
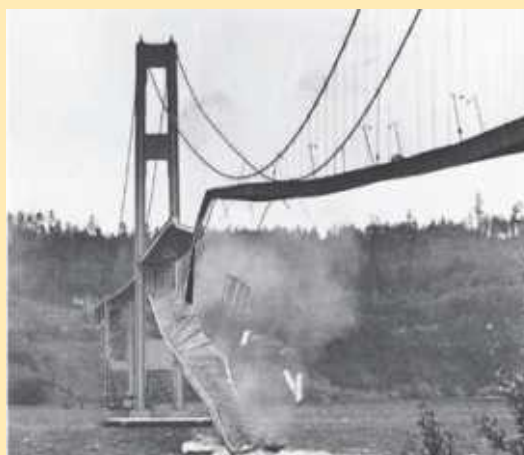


Fig. 14.21: Rubens tube apparatus

When sound waves of resonance frequency are produced by frequency generator, a standing wave is formed inside the tube. The standing wave creates points with

Point to ponder!



The resonance effect is very important in the design of bridges and other civil engineering projects. In July, 1940 the newly constructed Tacoma Narrows Bridge (Washington) was opened for traffic. Only four months after this, a mild wind set up the bridge in resonant vibrations. In a few hours the amplitude became so large that the bridge could not stand the stress and a part broke off and went into the water below.

Interesting Information



A Rubens tube also known as standing wave flame tube or simply flame tube is a physics apparatus for demonstrating acoustic standing waves in a tube invented by a German physicist **Heinrich Rubens in 1905.**

oscillating (higher and lower) pressure within the tube. Less gas will escape from the holes in the tube where pressure is low, hence, the flame will be lower at these points. Large quantity of gas will escape from holes in the tube where pressure is high, hence, flame will be high at these points.

These low and high flame reveal the nodes and antinodes of standing waves inside the tube as shown in Fig. (14.21).

Chladni Plate

A Chladni plate is a good example of resonance and it demonstrates a standing waves pattern. It is usually a circular or square shaped metal plate which is fixed horizontally to a vibrator and is driven transversally at its centre as shown in Fig. (14.22). When the plate is vibrated at the frequency of a resonant normal mode, the nodes and antinodes are formed over the surface of the plate.

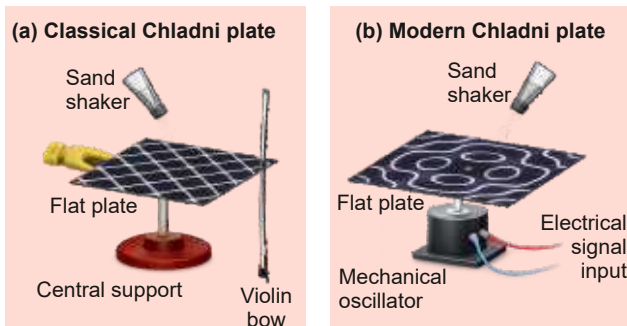


Fig. 14.22: Chladni plates

The metal plate does not vibrate at points where nodes are formed while the rest of the plate is vibrating with its maximum amplitude. In order to observe such pattern, sand is sprinkled evenly in the form of a thin layer over the plate. Hence, the sand will vibrate and shift from antinode to nodes because the nodes are stationary.

Acoustic Levitation

Acoustic levitation is a method for suspending matter in air against gravity using radiation pressure from high intensity sound waves. This method explains the relations between resonance and standing waves. In which the high frequency sound waves (ultrasonic) are used to suspend the small particles such as bits of styrofoam or droplets of water against gravity in air without any physical contact.

Working: Basically, an acoustic levitator consists of a source of sound (transducer) and a reflector. The high frequency sound waves from the transducer is bounced off by the reflector. The interaction between the incident and reflected waves forms a standing wave, where the pressure of sound in the area of node is minimum and pressure is maximum in the area of antinode. Such difference between higher and lower pressure exerts a force on small particles of

Interesting Information



Chladni plates, invented by the physicist, musician and musical instrument maker **Ernst Chladni (1756-1827)** in the late 18th century are used to demonstrate the complex pattern of standing waves vibrations.

Styrofoam against the gravity. As a result, the particles are trapped at node as shown in Fig. (14.23). In this case, the nodes of the standing waves act as acoustic traps.

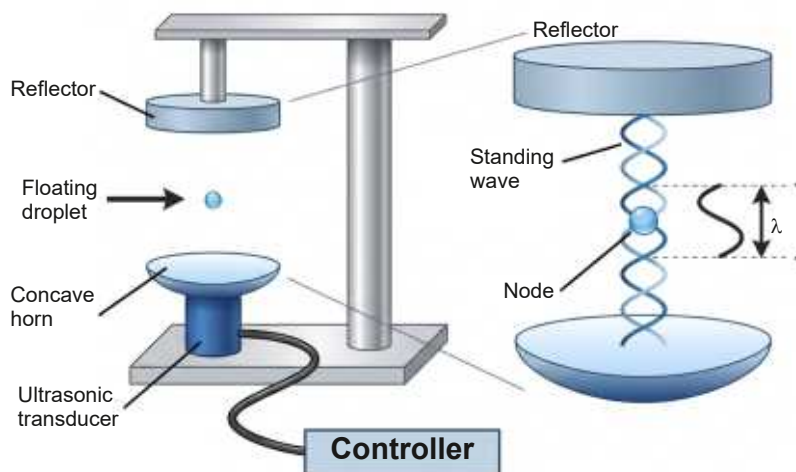


Fig. 14.23: Acoustic levitation setup

QUESTIONS

Multiple Choice Questions

Choose the correct answer.

14.1 While determining time period of simple pendulum, we keep the amplitude:

- (a) large (b) small (c) maximum (d) zero

14.2 In simple harmonic motion, the acceleration of a body is zero, when:

- (a) velocity is zero
 (b) displacement is zero
 (c) both velocity and displacement are zero
 (d) both velocity and displacement are maximum

14.3 If a tunnel is bored through the centre of Earth and a stone is dropped into it, then:

- (a) stone will stop at the centre of Earth
 (b) stone will move out from other side of the tunnel
 (c) stone will perform S.H.M.
 (d) none of the above

14.4 Which of the following variables has zero value at the extreme position in S.H.M?

- (a) Acceleration (b) Speed
 (c) Displacement (d) Angular frequency

14.5 Which of the following force is responsible for S.H.M?

- (a) Applied force (b) Frictional force
 (c) Restoring force (d) Elastic force

- 14.6 The velocity of a particle moving with S.H.M at the mean position is:
(a) zero (b) minimum (c) maximum (d) none of these
- 14.7 A child swinging on a swing in sitting position, stands up, then time period of the swing will:
(a) increase (b) decrease
(c) remain same (d) become zero
- 14.8 The total energy of a particle executing S.H.M is proportional to:
(a) frequency of oscillation
(b) velocity of motion
(c) amplitude of motion
(d) square of amplitude of motion
- 14.9 The period of simple pendulum increases how many times if its length is doubled?
(a) $\sqrt{2}$ times (b) 2 times
(c) 4 times (d) 8 times

Short Answer Questions

- 14.1 Define simple harmonic motion and write its equation.
- 14.2 State the basic conditions for frictionless system to execute simple harmonic motion.
- 14.3 Differentiate between free and forced oscillations.
- 14.4 Why we cannot construct an ideal simple pendulum? Explain briefly.
- 14.5 How does sharpness of resonance occur?
- 14.6 What would happen to the time period if the Earth suddenly stop rotating?
- 14.7 What is restoring force? What provides restoring force for simple harmonic oscillator in the case of simple pendulum and mass-spring system?
- 14.8 What role do shock absorbers play in oscillations of vehicles in a bumpy road? Explain briefly.

Constructed Response Questions

- 14.1 Why marching troops are asked to break their steps while crossing a bridge? Explain.
- 14.2 A girl is swinging in the sitting position. Describe, how will the period of swing be changed if she stands up?
- 14.3 Will a pendulum that keeps correct time at Karachi, be accurate at Murree or at the top of mount Everest? Explain.
- 14.4 A wire hangs from a dark high tower so that its upper end is not visible. How can we determine the length of a tower?
- 14.5 Will the period of a vibrating spring increase, decrease or remain constant by addition of more mass? Explain.

14.6 Can a simple pendulum vibrate at centre of the Earth? Explain.

Comprehensive Questions

- 14.1 Show that motion of mass attached with spring executes S.H.M. Also find an expression for its time period and frequency.
- 14.2 Show that motion of the projection of a particle moving around a circle with constant speed is S.H.M and derive a relation for its instantaneous velocity.
- 14.3 What is a simple pendulum? Show that motion of a simple pendulum is S.H.M. Also derive an expression for time period of a simple pendulum, on what factors does it depend?
- 14.4 Derive expressions for *K.E.* and *P.E.* of the mass-spring system executing S.H.M. Also prove that its total energy remains same.
- 14.5 Explain resonance. Give some of its advantages and disadvantages.
- 14.6 Discuss physical significance of phase in simple harmonic motion, and explain how does it relate to timing and synchronization of oscillation?

Numerical Problems

- 14.1 A body is executing S.H.M of amplitude 0.05 m and frequency 10 Hz. Find the value of (i) maximum acceleration (ii) maximum velocity
(Ans: $1.97 \times 10^2 \text{ m s}^{-2}$, 3.14 m s^{-1})
- 14.2 A mass of 0.5 kg is suspended from a spring. The spring is stretched 0.098 m. Find:
(i) the period of oscillation of mass when it is given a small displacement.
(ii) frequency of oscillation of spring. (Ans: 0.628 s, 1.59 Hz)
- 14.3 A ball of mass 0.5 kg is attached to an elastic spring. The spring is compressed by 10 cm by a force of 45 N. Calculate the energy stored in the spring. What will be the maximum velocity of the ball if the compressing force is suddenly removed?
(Ans: 2.25 J, 3 m s^{-1})
- 14.4 A 8.0 kg body executes S.H.M with amplitude 30 cm. The restoring force is 60 N, when the displacement is 30 cm. Find:
(i) time period (ii) acceleration, speed, *K.E.* and *P.E.* when displacement is 12 cm.
(Ans: 1.3 s, 3.0 m s^{-2} , 1.4 m s^{-1} , 7.6 J, 1.44 J)
- 14.5 A block of mass 5 kg is dropped from a height of 0.8 m onto a spring of spring constant 1960 N m^{-1} . Find the displacement through which the spring will be compressed.
(Ans: 0.2 m)

Physical Optics

Students Learning Outcomes

After studying this chapter, the students will be able to:

- ◇ explain experiments that demonstrate two-source interference for sound, light and microwaves.
- ◇ describe the conditions required if two-source interference fringes are to be observed.
- ◇ use $\Delta y = L\lambda / d$ for double-slit interference using light to solve problems.
- ◇ use $d \sin \theta = n\lambda$ to solve problems.
- ◇ describe the use of a diffraction grating to determine the wavelength of light [the structure and use of the spectrometer are not included].

It has been observed that light has been a dilemma since centuries. Basically, light is a type of energy which produce sensation of vision. Now, the question is how light transmits energy from one point to another as a wave or stream of particles.

In about 1678, Huygen, an eminent Dutch scientist, proposed that light energy from a luminous source travels in space by means of wave motion. The experimental evidence in support of wave theory in Huygen's time was not convincing. However, with the discovery of interference of light, Young's interference experiment performed in 1801 proved wave nature of light and thus established the Huygen's wave theory. In 1860, Maxwell proposed that light consists of electromagnetic waves. According to him, the oscillating electric charges can produce changing electric and magnetic fields that mutually regenerate each other and can propagate through vacuum. In early twentieth century, the quantum mechanical explanation revealed the dual nature of light, both wave and particle characteristics, depending on the experiment. It propagates as an electromagnetic wave but interacts with matter as discrete, massless particles called photons. However, the wave theory of light is still valid and explains phenomena of light successfully.

In this chapter, we will study the properties of light, associated with its wave nature such as interference and diffraction.

15.1 WAVEFRONT

To understand how energy is spread out in the form of waves, we need to understand the concept of wavefront.

A wavefront is an imaginary surface that connects all the points where the waves have the same phase of vibration.

Wavefronts can be of three types:

1. **Plane wavefronts:** Represented as a straight line like light waves from the distant stars and water waves in a ripple tank.
2. **Circular wavefronts:** Represented by circles as formed in water inside a pond.
3. **Spherical wavefronts:** Formed in three-dimensional space by a sound or light wave.

For a point source in space, the wavefronts will always be spherical i.e., they spread out in all directions. For very large distances, like light rays coming from the Sun, these wavefronts can be assumed as planes. The shortest distance between wavefronts is always equal to wavelength of the wave. Direction of transfer of energy of wave is always perpendicular to any point on the surface of a given wavefront.

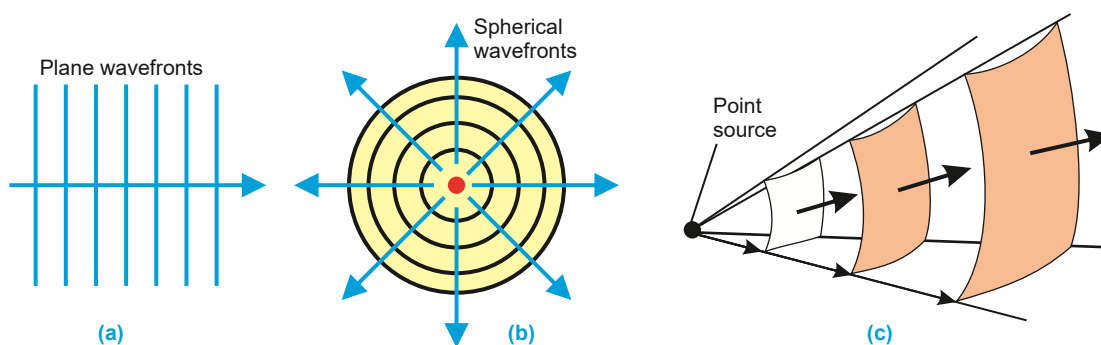


Fig.15.1(a,b,c): Plane, circular and spherical wavefronts

If the source is periodic, it produces a succession of wavefronts, all of the same shape. Plane wavefronts are produced by a plane source or by any source at distant point. A line source for example a vibrator in a ripple tank creates wavefronts that are plane and straight in two dimensions. A line at right angle to a wavefront which shows its direction of travel is called a ray, as shown in Fig. 15.1 (a). Circular wavefronts in two dimensions and spherical wavefronts in three dimensions from any point source are shown in Figs. 15.1 (b and c).

Here in this chapter, we will assume plane wavefronts from different sources to study interference and diffraction phenomenon. If a source of light is placed at the focal point of a lens, plane wavefronts are formed on the other side of the lens.

15.2 HUYGEN'S PRINCIPLE

Dutch Scientist Christiaan Huygens, used geometrical method for finding, from the known shape of a wavefront at some instant, the shape of the wavefront at some later time. Huygens assumed that:

Every point of a wavefront may be considered the source of secondary wavelets that spread out in all directions with a speed equal to the speed of propagation of the wave.

The new wavefront at a later time is then found by constructing a surface tangent to the secondary wavelets or, as it is called, the envelope of the wavelets. Figure 15.2

illustrates Huygens's principle. The original wavefront AB is travelling outward from a source, as indicated by the arrows. We want to find the shape of the wavefront after a time interval t . We assume that v , the speed of propagation of the wave, is the same at all points. Then, in time t the wavefront travels a distance vt . We construct several circles (traces of spherical wavelets) with radius vt , centered at points along AB. The trace of the envelope of these wavelets, which is the new wavefront, is the curve CD.

All the results that we obtain from Huygens' principle can also be obtained from Maxwell's equations, but Huygens's simple model is easier to use.

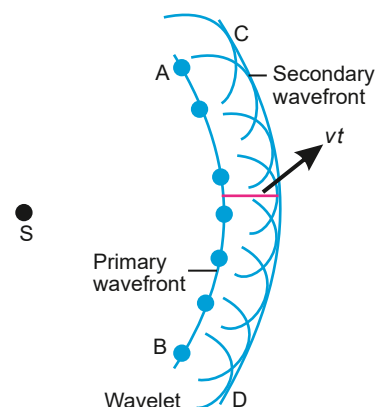


Fig. 15.2: Formation of new wavefront from a previous wavefront

15.3 INTERFERENCE

Adding waves of different wavelengths and amplitudes results in complex waves. We can find some interesting effects if we consider what happens when two waves of the same wavelength overlap at a point. Two coherent waves, when meet in a region, they reinforce each other at some points and cancel out each other at some other points forming a pattern called interference.

Consider two loudspeakers producing sounds of same frequency and wavelength. Figure 15.3 shows how interference arises. The loudspeakers are emitting waves that are in phase because both are connected to the same source. At each point in front of the loudspeakers, waves are arriving from the two loudspeakers. At some points, the two waves arrive in phase with one another and with equal amplitude (Fig. 15.3-a). The principle of superposition predicts that the resultant wave has twice the amplitude of a single wave. We hear a louder sound.

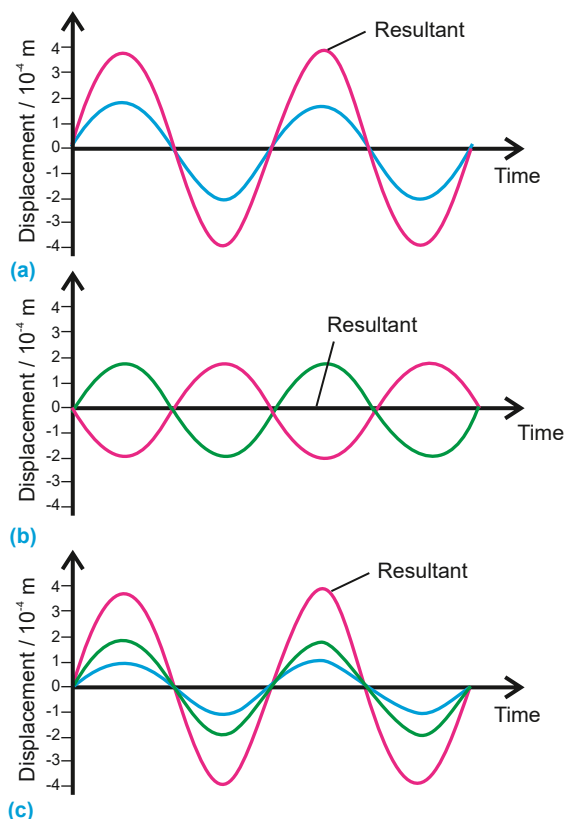


Fig. 15.3: Adding waves by the principle of superposition. Blue and green waves of the same amplitude may give a constructive interference (b) destructive interference, according to the phase difference between them. (c) Waves of different amplitudes can also interfere constructively.

The basic conditions for two waves to produce interference are:

- (i) waves must be of same type i.e. sound waves can only produce interference with other sound waves.
- (ii) waves must be coherent having a constant phase difference.
- (iii) waves must be moving in the same direction.

Coherence

We are surrounded by many types of waves, for example; visible light, infrared radiations, radiowaves and sound waves. There are waves coming at us from all directions. So, why do we not observe interference patterns all the time? Why do we need specialised equipment in a laboratory to observe these effects?

In fact, we can see interference of light occurring in everyday life. For example, you may have noticed haloes of light around street lamps or the moon on a foggy night.

You may have noticed bright and dark bands of light if you look through fabric at a bright source of light. These are interference effects.

We usually need specially arranged conditions to produce interference effects that we can measure. Think about the demonstration with two loudspeakers. If they were connected to different signal generators with slightly different frequencies, the sound waves might start off in phase with one another, but they would soon go out of phase (Fig.15.4). We would hear loud, then soft, then loud again. The interference pattern would keep shifting around the room.

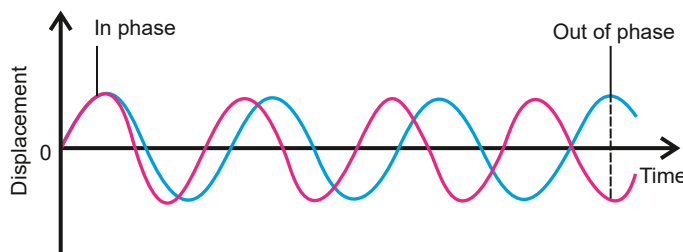


Fig.15.4: Waves have slightly different wavelengths, therefore, they move in, and out of phase with one another.

By connecting the two loudspeakers to the same signal generator, we can be sure that the sound waves that they produce are constantly in phase with one another. We say that they act as two coherent sources of sound waves. Coherent sources emit waves that have a constant phase difference. Note that the two waves can only have a constant phase difference if their frequencies are the same and remains constant.

Now think about the laser experiment. Could we have used two lasers producing exactly the same frequency and hence wavelength of light? Figure 15.5 (a)

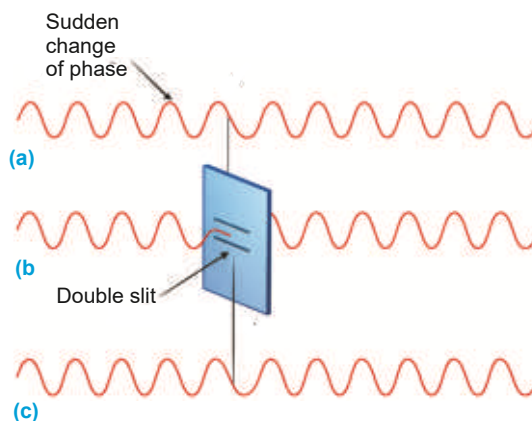


Fig.15.5: Waves must be coherent if they are to produce a clear interference pattern.

represents the light from a laser. We can think of it as being made up of many separate bursts of light. We cannot guarantee that these bursts from two lasers will always be in phase with one another.

This problem is overcome by using a single laser and dividing its light using the two slits (Fig. 15.5-b). The slits act as two coherent sources of light. They are constantly in phase with one another (or there is a constant phase difference between them).

If they were not coherent sources, the interference pattern would be constantly changing, far too fast for our eyes to detect. We would simply see a uniform band of light, without any definite bright and dark regions. From this, we should be able to see that, in order to observe interference, we need two coherent sources of waves.

15.4 INTERFERENCE OF MICROWAVES

Using 2.8 cm wavelength microwave equipment (Fig. 15.6), we can observe an interference pattern. The microwave transmitter is directed towards the double gap in a metal barrier. The microwaves are diffracted at the two gaps so that they spread out into the region beyond, where they can be detected using the probe receiver. By moving the probe around, it is possible to detect regions of high intensity (constructive interference) and low intensity (destructive interference). The probe may be connected to a meter, or to an audio amplifier and loudspeaker to give an audible output. Because microwaves have wavelengths comparable to the metal barrier (apparatus) size, the pattern is easily observable and measurable. The path difference is given by

$$\Delta x = d \sin \theta$$

i.e., For Maxima (bright region) $\Delta x = m\lambda$

For Minima (dark region) $\Delta x = (m + 1/2)\lambda$

Interference of Light Waves

Here is one way to show the interference effects produced by light. A simple arrangement involves directing the light from a laser through two slits (Fig. 15.7). The slits are two clear lines on a black slide, separated by a few millimetres. Where the light falls on the screen, a series of equally spaced dots of light are seen. These bright dots are referred to as interference 'fringes', and they are regions where light

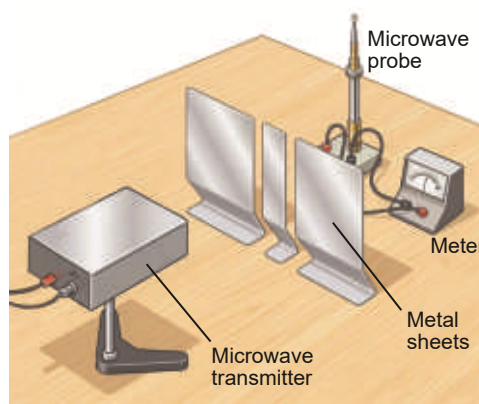


Fig. 15.6: Microwaves can also be used to show interference effects

For your information

A microwave oven has a metal grid in the door to keep microwaves in and let light out.

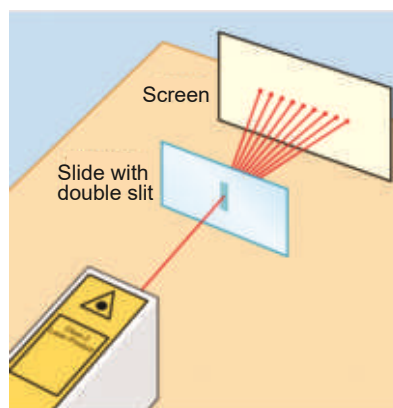


Fig. 15.7: Interference of light beams from two slits.

waves from the two slits are arriving in phase with each other, i.e., there is constructive interference. The dark regions in between are the result of destructive interference.

These bright and dark fringes are the equivalent of the loud and quiet regions that we detected if we investigated the interference pattern of sounds from the two loudspeakers. Bright fringes correspond to loud sound, dark fringes to quiet sound or silence. We can check that light is indeed reaching the screen from both slits as follows: Mark a point on the screen where there is a dark fringe. Now carefully cover up one of the slits so that light from the laser is only passing through one slit. We should find that the pattern of interference fringes disappears. Instead, a broad band of light appears across the screen. This broad band of light is the diffraction pattern produced by a single slit. The point that was dark is now bright. Cover up the other slit instead, and we shall see the same effect. We have now shown that light is arriving at the screen from both slits, but at some points (the dark fringes), the two beams of light cancel each other out.

15.5 YOUNG'S DOUBLE-SLIT EXPERIMENT

We will take a close look at a famous experiment which Thomas Young performed in 1801. He used this experiment to show the wave nature of light. A beam of light is shown on a pair of parallel slits placed at right angles to the beam. Light diffracts and spreads outwards from each slit into the space beyond; the light from the two slits overlaps on a screen. An interference pattern of bright and dark bands called 'fringes' is formed on the screen.

In order to observe interference, we need two sets of waves. The sources of the waves must be coherent, the phase difference between the waves emitted at the sources must remain constant. This also means that the waves must have the same wavelength. Today, this is readily achieved by passing a single beam of laser light through the two slits. A laser produces intense coherent light. As the light passes through the slits, it is diffracted so that it spreads out into the space beyond (Fig.15.8). Now we have two overlapping sets of waves, and the pattern of fringes on the screen shows us the result of their interference. How does this pattern arise? We will consider three points on the screen (Fig.15.9-a,b,c), and work out what we would expect to observe at each point.

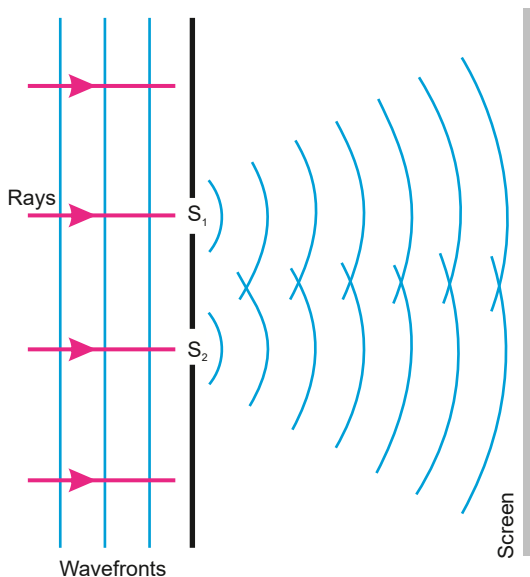


Fig.15.8: Double-slit arrangement

Point A

This point is directly opposite the mid-point of the double slit. Two rays of light arrive at A, one from slit S_1 and the other from slit S_2 . Point A is equidistant from the two slits, and so

the two rays of light have travelled the same distance. The path difference between the two rays of light is zero. If we assume that they were in phase with each other when they left the slits, then they will be in phase when they arrive at A. Hence, they will interfere constructively, and we will observe a bright fringe at point A.

Point B

This point is slightly to the side of point A, and is the mid-point of the first dark fringe. Again, two rays of light arrive at B, one from each slit. The light from slit 1 has to travel slightly further than the light from slit 2, and so the two rays are no longer in phase. Since point B is at the mid-point of the dark fringe, the two rays must be in antiphase (phase difference of 180°). The path difference between the two rays of light must be half a wavelength and so the two rays interfere destructively.

Point C

This point is the mid-point of the next bright fringe. Again, ray from S_1 has travelled further than ray from S_2 ; this time, it has travelled an extra distance equal to a whole wavelength λ . The path difference between the rays of light is now a whole wavelength. The two rays are in phase at the screen. They interfere constructively and we see a bright fringe. The complete interference pattern (Fig. 15.9) can be explained in this way.

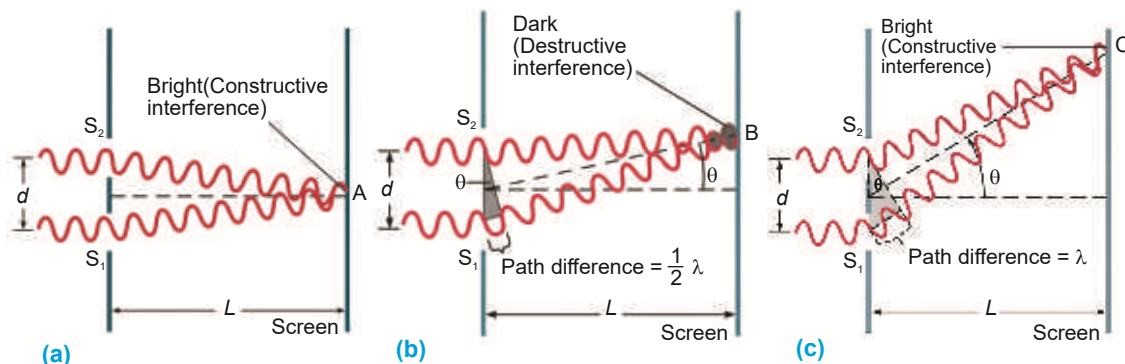


Fig. 15.9:

(a) for $\theta = 0^\circ$ Central bright region. (b) For first order dark region (c) For first order bright region

To simplify the analysis of Young's experiment, we assume that the distance L from the slits to the screen is so large in comparison to the distance d between the slits that the lines from S_1 and S_2 , as in Fig. 15.9(c). This is usually the case for experiments with light; the slit separation is typically a few millimetres, while the screen may be a metre or more away. The path difference is then given by

$$\text{Path difference} = d \sin \theta \quad \dots\dots\dots (15.1)$$

where θ is the angle between a line from slits to screen and the normal to the plane of the slits.

Constructive and Destructive Interference

We have learnt that constructive interference (reinforcement) occurs at points where the path difference is an integer number of wavelengths, $m\lambda$, where $m = 0, 1, 2, 3, \dots$. So the bright regions on the screen in Fig. 15.10 occur at angles θ for which:

$$d \sin \theta = m\lambda \quad \text{where } m = 0, 1, 2, 3, \dots \quad (15.2)$$

Similarly, destructive interference (cancellation) occurs, forming dark regions on the screen, at points for which the path difference is a half-integer number $(m + \frac{1}{2})$ of wavelengths.

$$d \sin \theta = (m + \frac{1}{2}) \lambda \quad \text{where } (m = 0, 1, 2, 3, \dots)$$

Thus, the pattern on the screen of Fig. 15.10 is a succession of bright and dark bands, or interference fringes, parallel to the slits S_1 and S_2 . A photograph of such a pattern is shown in Fig. 15.10.

We can derive an expression for the positions of the centres of the bright bands on the screen. In Fig. 15.10, y is measured from the centre of the pattern, corresponding to the distance from the centre of the pattern. Let y_m be the distance from the centre of the pattern ($\theta = 0$) to the centre of the m th bright fringe. Let θ be the corresponding value of θ ; then

$$y_m = L \tan \theta \quad (15.3)$$

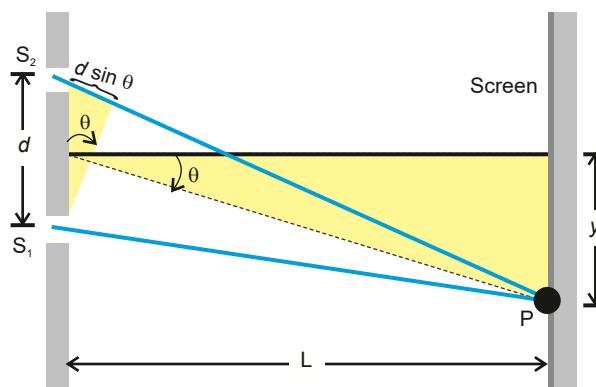


Fig.15.10: Schematic double slit arrangement

In such experiments, the distances y_m are often much smaller than the distance L from the slits to the screen. Hence θ is very small, so, $\sin \theta = \tan \theta$ and:

$$y_m = L \sin \theta \quad (15.4)$$

Hence, for small angles, positions of m th bright fringe is given by

$$y_m = \frac{mL\lambda}{d} \quad (15.5)$$

m th dark fringe is given by; $y_m = (m + \frac{1}{2}) \frac{L\lambda}{d}$ for the first order dark fringe $m = 0$.

The distance between two consecutive bright fringes is called fringe width or fringe spacing and can be calculated as:

$$\text{Let } m\text{th fringe } y_m = \frac{mL\lambda}{d} \quad \text{and} \quad (m+1)\text{th fringe is } y_{m+1} = \frac{(m+1)L\lambda}{d}$$

$$\text{Hence, fringe spacing } \Delta y = y_{m+1} - y_m \quad \text{or} \quad \Delta y = \frac{(m+1)L\lambda}{d} - \frac{mL\lambda}{d}$$

$$\Delta y = \frac{L\lambda}{d} \quad (15.6)$$

The distance between adjacent bright bands in the pattern is inversely proportional to the separation d between the slits. The closer together the slits are, the more the pattern spreads out. When the slits are far apart, the bands in the pattern are closer together.

While we have described the experiment that Young performed with visible light, the results given in the above equations are valid for any type of wave, provided that the resultant wave from two coherent sources is detected at a point that is far away in comparison to the separation d .

Example 15.1 The distance between two consecutive slits in Young's experiment is 0.25 cm. Screen is at a distance of 100 cm from slits. Calculate distance of third order dark fringe from centre of screen, if wavelength of light used is $\lambda = 519 \text{ nm}$.

Solution $d = 0.25 \text{ cm} = 0.25 \times 10^{-2} \text{ m}$, $L = 100 \text{ cm} = 1.0 \text{ m}$
 $\lambda = 519 \text{ nm} = 519 \times 10^{-9} \text{ m}$
 For third dark fringe $m = 2$

$$y = \left(m + \frac{1}{2}\right) \frac{L\lambda}{d}$$

$$\text{So } y = \left(2 + \frac{1}{2}\right) \frac{1.0 \text{ m} \times 519 \times 10^{-9} \text{ m}}{0.25 \times 10^{-2} \text{ m}} = 5.9 \times 10^{-4} \text{ m}$$

15.6 INTERFERENCE IN THIN FILMS

A film whose thickness is of the order of wavelength of light is called thin film.

We often see bright bands of colour when light reflects from a thin layer of oil floating on water or from a soap bubble. These are the results of interference. Light waves are reflected from the front and back surfaces of such thin films, and constructive interference between the two reflected waves (with different path lengths) occurs in different places for different wavelengths. Fig. 15.11 shows the situation. Light shining on the upper surface of a thin film with thickness t is partly reflected at the upper surface (path abc) of thin film. Light transmitted through the upper surface is partly reflected at the lower surface (path $abdef$). The two reflected waves come together at point P on the retina of the eye. Depending on the phase relationship, they may interfere constructively or destructively. Different colours have different wavelengths, so the interference may be constructive for some colours and destructive for others.

That is why we see coloured patterns in thin film as shown in Fig. 15.12 (which shows thin films of soap solution that makes up the bubble walls). In similar fashion, thin film of floating on water shows same pattern. The complex shapes of the coloured patterns result from variations in the thickness of the film.

Let us look at a simplified situation in which monochromatic light reflects from two nearly parallel surfaces at nearly normal incidence. Figure 15.13 shows two plates of glass

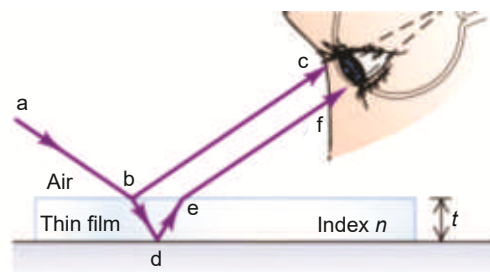


Fig.15.11: Interference in thin film



Fig.15.12: White light falling on soap bubble

separated by a thin wedge, or film of air. We want to consider interference between the two light waves reflected from the surfaces adjacent to the air wedge. Here reflections also occur at the top surface of the upper plate and the bottom surface of the lower plate; to keep our discussion simple, we would not include these.) The situation is the same as in Fig. 15.13 except that the thin film thickness is not uniform. The path difference between the two waves is just twice the thickness t of the thin air film at each point. At points where $2t$ is an integer number of wave lengths, we expect to see constructive interference and a bright area; where it is a half-integer number of wavelengths, we expect to see destructive interference and a dark area. Where the plates are in contact, there is practically no path difference, and we expect a bright area.

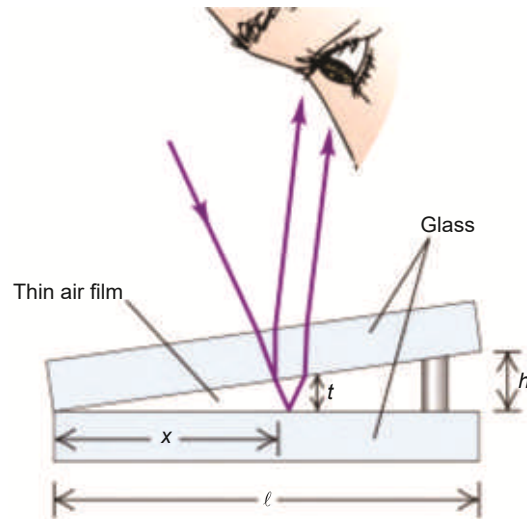


Fig.15.13: Thin air film between two glass plates

When we carry out the experiment, the bright and dark fringes appear, but they are interchanged. Along the line where the plates are in contact, we find a dark fringe, not a bright one. This suggests that one or the other of the reflected waves has undergone a half cycle phase shift during its reflection. In that case, the two waves that are reflected at the line of contact are a half-cycle out of phase even though they have the same path length. In fact, this phase shift can be predicted from Maxwell's equations and the electromagnetic nature of light.

15.7 NEWTON'S RINGS

Figure 15.14 shows the plano-convex lens in contact with a plane glass plate. A thin film of air is formed between the two surfaces.

When we view the setup in which thin air film is enclosed between plano-convex lens and a glass plate with monochromatic light, alternate bright and dark circular interference fringes are formed. These were studied by Newton and are called Newton's rings.

Newton's rings are formed when a plano-convex lens is placed on a flat glass plate, a thin air film forms between them. Light from a monochromatic source reflects off both the top and bottom surfaces of this air film. The reflected rays interfere constructively (bright rings) or destructively (dark rings) depending on the path difference and phase shift between them. Newton's rings can be used to calculate the wavelength of monochromatic light and to determine the radius of curvature of lenses and is used in precision optics to detect surface irregularities. Alternate bright and dark circular

interference fringes are shown in the Fig. 15.15.

Do you know?

Different colours seen on the surface of soap bubble is due to interference of white light falling and reflecting from it.

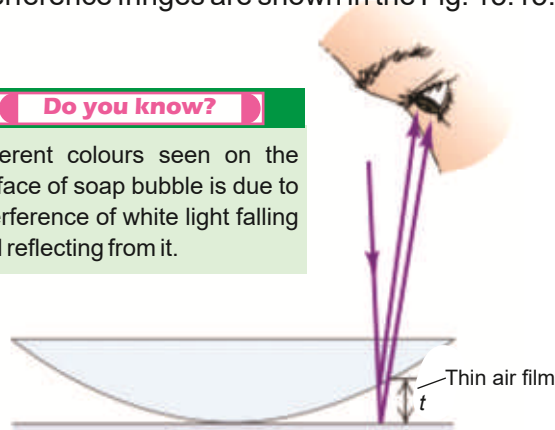


Fig.15.14: Plano-convex lens in contact with glass plate

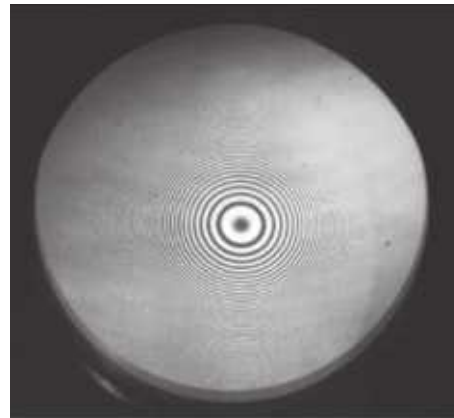


Fig.15.15: Circular interference fringes

15.8 DIFFRACTION GRATING

If we increase the number of slits in an interference experiment (while keeping the spacing of adjacent slits constant) it gives interference patterns in which the maxima are in the same positions, but progressively narrower, than with two slits. Because these maxima are so narrow, their angular position, and hence the wavelength, can be measured to very high precision. This effect has many important applications.

An array of a large number of parallel slits, all with the same width d and spaced equal distances d between centres, is called a **diffraction grating**. The first one was constructed by Fraunhofer using fine wires (Fig. 15.16). Gratings can be made by using a diamond point to scratch many equally spaced grooves on a glass or metal surface, or by photographic reduction of a pattern of black and white stripes on paper. For a grating, what we have been calling slits are often called rulings or lines.

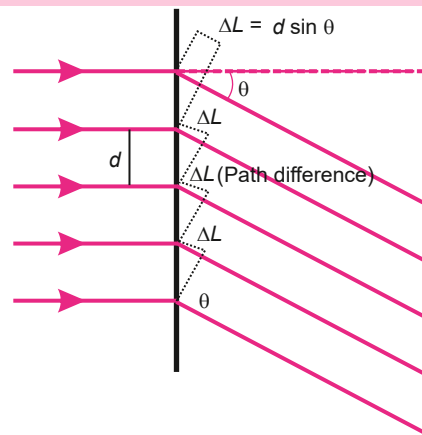


Fig.15.16: Diffraction grating

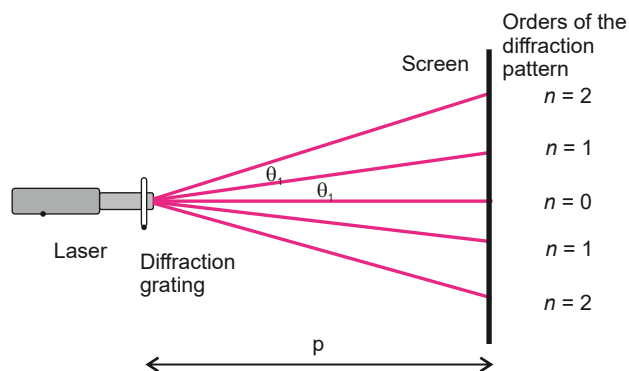


Fig.15.17: Orders of pattern

The property of bending of light around obstacles and spreading of light waves into the geometrical shadow of an obstacle is called diffraction.

Figure 15.17 is a cross section of a transmission grating; the slits are perpendicular to the plane of the page, and an interference pattern is formed by the light that is transmitted through the slits. The diagram shows only five slits; an actual grating may contain several thousands. The spacing d between centres of adjacent slits is called the grating spacing. A plane monochromatic wave is incident normally on the grating from the left side. We assume far-field (Fraunhofer) conditions; that is, the pattern is formed on a screen that is far enough away that all rays emerging from the grating and going to a particular point on the screen can be considered to be parallel.

The principal intensity maxima with multiple slits occur in the same directions as for the two-slit pattern. These are the directions for which the path difference for adjacent slits is an integer number of wavelengths. So, the positions of the maxima are once again:

$$d \sin \theta = n\lambda \quad \text{..... (15.7)}$$

where $n = 0, \pm 1, \pm 2, \pm 3, \text{.....}$

Determination of the wavelength using diffraction grating

To measure the wavelength of a beam of light if we know the d (i.e., number of lines per mm or per cm) e.g., if 600 lines per mm, then

$$d = \frac{1}{600} \text{ mm} = 1.67 \times 10^{-6} \text{ m}$$

Illuminate the grating with the monochromatic light or spectral line. If θ is the diffraction angle for a particular order " n " (usually first or second order) then

$$\lambda = \frac{d \sin \theta}{n}$$

When a grating containing hundreds or thousands of slits is illuminated by a beam of parallel rays of monochromatic light, the pattern is a series of very sharp lines at angles determined by this formula. The $n = \pm 1$ lines are called the first-order lines, the $n = \pm 2$ lines, the second-order lines, and so on. If the grating is illuminated by white light with a continuous distribution of wavelengths, each value of n corresponds to a continuous spectrum in the pattern. The angle for each wavelength is determined by Eq. (15.7) for a given value of n , long wavelengths (the red end of the spectrum) lie at larger angles (that is, are deviated more from the straight-ahead direction) than do the shorter wavelengths at small angles (blue end of the spectrum).

Gratings for use with visible light of wavelength from 400 to 700 nm, usually have about 1000 slits per millimetre; the value of d is the reciprocal of the number of slits per unit length, so d is of the order of 1000 nm.

15.9 DIFFRACTION OF X-RAY BY CRYSTALS AND BRAGG'S LAW

X-ray is a type of electromagnetic radiation of much shorter wavelength, about 10^{-10} m. In order to observe the effects of diffraction, the grating spacing must be of the order of the wavelength of the radiation used. The regular array of atoms in a crystal forms a natural diffraction grating with spacing that is typically $\approx 10^{-10}$ m. The scattering of X-rays from the

atoms in a crystalline lattice gives rise to diffraction effects very similar to those observed with visible light incident on ordinary grating.

The study of atomic structure of crystals by X-ray was initiated in 1914 by W.H. Bragg and W.L Bragg with remarkable achievements. They found that a monochromatic beam of X-ray was reflected from a crystal plane as if it acted like mirror. To understand this effect, a series of atomic planes of constant interplanar spacing d parallel to a crystal face are shown by lines PP' , $P_1P'_1$, $P_2P'_2$, and so on (Fig. 15.18).

Suppose an X-rays beam is incident at an angle θ on one of the planes. The beam can be reflected from both the upper and the lower planes of atoms. The beam reflected from lower plane travels some extra distance as compared to the beam reflected from the upper plane. The effective path difference between the two reflected beams is $2d \sin \theta$. Therefore, for reinforcement, the path difference should be an integral multiple of the wavelength. Thus,

From figure 15.18:

$$\text{Path difference} = BC + B'C$$

$$\text{In triangle } ABC, \frac{BC}{AC} = \sin \theta$$

Assuming distance d between planes corresponds to a side related to AC , so that

$$BC = d \sin \theta.$$

$$\text{In triangle } AB'C, \frac{B'C}{AC} = \sin \theta$$

Therefore, similarly, $B'C = d \sin \theta$

Hence,

$$\begin{aligned} \text{Total path difference} &= BC + B'C \\ &= d \sin \theta + d \sin \theta \end{aligned}$$

$$\text{Total path difference} = 2d \sin \theta$$

$$2d \sin \theta = n\lambda \dots\dots\dots (15.8)$$

The value of n is referred to as the order of reflection. The Eq.15.8 is known as the Bragg equation. It can be used to determine interplanar spacing between similar parallel planes of a crystal if X-rays of known wavelength are allowed to diffract from the crystal.

X-rays diffraction has been very useful in determining the structure of biologically important molecules such as hemoglobin, which is an important constituent of blood, and double helix structure of DNA.

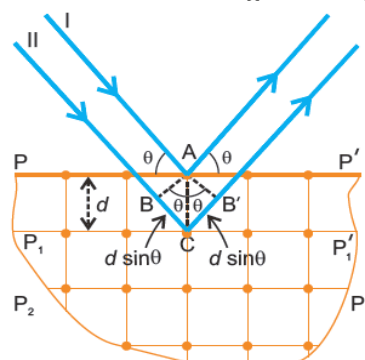


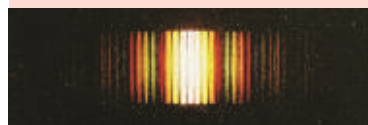
Fig.15.18: Diffraction of X-rays by crystals

Interesting Application

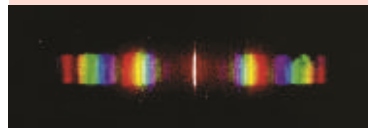


Diffraction of radiowaves

Interesting Information



The spectrum of white light due to diffraction grating of 100 slits.



The spectrum of white light due to diffraction grating of 2000 slits.

Example 15.2 For a NaCl crystal, the interplanar spacing for (1 00) planes is $d = 0.282$ nm. Copper K_α X-rays have $\lambda = 0.154$ nm. Find the Bragg's angle θ for first-order ($n = 1$) reflection.

Solution $2d\sin\theta = n\lambda \Rightarrow \sin\theta = \frac{n\lambda}{2d}$

$$= \frac{1 \times 0.154 \times 10^{-9} \text{ m}}{2 \times 0.282 \times 10^{-9} \text{ m}} = 0.273$$

$$\theta = \sin^{-1}(0.273) = 15.9^\circ$$

Example 15.3 For the NaCl at what angle would the second-order ($n = 2$) reflection occur?

Solution $n = 2, \theta = ?, \lambda = 0.154 \times 10^{-9} \text{ m}, d = 0.282 \times 10^{-9} \text{ m}$

$$\sin\theta = \frac{2 \times 0.154 \times 10^{-9} \text{ m}}{2 \times 0.282 \times 10^{-9} \text{ m}} = 0.546 \Rightarrow \theta \approx 33.1^\circ$$

Example 15.4 In an X-ray diffraction experiment using $\lambda = 0.071$ nm ($\text{Mo } K_\alpha$), a first-order reflection from a crystal is observed at $\theta = 12.0^\circ$. Calculate the interplanar spacing d .

Solution $n = 1, \lambda = 0.071 \times 10^{-9} \text{ m}, \theta = 12.0^\circ$

$$d = \frac{n\lambda}{2\sin\theta} = \frac{1 \times 0.071 \times 10^{-9} \text{ m}}{2\sin 12^\circ}$$

$$d = \frac{0.071 \times 10^{-9} \text{ m}}{2 \times 0.2079} = 0.171 \text{ nm}$$

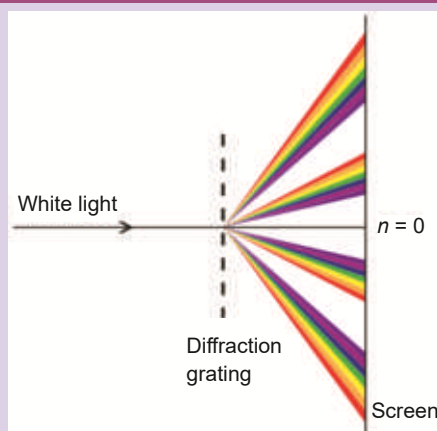
This spacing corresponds to a particular set of planes in the crystal lattice.

Do you know?



A CD ROM acts as a reflection diffraction grating. White light is reflected and diffracted at its surface, producing a display of spectral colours.

Tidbit



A diffraction grating can split white light into its constituent colours. However central fringe remain white.

QUESTIONS

Multiple Choice Questions

Choose the correction answer.

- 15.1 When the light from two lamps falls on a screen, no interference pattern can be obtained. Why is this?
- (a) The lamps are not point sources
 - (b) The lamps emit light of different amplitudes
 - (c) The light from the lamps is not coherent
 - (d) The light from the lamps is white
- 15.2 Monochromatic light illuminates two narrow parallel slits. The interference pattern which results is observed on a screen some distance beyond the slits. Which change increases the separation between the dark lines of the interference pattern?
- (a) Decreasing the distance between the screen and the slits
 - (b) Increasing the distance between the slits
 - (c) Monochromatic light of higher frequency
 - (d) Using monochromatic light of longer wavelength
- 15.3 Using monochromatic light, interference fringes are produced on a screen placed at a distance L from a pair of slits of separation d . The separation of the fringes is y . Both d and L are now doubled. What is the new fringe separation?
- (a) $y/2$ (b) y (c) $2y$ (d) $4y$
- 15.4 Which of the following is not a necessary condition for two waves to produce interference?
- (a) Waves must be of same type
 - (b) Waves must be coherent
 - (c) Waves must have same amplitude
 - (d) Waves must be moving in same direction
- 15.5 The colours observed on a soap bubble are due to which phenomenon of light?
- (a) Reflection (b) Refraction
- (c) Interference in thin film (d) Diffraction
- 15.6 Continuous water waves are diffracted through a gap in a barrier in a ripple tank. Which change will cause the diffraction of the waves to increase?
- (a) Increasing the frequency of the waves
 - (b) Increasing the width of the gap
 - (c) Reducing the wavelength of the waves
 - (d) Reducing the width of the gap

- 15.7 Which property of a light wave can be determined using a diffraction grating?
(a) Amplitude (b) Intensity (c) Speed (d) Wavelength
- 15.8 The central maximum in a diffraction grating is called:
(a) zero order (b) first order (c) second order (d) higher order

Short Answer Questions

- 15.1 Can Huygen's principle be applied on sound and water waves? Explain briefly.
- 15.2 Why interference pattern is not produced due to headlights of a car? Explain briefly.
- 15.3 Why interference pattern is observed on soap bubbles but not on window glass or windscreens of cars?
- 15.4 If the whole apparatus of Young's double-slit experiment is submerged in water, what will be the effect on fringes?
- 15.5 Compare the double-slit experimental arrangements for light and sound waves, similarities and differences.
- 15.6 For a diffraction grating, what are the advantages of using: (a) many slits (b) narrow slits (c) closely spaced slits?
- 15.7 How the number of orders in a given diffraction grating can be increased?
- 15.8 How the pattern would change if red light is replaced with blue light in a given diffraction grating?
- 15.9 Why central spot in Newton's rings is dark?

Constructed Response Questions

- 15.1 Light is sometimes described as rays and sometimes as waves. Construct diagrams to explain differences in both.
- 15.2 Suppose white light is falling on a double slit arrangement, such that one slit is covered with red filter (650 nm) and other is with blue filter (450 nm), explain what would be observed on the screen?
- 15.3 What will be the effect on width and brightness of the fringes in a double-slit interference experiment, if we decrease the width of one slit to half keeping all other factors same?
- 15.4 Why can we readily observe diffraction effects for sound waves and water waves, but not for light waves? Is this because light travels so much faster than these other waves? Explain.
- 15.5 What is the significance of path difference in deriving Bragg's equation? Explain.

Comprehensive Questions

- 15.1 Define interference of light. Explain the concept of coherent sources. Why are coherent sources necessary for sustained interference?

- 15.2 Describe the arrangement of Young's double slit experiment. Derive the expression for fringe width. How does it demonstrate the wave nature of light?
- 15.3 Explain the formation of interference in thin films with the help of bright and dark fringes.
- 15.4 Explain X-ray diffraction and derive Bragg's equation.

Numerical Problems

- 15.1 Coherent light of wavelength 500 nm is incident on two very narrow and closely spaced slits. The interference pattern is observed on a very tall screen that is 2.00 m from the slits. Near the center of the screen the separation between two adjacent interference maxima is 3.53 cm. What is the distance on the screen between the $m = 49$ and $m = 50$ maxima? (Ans: 3.53 cm)
- 15.2 Young's experiment is performed with light from excited helium atoms $\lambda = 502$ nm. Fringes are measured carefully on a screen 1.20 m away from the double slit, and the center of the 20th fringe (not counting the central bright fringe) is found to be 10.6 mm from the center of the central bright fringe. What is the separation of the two slits? (Ans: 1.14×10^{-3} m)
- 15.3 Coherent light with wavelength 450 nm falls on a pair of slits. On a screen 1.80 m away, the distance between dark fringes is 3.90 mm. What is the slit separation? (Ans: 2.08×10^{-4} m)
- 15.4 Two slits spaced 0.450 mm apart are placed 75.0 cm from a screen. What is the distance between the second and third dark lines of the interference pattern on the screen when the slits are illuminated with coherent light with a wavelength of 500 nm? (Ans: 8.33×10^{-4} cm)
- 15.5 Two very narrow slits are spaced 1.80 mm apart and are placed 35.0 cm from a screen. What is the distance between the first and second dark lines of the interference pattern when the slits are illuminated with coherent light of wavelength 550 nm? (Ans: 1.07×10^{-4} m)
- 15.6 Monochromatic light is at normal incidence on a plane transmission grating. The first-order maximum in the interference pattern is at an angle of 11.3° . What is the angular position of the fourth-order maximum? (Ans: 51.7°)
- 15.7 If a diffraction grating produces a third-order bright spot for red light of wavelength 700 nm at 65° from the central maximum, at what angle will the second-order bright spot be for violet light of wavelength 400 nm? (Ans: 20.2°)

Electrostatics

Students Learning Outcomes

After studying this chapter, the students will be able to:

- ◆ define and calculate electric potential [At a point as the work done per unit positive charge in bringing a small test charge from infinity to the point. Use for the electric $V_r = \frac{1}{4\pi\epsilon_0} \frac{q}{r}$ potential in the field due to a point charge].
- ◆ use the fact that the electric field at a point is equal to the negative of potential gradient at that point.
- ◆ state how the concept of electric potential leads to the electric potential energy of two point charges and use $V = \frac{W}{q_0}$.
- ◆ define and calculate capacitance [as applied to both isolated spherical conductors and to parallel plate capacitors].
- ◆ Derive and apply formulae for the combined capacitance of capacitors in series and in parallel.
- ◆ use the capacitance formula for capacitors in series and in parallel.
- ◆ determine the electric potential energy stored in a capacitor from the area under the potential charge graph [Use to solve $P.E. = \frac{1}{2} CV^2$ physics related problems].
- ◆ analyze graphs of the variation with time of potential difference, charge and current for a capacitor discharging through a resistor [use of time constant for a capacitor discharging through a resistor $RC = \tau$].
- ◆ Use equations of the form $x = x_0 e^{-\frac{t}{RC}}$ [where x could represent current, charge or potential difference for a capacitor discharging through a resistor].
- ◆ list the use of capacitors in various household appliances [such as in flash guns, refrigerators, electric fans, rectification circuits, etc.].

Electrostatics or static electricity refers to the phenomena arising from the stationary electric charges. These include forces exerted by charges on each other, electric field and electric potential produced by them and electrical energy stored in space and in various devices.

We have already studied about the electric field produced by a charge around it. Any other charge placed in this field experiences a force. Quantitatively, the electrostatic force can be found by applying Coulomb's law $F_e = \frac{1}{4\pi\epsilon_0} \frac{q_1 q_2}{r^2}$.

The direction of this force depends on the nature of charges whether positive or negative. In this chapter, we will study in detail the electric potential produced by a charge at any point in its field. In addition, we will see that electrical energy and charge can be stored in a widely used device known as a capacitor. We will also become acquainted with some important uses of capacitors.

16.1 ELECTRIC POTENTIAL DIFFERENCE

Recalling the electric potential difference studied in the previous class, the electric potential difference is the work done per unit charge in moving it between two points against the electric field (Fig. 16.1) as:

$$\Delta V = V_B - V_A = \frac{W_{AB}}{q_o} = \frac{\Delta U}{q_o} \dots\dots\dots(16.1)$$

Its SI unit is volt, where

$$1V = \frac{1 \text{ J}}{1 \text{ C}} \dots\dots\dots(16.2)$$

It can be converted into other units such as kilovolts (kV) or millivolts (mV). The relationship between electric potential difference and the electric field is:

$$\Delta V = -\mathbf{E} \cdot \Delta \mathbf{d} \dots\dots\dots(16.3)$$

It is given by the fact that the electric field represents the rate of change of potential with distance called gradient of potential and expressed as:

$$\mathbf{E} = -\frac{\Delta V}{\Delta d} \dots\dots\dots(16.4)$$

where the negative sign indicates that the field points from higher to lower potential.

In the case of uniform electric field, this relation simplifies to:

$$\mathbf{E} = \frac{V}{d} \dots\dots\dots(16.5)$$

This shows that the potential difference between two points is directly proportional to the electric field strength and the distance between them. Thus, a stronger electric field causes more rapid decrease in potential, providing a clear physical link between field intensity and voltage variation in space.

Example 16.1 A 6 volts battery is connected to two parallel copper plates separated by a distance of 0.3 cm. Find the electric field intensity between them.

Solution

Potential difference between the plates $\Delta V = 6 \text{ V}$

Distance between the plates $\Delta d = 0.3 \text{ cm} = 0.003 \text{ m}$

Using the equation; $E = \frac{\Delta V}{\Delta d}$

Putting the values; $E = \frac{6 \text{ V}}{0.003 \text{ m}} = 2000 \text{ V m}^{-1}$ or 2000 N C^{-1}

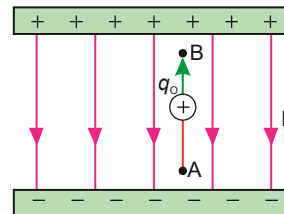


Fig. 16.1: Uniform electric field between two charged parallel plates

For your information

Although we have assumed that the points at infinity are at zero potential, yet the Earth is assumed to be at zero potential practically. Therefore, absolute value of the electric potential at a point is measured with respect to the Earth.

16.2 ELECTRIC POTENTIAL

In order to give a concept of electric potential at a point in an electric field due to a point charge, we must have a reference to which we assign zero electric potential. This point

is usually taken at infinity.

If we take A to be at infinity and choose $V_A = 0$, the electric potential at B will be $V_B = W_B / q_0$ or dropping the subscripts:

$$V = \frac{W}{q_0} \dots\dots\dots(16.6)$$

which states that;

The electric potential at any point in an electric field is equal to work done in bringing a unit positive charge from infinity to that point keeping it in electrostatic equilibrium.

It is to be noted that potential at a point is still potential difference between the potential at that point and potential at infinity. Both potential and potential differences are scalar quantities because both W and q_0 are scalars.

Let us derive an expression for the potential at a certain point in the field of a positive point charge q . This can be accomplished by bringing a unit positive charge from infinity to that point keeping it in electrostatic equilibrium. The target can be achieved by using Eq.16.4 in the form $\Delta V = -E\Delta r$, provided electric intensity E remains constant. However, in this case, E varies inversely as square of distance from the point charge, it no more remains constant, so we use basic principles to compute the electric potential at a point. The field is radial as shown in Fig. 16.2.

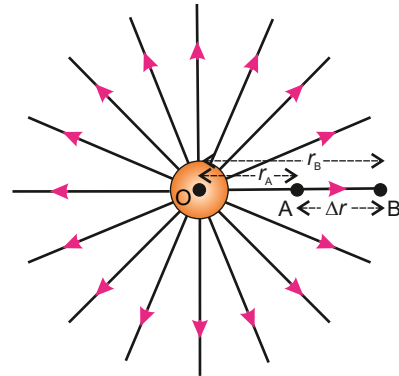


Fig. 16.2: Electric field due to a point charge q

Let us take two points A and B, infinitesimally close to each other, so that E remains almost constant between them. The distances of points A and B from q are r_A and r_B respectively and distance of mid-point of space interval between A and B is r from q . Then according to Fig. 16.2

$$r_B = r_A + \Delta r \dots\dots\dots(16.7)$$

$$\Delta r = r_B - r_A \dots\dots\dots(16.8)$$

As r represents mid-point of interval between A and B, so

$$r = \frac{r_A + r_B}{2} \dots\dots\dots(16.9)$$

The magnitude of electric intensity at this point is:

$$E = \frac{1}{4\pi\epsilon_0} \frac{q}{r^2} \dots\dots\dots(16.10)$$

As the points A and B are very close, therefore, as a first approximation, we can take the arithmetic mean to be equal to geometric mean which gives:

$$r^2 = r_A r_B \dots\dots\dots(16.11)$$

Do you know?



Bug-zappers with their characteristic bluish glow, attract certain insects and electrocute them with very high voltage of electricity.

Thus, Eq. 16.10 can be written as:

$$E = \frac{1}{4\pi\epsilon_0} \frac{q}{r_A r_B} \dots\dots\dots (16.12)$$

Now, if a unit positive charge is moved from B to A, the work done is equal to the potential difference between A and B.

$$V_A - V_B = -E(r_A - r_B)$$

$$V_A - V_B = E(r_B - r_A) \dots\dots\dots (16.13)$$

Substituting the value of E from Eq. 16.12., we have

$$V_A - V_B = \frac{1}{4\pi\epsilon_0} \left(\frac{r_B - r_A}{r_A r_B} \right) \dots\dots\dots (16.14)$$

$$V_A - V_B = \frac{1}{4\pi\epsilon_0} \left(\frac{1}{r_A} - \frac{1}{r_B} \right) \dots\dots\dots (16.15)$$

To calculate absolute potential or potential at A, point B is assumed to be infinity point so that $V_B = 0$, and, hence

$$\frac{1}{r_B} = \frac{1}{r_\infty} = \frac{1}{\infty} = 0$$

This gives
$$V_A = \frac{1}{4\pi\epsilon_0} \frac{q}{r_A} \dots\dots\dots (16.16)$$

The general expression for electric potential V_r at a distance r from q is:

$$V_r = \frac{1}{4\pi\epsilon_0} \left(\frac{q}{r} \right) \dots\dots\dots (16.17)$$

The electric potential will be negative if the charge q is negative. However, if there are more than one charge at different locations, then each charge contributes to the total electric potential. At any point, the total electric potential is the algebraic sum of the individual potentials created by all the charges.

Example 16.2 Find the electric potential at a point 1.2 m from a charge of 5.0×10^{-8} C.

Solution $q = 5.0 \times 10^{-8}$ C, $r = 1.2$ m, $V = ?$

Using the equation; $V = \frac{1}{4\pi\epsilon_0} \frac{q}{r}$

Putting the values; $V = \frac{(9 \times 10^9 \text{ N m}^2 \text{C}^{-2}) (5.0 \times 10^{-8} \text{ C})}{1.2 \text{ m}}$

$$V = 375 \text{ V}$$

Example 16.3 Two point charges $+8 \mu\text{C}$ and $-4 \mu\text{C}$ are placed 60 cm apart as shown in Fig. 16.3. Where is the electric potential zero on the line passing through the charges?

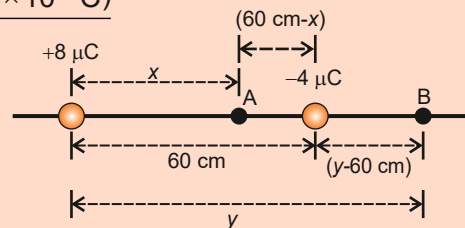


Fig. 16.3

Solution

The electric potential at any point will be zero where the potential due to positive charge will be exactly equal in magnitude to that due to negative charge. Since the magnitude of positive charge is greater than that of negative charge, so there will be two zero potential points. Both the points should be closer to the smaller negative charge. That is; one point A between the two charges and the other point B outside the negative charge.

Let the point A lies at a distance x from the positive charge, then using $V = \frac{1}{4\pi\epsilon_0} \frac{q}{r}$, we have, separation between charges = 60 cm = 0.6 m

$$\text{Electric potential at A due to positive charge} = \frac{1}{4\pi\epsilon_0} \left(\frac{8 \times 10^{-6} \text{ C}}{x} \right)$$

$$\text{Electric potential at A due to negative charge} = \frac{1}{4\pi\epsilon_0} \left(\frac{4 \times 10^{-6} \text{ C}}{0.6 \text{ m} - x} \right)$$

As both the potentials balance each other, therefore,

$$\frac{1}{4\pi\epsilon_0} \left(\frac{8 \times 10^{-6} \text{ C}}{x} \right) = \frac{1}{4\pi\epsilon_0} \left(\frac{4 \times 10^{-6} \text{ C}}{0.6 \text{ m} - x} \right)$$

$$\begin{aligned} 4.8 \text{ m} - 8x &= 4x \\ 12x &= 4.8 \text{ m} \\ x &= 0.4 \text{ m} = 40 \text{ cm} \end{aligned}$$

$$\text{Electric potential at B due to positive charge} = \frac{1}{4\pi\epsilon_0} \left(\frac{8 \times 10^{-6} \text{ C}}{y} \right)$$

$$\text{Electric potential at B due to negative charge} = \frac{1}{4\pi\epsilon_0} \left(\frac{4 \times 10^{-6} \text{ C}}{y - 0.6 \text{ m}} \right)$$

Applying balance condition,

$$\frac{1}{4\pi\epsilon_0} \left(\frac{4 \times 10^{-6} \text{ C}}{y - 0.6 \text{ m}} \right) = \frac{1}{4\pi\epsilon_0} \left(\frac{8 \times 10^{-6} \text{ C}}{y} \right)$$

$$4y = 8(y - 0.6 \text{ m})$$

$$y = 2(y - 0.6 \text{ m})$$

$$y = 2y - 1.2 \text{ m}$$

$$y = 1.2 \text{ m}$$

$$y = 120 \text{ cm}$$

For your information

The energy used to operate an electric car comes from the energy stored in the car's battery pack. The voltage of the pack and the magnitude of the charge determine how much energy is consumed from the battery.

16.3 ⚡ ELECTRIC POTENTIAL ENERGY

We have learnt that a charge placed in an electric field possesses some electrical potential energy. If a charge q_0 is placed at any point in the electric field due to a charge q and allowed to move it freely, it will be repelled away by the charge q . Equation (16.17) shows that the potential decreases with the increase of distance. Therefore, if a positive charge is allowed to move freely in an electric field, it will always accelerate from a point of higher potential to a point of lower potential. If the potential difference between these points be ΔV , the potential energy of the charge will decrease by $\Delta U = q_0 \Delta V$. This decrease in potential energy of the charge will appear as its kinetic energy.

If m is the mass of the body carrying charge q_0 , then the body will reach the point at lower potential energy with velocity v .

$$\frac{1}{2}mv^2 = q_0 \Delta V \dots\dots\dots (16.18)$$

A negative charge allowed to move freely in an electric field, moves from a point of lower potential to a point of higher potential.

Equation (16.18) suggests that energy can also be expressed as the product of potential difference and charge. In atomic physics, we use one such unit of energy known as electron-volt (eV). It is the energy acquired or lost by an electron which has been accelerated through a potential difference of 1 volt.

Hence $1 \text{ eV} = \text{charge on an electron} \times 1 \text{ volt}$

or $1 \text{ eV} = 1.6 \times 10^{-19} \text{ J}$

Electric Potential Energy of Two Point Charges

The concept of electric potential directly leads to the expression for electric potential energy of two point charges. Consider two points A and B distance r apart (Fig.16.4).

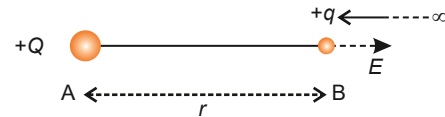


Fig. 16.4: Two charges Q and q are separated by a distance r

The charge Q has been placed at point A, an electric potential V is created by it at point B.

$$V = \frac{1}{4\pi\epsilon_0} \frac{Q}{r} \dots\dots\dots (16.19)$$

If an other charge q is brought from infinity to point B, work has to be done in moving it against the electric field. This work is:

$$W_B = q \times (\text{Electric potential})$$

$$W_B = qV \dots\dots\dots (16.20)$$

Putting the value of V in equation 16.20, we have

$$W_B = q \left(\frac{1}{4\pi\epsilon_0} \frac{Q}{r} \right)$$

$$\text{or } W_B = \frac{1}{4\pi\epsilon_0} \frac{Qq}{r} \dots\dots\dots (16.21)$$

Total work done in placing the two charges at point A and B (dropping the subscript) is:

$$\text{or } W = \frac{1}{4\pi\epsilon_0} \frac{Qq}{r} \dots\dots\dots (16.22)$$

Therefore, by definition, the work done in placing these two charges at points A and B is stored as electric potential energy (U) in the system which is given by

$$U = \frac{1}{4\pi\epsilon_0} \frac{Qq}{r} \dots\dots\dots(16.23)$$

Example 16.4 Two point charges are separated by a distance of 4.0 cm. If the charges are $8\text{ }\mu\text{C}$ and $5\text{ }\mu\text{C}$, what is the electric potential energy of the two charges?

Solution $r = 4 \times 10^{-2}\text{ m}$, $Q = 8\text{ }\mu\text{C} = 8 \times 10^{-6}\text{ C}$, $q = 5\text{ }\mu\text{C} = 5 \times 10^{-6}\text{ C}$

Using the formula; $U = \frac{1}{4\pi\epsilon_0} \frac{Qq}{r}$

$$U = 9 \times 10^9 \text{ N m}^2 \text{ C}^{-2} \left[\frac{(8 \times 10^{-6} \text{ C})(5 \times 10^{-6} \text{ C})}{4 \times 10^{-2} \text{ m}} \right]$$

$$U = 9 \text{ J}$$

16.4 CAPACITORS

A capacitor is a device that can store electrical energy in the form of electric field. We will discuss here two types of such devices:

- (i) Isolated spherical conductors
- (ii) Parallel plate capacitors

Isolated Spherical Conductors

A spherical conductor placed in space such that it is far from other conductors or influence is known as an isolated spherical conductor.

The ability of an isolated hollow spherical conductor to store electric charge is called capacitance. It is denoted by C .

The capacitance of an isolated hollow spherical conductor is defined as the electric charge stored on the sphere per unit electric potential at its surface.

Mathematically, we can write as:

$$C = \frac{Q}{V} \dots\dots\dots(16.24)$$

where Q is the amount of charge stored and V is the electric potential created at the surface of sphere due to charge Q (Fig. 16.5). The SI unit of capacitance is farad (F).

For a spherical conductor, 1 farad is the capacitance of a sphere that stores 1 coulomb of charge when its electric potential is 1 volt, assuming it is isolated in free space.

$$1 \text{ farad} = \frac{1 \text{ coulomb}}{1 \text{ volt}}$$

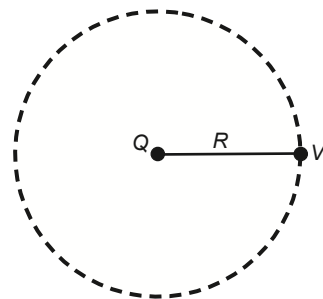


Fig. 16.5: Isolated spherical conductor with charge at the centre

The electric potential V at the surface of the sphere of radius R is given by

$$V = \frac{1}{4\pi\epsilon_0} \frac{Q}{R}$$

Substituting the value of V into Eq. 16.24, we have

$$C = \frac{Q}{\left(\frac{1}{4\pi\epsilon_0} \frac{Q}{R} \right)}$$

or $C_{\text{vac}} = 4\pi\epsilon_0 R \dots\dots\dots (16.25)$

However, the capacitance of a spherical conductor in a medium of dielectric constant ϵ_r is:

$$C_{\text{med}} = 4\pi\epsilon_0 \epsilon_r R \dots\dots\dots (16.26)$$

Equation (16.25) shows that the capacitance of a spherical conductor is independent of the charge or potential. It depends on the radius of the sphere and medium. The capacitance increases with increase in radius of the sphere. This means larger spheres store more charge at the same surface potential. The presence of a dielectric medium also increases the capacitance of the device.

16.5 PARALLEL PLATE CAPACITORS

A parallel plate, capacitor consists of two conductors placed near one another separated by vacuum, air or any other insulator, known as dielectric. Usually, the conductors are in the form of parallel plates, and the capacitor is known as parallel plate capacitor. When the plates of such a capacitor are connected to a battery of voltage V (Fig.16.6). It establishes a potential difference of V volts between the two plates and the battery places a charge $+Q$ on the plate connected with its positive terminal and a charge $-Q$ on the other plate, connected to its negative terminal. Let Q be the magnitude of the charge on either of the plates. It is found that;

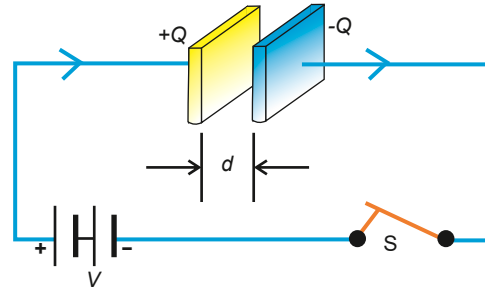


Fig. 16.6: Parallel plate capacitor

$$Q \propto V \quad \text{or} \quad Q = CV \dots\dots\dots (16.27)$$

$$\text{or} \quad C = \frac{Q}{V} \dots\dots\dots (16.28)$$

The proportionality constant C is called the capacitance of the capacitor. As we shall see later, it depends upon the geometry of the plates and the medium between them. It is a measure of the ability of capacitor to store charge.

The capacitance of a capacitor can be defined as the amount of charge on one plate necessary to raise the potential of that plate by one volt with respect to the other.

Capacitance of a Parallel Plate Capacitor

Consider a parallel plate capacitor consisting of two parallel metal plates, each of area A , separated by a distance d as shown in Fig.16.6. The distance d is small so that the electric field E between the plates is uniform and confined almost entirely in the region between the plates. Let initially the medium between the plates be air or vacuum. Then, according to Eq.16.28

$$C_{\text{vac}} = \frac{Q}{V} \dots\dots\dots (16.29)$$

where Q is the charge on the capacitor and V is the potential difference between the parallel plates. From Eq. (16.5), the magnitude E of electric intensity is related with the distance d as:

$$E = \frac{V}{d} \dots\dots\dots (16.30)$$

As electric intensity between parallel plate capacitor is: $E = \frac{\sigma}{\epsilon_0}$ (Gauss's law)

We can define surface charge density as:

$$\sigma = \frac{Q}{A} \text{ or } Q = \sigma A$$

$$\text{Thus } E = \frac{Q}{A\epsilon_0} \dots\dots\dots (16.31)$$

From Eqs. (16.30) and (16.31)

$$\frac{V}{d} = \frac{Q}{A\epsilon_0}$$

$$\text{or } \frac{Q}{V} = \frac{A\epsilon_0}{d}$$

$$\text{As } C_{\text{vac}} = \frac{Q}{V}, \text{ therefore,}$$

$$C_{\text{vac}} = \frac{A\epsilon_0}{d} \dots\dots\dots (16.32)$$

If an insulating material, called dielectric, of relative permittivity ϵ_r , is introduced between the plates, the capacitance of capacitor is enhanced by the factor ϵ_r . Capacitors commonly have some dielectric medium, thereby ϵ_r is also called as dielectric constant.

Following experiment gives the effect of insertion of dielectric between the plates of a capacitor.

Consider a charged capacitor whose plates are connected to a voltmeter (Fig. 16.7-a). The deflection of the meter is a measure of the potential difference between the plates. When a dielectric material is inserted between the plates, reading drops indicating a decrease in the potential difference between the plates (Fig. 16.7 b). From the definition; $C = Q/V$, since V decreases while Q remains constant, the value of C increases. Then, Eq.16.32 becomes:

$$C_{\text{med}} = \frac{A\epsilon_0\epsilon_r}{d} \dots\dots\dots(16.33)$$

Equation 16.32 shows the dependence of a capacitor upon the area of plates, the separation between the plates and medium between them.

Dividing Eq.16.33 by Eq.16.32, we have an expression for dielectric constant as:

$$\epsilon_r = \frac{C_{\text{med}}}{C_{\text{vac}}} \dots\dots\dots(16.34)$$

From Eq. 16.34 dielectric co-efficient or dielectric constant is defined as:

The ratio of the capacitance of a parallel plate capacitor with an insulating substance as medium between the plates to its capacitance with vacuum (or air) as medium between them.

The reason for decrease in potential difference between the plates of the capacitor is the polarization of the atoms of the dielectric medium. A dielectric is an insulator. It has no free electrons, i.e., all the electrons are bound to their respective positively charged nuclei. When such an atom is subjected to an electric field, its electrons are displaced slightly relative to the nucleus in a direction opposite to the electric field as shown in Fig. 16.7 (c). This phenomena is known as polarization of the dielectric and the distorted atom is said to be polarized.

When a slab of dielectric is placed between the two oppositely charged plates, its atoms get polarized due to the electric field in between the two plates as shown in Fig. 16.7 (d). As a result, charges appear on those surfaces of the dielectric slab which are in contact with the metal plates. The polarity of these charges is opposite to that of the charge on the adjacent plate. This effectively decreases the charge on the plates. Consequently, the electric intensity and potential difference between the plates decrease.

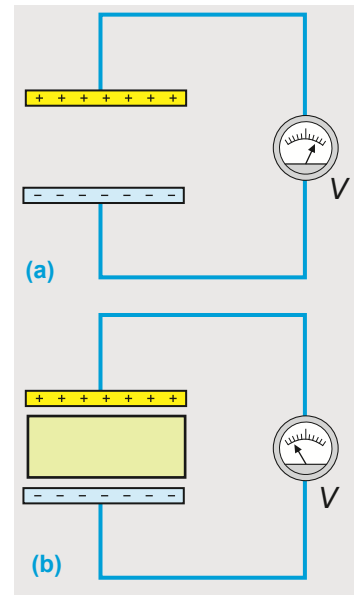


Fig. 16.7: (a) Without dielectric (b) With dielectric

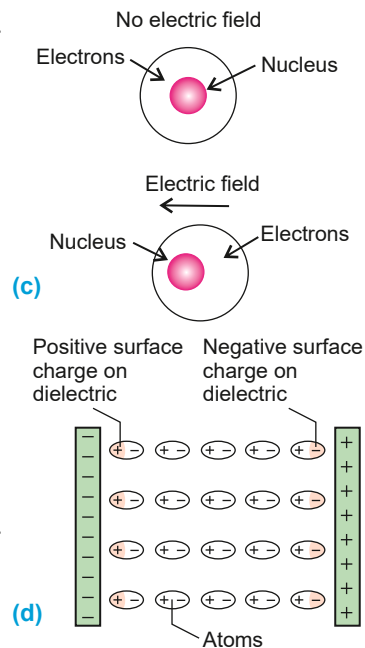


Fig. 16.7: Electric polarization of a dielectric

16.6 COMBINED CAPACITANCE OF CAPACITORS

When two or more capacitors are connected together in a circuit, it is known as combination of capacitors.

It is done to obtain a desired capacitance. In this process, multiple connections of capacitors behave as a single equivalent capacitor. Usually, it is named as combined capacitance. Two common types of combination are:

- (i) Parallel combination
- (ii) Series combination

Parallel Combination

In this method, the plates of one side of all the capacitors are connected to the negative terminal and plates of the other side to the positive terminal of the battery as shown in Fig. 16.8(a). The plates are charged by the battery. In this combination, the potential difference V across the plates of each capacitor is the same. The charge on each will be different depending upon its capacitance. If Q_1 , Q_2 , and Q_3 are the charges on the capacitors C_1 , C_2 and C_3 respectively, then the amount of charge Q on equivalent capacitor C_e would be the same as all the three capacitors hold together under the same potential difference V (Fig. 16.8-b). So,

$$Q = Q_1 + Q_2 + Q_3$$

Putting $Q_1 = C_1 V$, $Q_2 = C_2 V$

$$Q_3 = C_3 V \text{ and } Q = C_e V$$

we have $C_e V = C_1 V + C_2 V + C_3 V$

$$\text{or } C_e = C_1 + C_2 + C_3 \dots\dots\dots (16.35)$$

When capacitors are connected in parallel, their combined capacitance is equal to the sum of their individual capacitances.

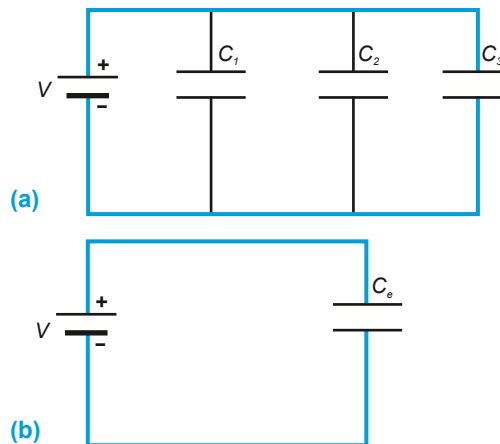


Fig. 16.8: (a) Parallel combination
(b) Equivalent capacitance

Example 16.5 A capacitor of 80 pF capacitance is charged to potential difference of 48 volts. Its plates are then connected in parallel to another capacitor. The potential difference across the combination is found to be 30 volts. What is the capacitance of the second capacitor?

Solution

Before connecting the capacitor C_1 to C_2 ; using formula $Q = CV$

Charge on C_1 is: $Q_1 = C_1 V = 80 \times 10^{-12} \text{ F} \times 48 \text{ V}$

$$\text{or } Q_1 = 3840 \times 10^{-12} \text{ C}$$

After connecting C_2 in parallel to C_1 , $V_2 = 30 \text{ V}$

Now charge on C_1 is; $Q'_1 = C_1 V_2$

$$= 80 \times 10^{-12} \text{ F} \times 30 \text{ V}$$

$$= 2400 \times 10^{-12} \text{ C}$$

As total charge remains constant, so charge shifted to C_2 is:

$$Q_2 = Q_1 - Q'_1 = 3840 \times 10^{-12} \text{ C} - 2400 \times 10^{-12} \text{ C}$$

$$Q_2 = 1440 \times 10^{-12} \text{ C}$$

To determine the capacitance C_2 of the second capacitor, we use the formula:

$$C = \frac{Q}{V}$$

$$\text{or } C_2 = \frac{Q_2}{V_2}$$

$$\text{or } C_2 = \frac{1440 \times 10^{-12} \text{ C}}{30 \text{ V}} = 48 \times 10^{-12} \text{ F} = 48 \text{ pF}$$

Series Combination

The series combination is one after the other in a line. The circuit is shown in Fig. 16.9 (a). We can see that the second plate of capacitor C_1 is connected to the first plate of C_2 . Likewise, the second plate of C_2 is connected to the first plate of C_3 . The first plate of C_1 and the second plate of C_3 are connected to the battery.

The battery sends positive charge Q to the first plate of capacitor C_1 . As the second plate of C_1 is insulated from the first plate, so it cannot draw any charge from the battery. But the positive of the first plate induces an equal negative charge $-Q$ on the second plate. Since the first plate of capacitor C_2 is connected to the second plate of C_1 , therefore, a charge $+Q$ appears on the first plate of C_2 . Similarly, the capacitor C_3 is also charged. Thus, each capacitor receives a charge of magnitude Q on each of its plates.

Figure 16.9(b) shows that the potential difference across the combination of three capacitors is equal to that of battery terminals V . As C_1 , C_2 , and C_3 are connected in series, so potential difference across each of them will not be V . It will be different across each of them depending upon its capacitance. If V_1 , V_2 and V_3 are the values of potential difference across them, then

$$V = V_1 + V_2 + V_3 \dots \dots \dots (16.36)$$

If C_e is the capacitance of equivalent capacitor, the charge on it will also be equal to Q .

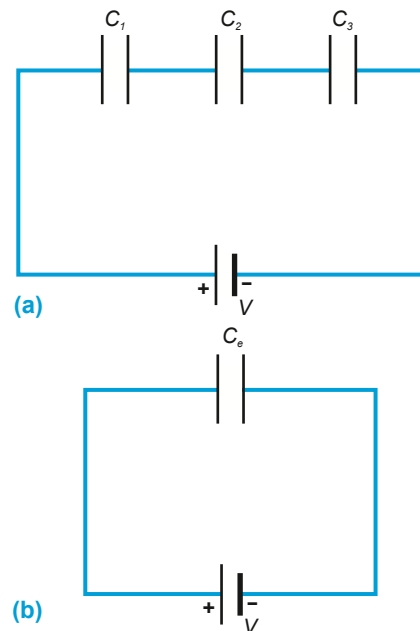


Fig. 16.9: (a) Series combination
(b) Equivalent capacitance

Also the potential difference across C_e will be V which is the potential difference across the combination.

Using the formula; $C = \frac{Q}{V}$ or $V = \frac{Q}{C}$, we have

$$V = \frac{Q}{C_e}, \quad V_1 = \frac{Q}{C_1}, \quad V_2 = \frac{Q}{C_2} \quad \text{and} \quad V_3 = \frac{Q}{C_3}$$

Putting these values in equation (16.36), we have

$$\frac{Q}{C_e} = \frac{Q}{C_1} + \frac{Q}{C_2} + \frac{Q}{C_3}$$

$$\text{or} \quad \frac{1}{C_e} = \frac{1}{C_1} + \frac{1}{C_2} + \frac{1}{C_3} \dots\dots\dots(16.37)$$

Thus,

The reciprocal of equivalent capacitance is equal to the sum of the reciprocals of individual capacitances of the capacitors connected in series.

Example 16.6 Two capacitors of capacitances $12 \mu\text{F}$ and $4 \mu\text{F}$ are connected in series to a 80 V battery.

- (i) What is the total charge stored in the combination?
- (ii) What is the potential difference across each capacitor?

Solution

For series combination;

$$\frac{1}{C_e} = \frac{1}{C_1} + \frac{1}{C_2}$$

Putting the values of C_1 , and C_2 , we have

$$\begin{aligned} \frac{1}{C_e} &= \frac{1}{12 \times 10^{-6} \text{ F}} + \frac{1}{4 \times 10^{-6} \text{ F}} \\ &= \frac{1+3}{12 \times 10^{-6} \text{ F}} = \frac{4}{12 \times 10^{-6} \text{ F}} = \frac{1}{3 \times 10^{-6} \text{ F}} \\ C_e &= 3 \times 10^{-6} \text{ F} \end{aligned}$$

- (i) Total charge on C_e will be given by the formula $Q = CV$,

$$\text{Therefore,} \quad Q = C_e V = 3 \times 10^{-6} \text{ F} \times 80 \text{ V}$$

$$Q = 2.4 \times 10^{-4} \text{ C}$$

- (ii) In series combination, charge on each capacitor will be the same Q . Therefore, potential difference across C_1 will be:

$$V_1 = \frac{Q}{C_1} = \frac{2.4 \times 10^{-4} \text{ C}}{12 \times 10^{-6} \text{ F}} = 20 \text{ V}$$

Potential difference across C_2 will be:

$$V_2 = \frac{Q}{C_2} = \frac{2.4 \times 10^{-4} \text{ C}}{4 \times 10^{-6} \text{ F}} = 60 \text{ V}$$

16.7 ENERGY STORED IN A CAPACITOR

A capacitor is a device to store charge. Alternatively, it is possible to think of a capacitor as a device for storing electrical energy.

A good evidence of the energy stored in a charged capacitor is that of its discharging. If we connect the two plates of charged capacitor by means of a thick wire, the charge flows from one plate to the other. As a result, the wire becomes hot. Since heat is a form of energy, so the capacitor possessed energy when it was charged, and it was this energy which appeared as heat.

Let us determine the electric potential energy stored in a charged capacitor graphically. We have already studied that the charge Q on the plates increases with the increase in potential V between the plates of the capacitor.

i.e., $Q \propto V$

Therefore, if we draw a V vs Q graph, it will be a straight line passing through the origin as shown in Fig. 16.10.

We can understand the charging of capacitor by dividing the area under the graph into very small increments of charge ΔQ each as shown in Fig. 16.10. Generally, the work required to move a small charge ΔQ through a potential difference V across the plates of capacitor is:

$$\Delta W = (\Delta Q)(V_1) \dots \dots \dots (16.38)$$

But $(\Delta Q)(V_1)$ is equal to the area of the rectangle shown red in the figure. Therefore total work done to charge the capacitor to a final value Q is equal to the sum of areas of all such rectangles under V vs Q graph. Fig. 16.10 shows that this sum is equal to the total area under the graph, i.e., area of the right angled triangle OAB .

$$\text{Area of triangle } OAB = \frac{1}{2} (\text{base}) (\text{height}) = \frac{1}{2} (OA) (AB)$$

If OA at this stage is equal to total charge Q and AB is equal to the final potential difference V , then

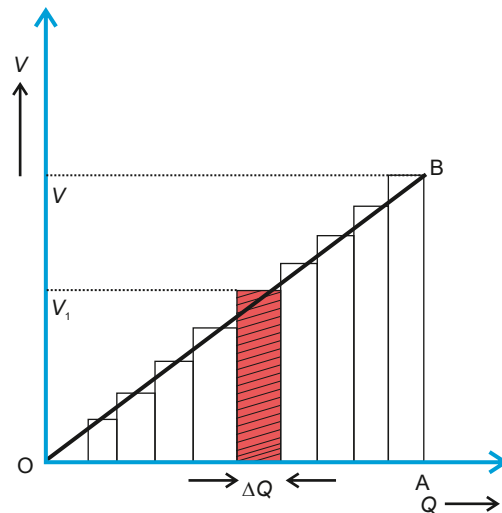


Fig. 16.10: Energy stored in the capacitor

Area of the triangle OAB = $\frac{1}{2} QV$

Therefore, energy stored in a charged capacitor is:

$$W = \frac{1}{2} QV \dots \dots \dots (16.39)$$

If we put the value of $Q = CV$ in Eq. (16.39), we have

$$\text{Energy stored} = W = \frac{1}{2} (CV)V$$

$$\text{or} \quad W = \frac{1}{2} CV^2 \dots \dots \dots (16.40)$$

$$\text{Therefore,} \quad P.E = \frac{1}{2} CV^2 \dots \dots \dots (16.41)$$

Energy can be stored in electric field between the plates, rather than the potential energy of the charges on the plates.

This relation can be obtained by substituting $V = Ed$ and $C = \frac{A\epsilon_0\epsilon_r}{d}$ in Eq. (16.41).

$$\begin{aligned} \text{Energy} &= \frac{1}{2} \left(\frac{A\epsilon_0\epsilon_r}{d} \right) (Ed)^2 \\ &= \frac{1}{2} \epsilon_0\epsilon_r E^2 \times (Ad) \end{aligned}$$

As (Ad) is the volume between the plates.

This equation is valid for any electric field strength.

$$\text{Energy density} = \frac{\text{Energy}}{\text{Volume}} = \frac{1}{2} \epsilon_0\epsilon_r E^2 \dots \dots \dots (16.42)$$

Example 16.7 Two capacitors of capacitances $4 \mu\text{F}$ and $8 \mu\text{F}$ are connected in series and charged by a 120 V power supply. Find the total charge and energy stored in the system.

Solution Equivalent capacitance C_e is given by $\frac{1}{C_e} = \frac{1}{C_1} + \frac{1}{C_2}$

$$\frac{1}{C_e} = \frac{1}{4 \times 10^{-6} \text{ F}} + \frac{1}{8 \times 10^{-6} \text{ F}}$$

Putting the values of C_1 and C_2 , we have

$$C_e = 2.67 \times 10^{-6} \text{ F}$$

and

$$Q = C_e V = (2.67 \times 10^{-6} \text{ F}) \times 120 \text{ V}$$

$$Q = 3.2 \times 10^{-4} \text{ C}$$

For your information



A device called as defibrillator uses the charge stored in a capacitor to deliver a controlled electric shock that can restore normal heart rhythm.

Energy stored in the system is given by

$$W = \frac{1}{2} QV$$

$$W = \frac{1}{2} (3.2 \times 10^{-4} \text{ C}) \times 120 \text{ V}$$

$$= 1.92 \times 10^{-2} \text{ J}$$

Example 16.8 A $16 \mu\text{F}$ capacitor is charged to a potential difference of 100 V . Determine the energy stored in it. If the potential difference is doubled by what factor does the energy increase?

Solution

Energy stored is given by

$$W = \frac{1}{2} CV^2$$

Putting the values of C and V , we have

$$W = \frac{1}{2} (16 \times 10^{-6} \text{ F}) \times (100 \text{ V})^2$$

$$= 8 \times 10^{-2} \text{ J}$$

Since energy depends on V^2 , therefore, if V is doubled, the energy increases by a factor of 4.

16.8 CHARGING AND DISCHARGING A CAPACITOR

Many electric circuits consist of both capacitors and resistors. Figure 16.11 shows a resistor-capacitor circuit called RC-circuit. When the switch S is set at terminal A , the R-C combination is connected to a battery of voltage V_0 which starts charging the capacitor through the resistor R .

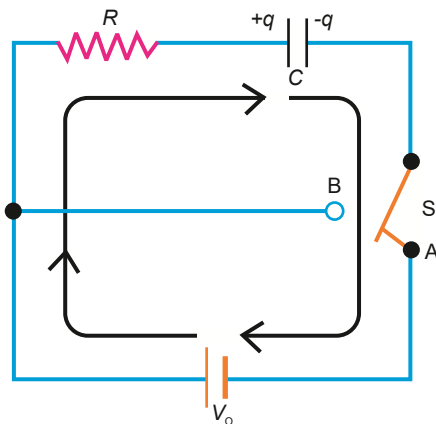


Fig. 16.11: Charging of a capacitor

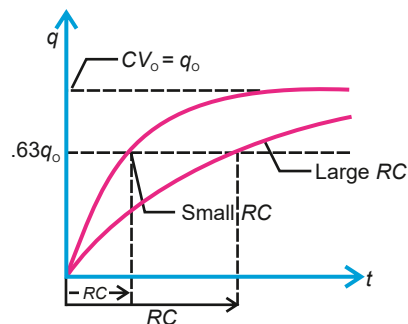


Fig. 16.12: Growth of charge through different combination of RC

The capacitor is not charged immediately, rather charges build up gradually to the equilibrium value of $q_0 = CV_0$. The growth of charge with time for different resistances is shown in Fig.16.12. According to this graph, $q = 0$ at $t = 0$ and increases gradually with time till it reaches its equilibrium value $q_0 = CV_0$. The voltage V across capacitor at any instant can be obtained by dividing q by C as: $V = q/C$.

How fast or how slow the capacitor is charging or discharging, depends upon the product of the resistance R and the capacitance C used in the circuit. As the unit of product RC is that of time, so this product is known as time constant and is defined as:

The time required by the capacitor to deposit 0.63 times the equilibrium charge q_0 .

The graphs of Fig.16.12 show that the charge reaches its equilibrium value sooner when the time constant is small.

Figure 16.13 illustrates the **discharging** of a capacitor through a resistor. In this figure, the switch S is set at point B , so the charge $+q$ on the left plate can flow anticlockwise through the resistance and neutralize the charge $-q$ on the right plate.

The graph in Fig.16.14 shows that discharging begins at $t = 0$ when $q = CV_0$ and decreases gradually to zero. Smaller values of time constant RC lead to a more rapid discharge.

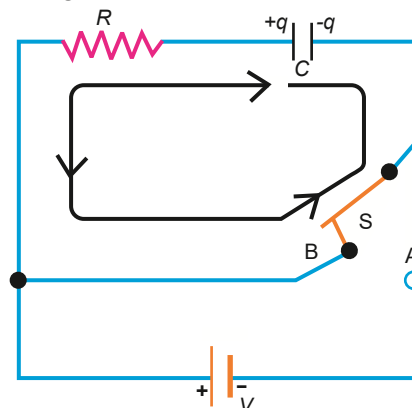


Fig. 16.13: Discharging of capacitor through resistance R

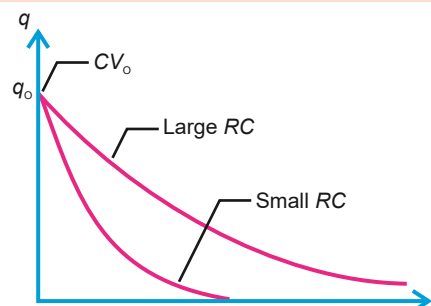


Fig. 16.14: Discharging through different combination of RC

Example 16.9 The time constant of a series RC -circuit is $t = RC$. Verify that an ohm times farad is equivalent to second.

Solution Ohm's law in terms of potential difference V , current and resistance R can be written as,

$$V = IR$$

Putting $I = q/t$, this equation transforms into the equation:

$$V = \frac{q}{t} R$$

$$\text{or } R = \frac{V \times t}{q}$$

According to equation; $q = CV$ or $C = q/V$

Multiplying this equation with above equation gives:

$$RC = \frac{V \times t}{q} \times \frac{q}{V} = t$$

16.9 EXPONENTIAL DECAY OF DISCHARGING CAPACITOR THROUGH RESISTOR

An RC-circuit for the discharging capacitor through a resistor was shown in Fig. 16.13. A graph, charge vs time for that is shown in Fig. 16.15. The charge Q decreases exponentially overtime t . In general, exponential means a process whose rate of change at any instant is proportional to the present value of the quantity, so it changes quickly at first and then more slowly.

The general exponential decay formula for a discharging capacitor in an RC-circuit is of the form:

$$x = x_0 \left[\exp \left(\frac{-t}{RC} \right) \right] \dots\dots\dots (16.43)$$

or $x = x_0 e^{\frac{-t}{RC}}$ where,

x = quantity (such as charge, voltage or current) at time t

x_0 = initial value of the quantity at $t = 0$

R = resistance in ohms (Ω)

C = capacitance in farads (F)

t = time in seconds

RC = time constant (τ) which determines how quickly the capacitor discharges.

e = In this equation e is a mathematical constant called Euler's constant which is also used as base of the natural logarithm ($\ln$). Its value is approximately 2.718.

The negative exponent ($-t/RC$) in the formula indicates that the quantity decreases exponentially over time.

Using Eq.16.43, the decrease in charge, voltage and current at any time t , can be described by the following equations:

$$Q = Q_0 e^{-t/RC}, V = V_0 e^{-t/RC} \text{ and } I = I_0 e^{-t/RC}.$$

In the graph for the decrease in all the three quantities, they start with a steep slope (fast discharge) and gradually flattens out as the time progresses, indicating a slower discharge rate. This is shown in Fig. 16.16 by the charge decay graph.

The initial charge is Q_0 at $t = 0$ and it drops to Q after a time interval t . Now, if at another stage during the discharging of capacitor, let the charge be Q' . After

the same interval of time t , it drops to Q' . We can see that the rate of decrease of charge

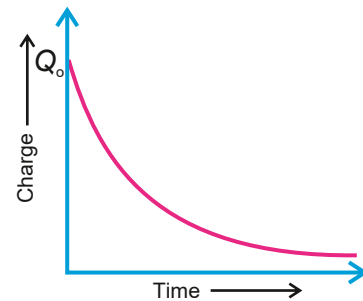


Fig. 16.15: Decay of charge with time

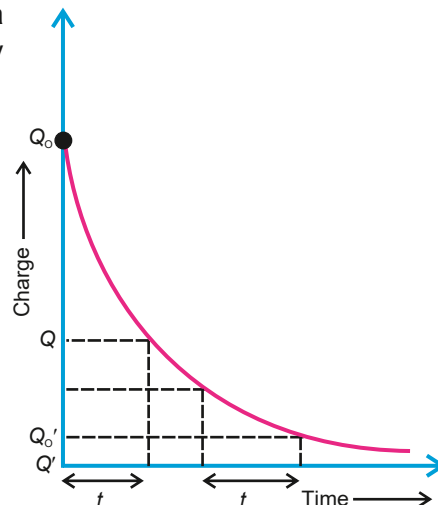


Fig. 16.16: Charge Decay graph

in the beginning was $(Q_0 - Q)$, but as the process goes on, the rate of decrease of charge becomes $(Q'_0 - Q')$ which is less than $(Q_0 - Q)$.

The decay of charge is very much similar to flow of water out of a tank through a narrow pipe. When the tank is full of water, its level and pressure is analogous to charge and voltage, and flow rate of water analogous to the current. In the beginning, the water level is high and it decreases rapidly and the flow of water through the pipe is fast. As the water level comes down, the rate of flow slows down. This flowing down is due to the decreasing pressure.

For your information

Exponential functions frequently and quantitatively describe a number of phenomena in physics, such as radioactive decay, in which the rate of change in a process or substance depends directly on its current value.

Role of Time Constant

The time constant $RC = \tau$ plays very important role in charging and discharge a capacitor through a resistance. Consider the exponential decay of the charge stored in a capacitor. It is given by the equation:

$$Q = Q_0 e^{-t/RC}$$

If we put $t = RC$, we will get the quantity of charge remained after 1 time constant.

Therefore, $Q = Q_0 e^{-RC/RC} = Q_0 e^{-1} = Q_0 \times 1/e$

Substituting the value of e , we have

$$Q = Q_0 \times \frac{1}{2.718} = 0.37Q_0$$

This gives the definition of time constant:

The time constant is the duration of time in which, the initial charge Q_0 drops to about $0.37 Q_0$.

This means that after one time constant, the charge drops to 37% of its initial value. Likewise the voltage and current also drops to the same extent. The exponential decay equation also shows that a smaller time constant ($\tau = RC$) means the circuit responds

faster, while a larger time constant means, it responds more slowly. This effect is shown in Fig. 16.17. Let us now determine $(t_{1/2})$ taken in terms of time constant RC during which the charge Q_0 drops to one half of its initial value. Putting the value of $Q = 1/2 Q_0$ in the exponential decay equation of the charge, we have

$$\frac{1}{2} Q_0 = Q_0 e^{-t_{1/2}/RC}$$

or

$$\frac{1}{2} = e^{-t_{1/2}/RC}$$

By definition of natural logarithm,

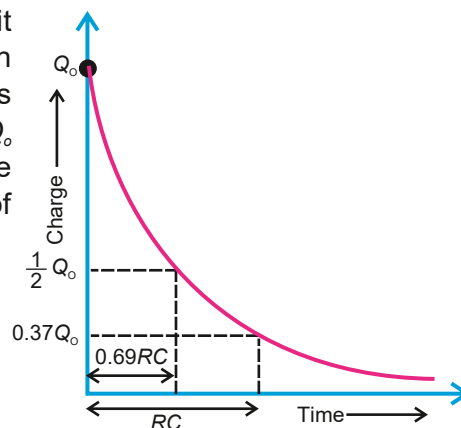


Fig. 16.17: Effect of RC on decay time

$$\ln\left(\frac{1}{2}\right) = \frac{-t_{1/2}}{RC}$$

$$\text{or} \quad -(RC) \left[\ln\left(\frac{1}{2}\right) \right] = t_{1/2} \dots \dots \dots (16.44)$$

$$\begin{aligned} \text{Now} \quad \ln\left(\frac{1}{2}\right) &= \ln(1) - \ln(2) \\ &= 0 - \ln(2) \end{aligned}$$

From calculator, the value of $\ln(2) = -0.69$

Putting the value of $\ln(1/2)$ in Eq. 16.44, we have

$$-RC(-0.69) = t_{1/2}$$

$$\text{or} \quad t_{1/2} = 0.69 RC \dots \dots \dots (16.45)$$

we can show this in the graph of Fig. 16.17.

Example 16.10 In an RC-circuit $R = 1 \text{ k}\Omega$, $C = 100 \text{ }\mu\text{F}$ and $V_o = 10 \text{ V}$. Determine the voltage across the capacitor after 0.2 s.

Solution

$$\begin{aligned} \text{The constant} \quad t &= RC = 1000 \Omega \times 100 \times 10^{-6} \text{ F} \\ &= 0.1 \text{ s} \end{aligned}$$

$$\text{Voltage after } 0.2 \text{ s} \quad = V(t) = ?$$

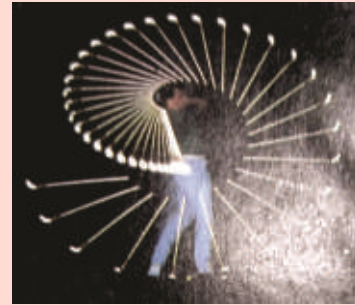
Using exponential decay equation for voltage:

$$V = V_o e^{\frac{-t}{RC}}$$

$$V = 10 \times e^{\frac{-0.2}{0.1}} = 10 \times e^{-2} = 10 \times \frac{1}{e^2}$$

$$\text{or} \quad V = 10 \times \frac{1}{(2.718)^2} = 1.35 \text{ V}$$

Interesting Information



This time-lapse photo of a golfer teeing off was obtained by using an electronic flash attachment with the camera. The energy for each “flash” comes from the electrical energy stored in a capacitor.

16.10 USE OF CAPACITORS

Capacitors are used in various household appliances such as electric fans, water drawing motors (Fig. 16.18), washing machines, refrigerators, flashguns, computer keyboards, touchscreens, etc. We will describe here a few of them with some detail.

Electric Fans and Electric Motors

A capacitor is connected to the starting of fan motor. When the power is turned ON, the capacitor starts charging, i.e., it accumulates



Fig. 16.18: Use of capacitors in electric fan and electric motor

charge. During the process of starting of fan motor, it provides additional power, helping the motor to overcome static friction and initial resistance due to inertia. In thin way, it enables the fan motor to start and begin spinning quickly. The same role is played by the capacitor in motors used in washing machines, electric blenders, choppers, grinders, water drawing motors, etc.

Refrigerators and Air-Conditioners

The compressors of a refrigerator and an air-conditioner typically use single phase motors, which can struggle to start on their own. The capacitors of high capacitance are used in them. Such a capacitor stores a big amount of electrical energy and releases it to the motor winding of the compressor. It acts like a “jump start” for the motor, providing a boost of power to get the motor rotating (Fig. 16.19).

Once the motor reaches a certain speed, a relay disconnects the capacitor as it is no longer needed. Commonly, this capacitor is called as “start capacitor.” Most refrigerators and air-conditioners also use a “run capacitor” which stays connected in the circuit to improve energy efficiency. However, some older models may use “start capacitors” that only function during start up.

Capacitors in conjunction with resistors can form a time delay circuit. It is widely used in refrigerators, air conditioners and other devices to create delay in triggering the power supply. This is a precautionary measure to protect the devices from power fluctuation.



Fig. 16.19: Compressor in Air-conditioner and refrigerator

Flashgun

In a flashgun, a capacitor stores and releases electrical energy to power the flash tube. A flash tube contains xenon gas in it. The capacitor accumulates charge from the battery of camera or some other device. When the shutter button is pressed, the capacitor is connected to the flash tube. The rapid discharge of the capacitor provides the high current needed to ionize the xenon gas inside the tube, causing it to emit a burst of bright light (Fig. 16.20).



Fig. 16.20: Flashgun in cameras

Filters as in Rectifier Circuits

Capacitors are widely used in the rectifier circuits for the smoothing of pulsating DC into smooth DC.

There is a vast use of capacitors in the medical field. Capacitors store electrical energy needed to deliver a life-saving shock to heart, restoring a normal rhythm in case of cardiac arrest. Capacitors are also used in imaging machines (CT, MRI, X-Ray, Ultrasound, etc.) helping to filter out noise, and improve image clarity. Similarly, capacitors filter out unwanted signals in EEG and ECG machines.

QUESTIONS

Multiple Choice Questions

Choose the correct answer.

- 16.1 What is the work done on an electron by a potential difference of 100 volts?
(a) 1.6×10^{-19} eV (b) 1.6×10^{-17} eV (c) 6.25×10^{-17} eV (d) 100 eV
- 16.2 A parallel plate capacitor of capacitance $50 \mu\text{F}$ is connected to a battery of 12 volts, the charge stored on each plate is:
(a) $60 \mu\text{C}$ (b) $4.2 \mu\text{C}$ (c) 6×10^{-4} C (d) 4×10^4 C
- 16.3 If a dielectric is inserted between the plates of a charged capacitor, the potential difference across the plates:
(a) increases (b) decreases (c) remains same (d) becomes zero
- 16.4 The potential energy of a system of two charges $8 \mu\text{C}$ and $-8 \mu\text{C}$ separated by a distance 2 m is:
(a) $\frac{8 \times 10^{-12}}{\pi \epsilon_0}$ J (b) $\frac{+16 \times 10^{-12}}{\pi \epsilon_0}$ J
(c) $\frac{-8 \times 10^{-12}}{\pi \epsilon_0}$ J (d) $\frac{-16 \times 10^{-12}}{\pi \epsilon_0}$ J
- 16.5 The electric potential at point A is 50 V and at B is 80 V. If the electric field between the charges is uniform, what can be said about the motion of positive charge from A to B?
(a) No work is done
(b) Work is done by the electric field
(c) Work is done against the electric field
(d) None of the above
- 16.6 The rate of charging or discharging the capacitor depends upon the product of resistance and:
(a) charge (b) potential difference (c) current (d) capacitance
- 16.7 In a charged capacitor, the energy resides in:
(a) the positive plate (b) the negative plate
(c) both positive and negative plates (d) the electric field between the plates
- 16.8 A capacitor of capacitance $1 \mu\text{F}$ is subjected to 4000 volts potential difference, the energy stored in it is:
(a) 8 J (b) 16 J
(c) 32 J (d) 64 J

16.9 The electric potential energy of a point charge q at some point in an electric field is equal to:

- (a) $\frac{1}{4\pi\epsilon_0} \frac{q}{r}$ (b) $\frac{1}{4\pi\epsilon_0} \frac{q}{r^2}$ (c) qV (d) V^2

16.10 The time constant of a capacitor discharging through a resistor is the time in seconds in which the initial charge Q_0 drops to:

- (a) $0.25 Q_0$ (b) $0.37 Q_0$ (c) $0.50 Q_0$ (d) $0.69 Q_0$

Short Answer Questions

- 16.1 If the absolute potential at a point is zero, what can you say about the electric intensity at that point?
- 16.2 Is electron volt, a unit of potential difference or energy? Explain briefly.
- 16.3 Define capacitance and write its units.
- 16.4 Define dielectric constant. Give its mathematical form.
- 16.5 How does the capacitance of a capacitor change if the separation between its plates is increased?
- 16.6 When will the energy stored in three capacitors be greater?
- (i) They are connected in series.
- (ii) They are connected in parallel.
- 16.7 What is time constant? Write its formula.

Constructed Response Questions

- 16.1 How is electric potential different from electric potential energy? Give an example to explain the difference.
- 16.2 How does the concept of electric potential help to understand the motion of charges in an electric field?
- 16.3 The electric field is negative of gradient of potential. What does the negative sign indicate physically? Illustrate your answer with an example.
- 16.4 When a capacitor is charging, no actual charge flows across the gap between the plates. But the current is still said to be flowing. Explain.
- 16.5 Two capacitors can be connected in parallel or in series. Explain how does total energy stored differ in two configurations.
- 16.6 A capacitor stores energy. Is this energy stored on the plates, in the dielectric or in the electric field? Explain how does a capacitor store energy.
- 16.7 Why does the capacitance increase on inserting a dielectric between the plates of a capacitor? Explain.

Comprehensive Questions

- 16.1 Derive an expression for the absolute electric potential at a point due to a point charge q .
- 16.2 What is a capacitor? Find the capacitance of a parallel plate capacitor. What is the role of a dielectric in the capacitor?
- 16.3 Derive the formula for series combinations of capacitors.
- 16.4 Determine the electric potential energy stored in a capacitor from the area under the potential-charge graph.
- 16.5 Explain the charging and discharging of a capacitor by drawing circuit diagrams and graphs between charge versus time. What is a time constant for a capacitor discharging through a resistor?
- 16.6 Describe some uses of capacitors in various household appliances.

Numerical Problems

- 16.1 A particle carrying a charge of $0.1 \mu\text{C}$ starts from rest in a uniform electric field of intensity 60 N C^{-1} . Find the kinetic energy acquired by the particle, when it has moved 2 m . (Ans: $1.2 \times 10^{-5} \text{ J}$)
- 16.2 What should be the potential difference through which an electron falls to acquire a speed of $2 \times 10^6 \text{ m s}^{-1}$? (Ans: 11.4 V)
- 16.3 The potential difference between points A and B is $V_A - V_B = 84 \text{ V}$, while the potential difference between points C and B is $V_C - V_B = 24 \text{ V}$. What should be the work done on moving a $40 \mu\text{C}$ charge from C to A? (Ans: $2.4 \times 10^{-3} \text{ J}$)
- 16.4 A parallel plate capacitor has a capacitance of $6 \mu\text{F}$ when filled with a dielectric. The area of each plate is 1.2 m^2 and the separation between the plates is $9 \times 10^{-6} \text{ m}$. What is the dielectric constant? (Ans: 5.1)
- 16.5 A $15 \mu\text{F}$ capacitor is charged till the potential difference across its plates becomes 300 V . Find the energy stored in it. (Ans: 0.675 J)
- 16.6 Two capacitors of $2.0 \mu\text{F}$ and $8.0 \mu\text{F}$ capacitances are connected in series and a potential difference of 300 volts is applied. Find the charge and potential difference for each capacitor. (Ans: $Q_1 = Q_2 = 4.8 \times 10^{-4} \text{ C}$, $V_1 = 240 \text{ V}$, $V_2 = 60 \text{ V}$)
- 16.7 A capacitor has a capacitance of $3 \mu\text{F}$. It is charged by a battery. When the potential difference between the plates is 240 volts , how many electrons have been removed from one plate and placed on the other plate? (Ans: $4.5 \times 10^{15} \text{ electrons}$)
- 16.8 Two positive point charges are held at two points 1.2 m apart. They are then moved towards each other so that their electric potential energy becomes double of its previous value. How much are they apart at their new positions? (Ans: 0.6 m)
- 16.9 A $100 \mu\text{F}$ capacitor is initially charged to 8 V . It is allowed to discharge through a 500Ω resistor. Find the voltage across the capacitor after 0.5 s . (Ans: $3.6 \times 10^{-4} \text{ V}$)

Alternating Current

Students Learning Outcomes

After studying this chapter, the students will be able to:

- ◆ analyze equations of the form $x = x_o \sin(\omega t)$ representing a sinusoidally alternating current or voltage.
- ◆ understand and describe the terms period, frequency and the peak value as applied to an alternating current or voltage.
- ◆ differentiate between root-mean-square (rms) and peak values [using $I_{\text{rms}} = \frac{I_o}{\sqrt{2}}$ and $V_{\text{rms}} = \frac{V_o}{\sqrt{2}}$ for a sinusoidal alternating current and voltage].
- ◆ know about the fact that the mean power in a resistive load is half the maximum power for a sinusoidal alternating current.
- ◆ understand the phase of AC and how phase lags and leads in AC circuits.
- ◆ describe impedance – the unseen force as vector summation of resistance of resistors and reactance of capacitors and inductors.
- ◆ apply the knowledge to calculate the reactance of capacitors and inductors.
- ◆ identify inductors as important components of AC circuits termed as chokes [devices which present a high resistance to alternating current].
- ◆ distinguish graphically between half-wave and full-wave rectification.
- ◆ explain the use of a single diode for the half-wave rectification of an alternating current.
- ◆ explain the use of four diodes (Bridge Rectifier) for the full-wave rectification of an alternating current.
- ◆ describe the effect of a single capacitor in smoothing current flow [including the effect of the values of capacitance and the load resistance]
- ◆ state and prove expressions for mutual inductance (M) and self-inductance (L), their unit henry.

17.1 ALTERNATING CURRENT AND ITS CHARACTERISTICS

An alternating current AC means a sinusoidally varying current which can be represented by time dependent sine or cosine functions. The output voltage of an AC generator is given by the following equation:

$$V = V_o \sin \omega t \dots\dots\dots (17.1)$$

where V represents magnitude of alternating voltage corresponding to time t , V_o represents maximum value of voltage and ω is the angular frequency of alternating voltage as shown in Fig 17.1. We can write Eq.17.1 for current as:

$$I = I_o \sin \omega t \dots\dots\dots (17.1-a)$$

where i is instantaneous value of alternating current at time t , i_0 is its maximum or peak value and ω is its angular frequency.

A periodically varying current or voltage is termed as alternating if every cycle occupies one time period and has two symmetrical half cycles, one positive and the other negative.

In case of current, the direction of current reverses after every half-cycle and in case of voltages, the polarity of potential difference reverses after every half cycle. The second condition to be an alternating one is that amplitude or the peak value of current or voltage i.e., the maximum value on both positive and negative sides remains constant in all cycles but changes occur most rapidly at the zero (crossover) points and most slowly at its peak (Fig. 17.1).

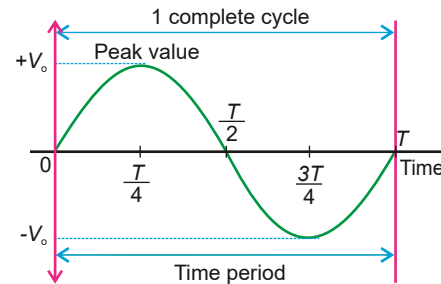


Fig.17.1: Waveform of alternating voltage

Some important terms related to alternating quantities are:

Waveform: The path traced by an alternating quantity, such as the voltage in Fig.17.1 plotted as a function of time.

Cycle: One complete set of positive and negative values of an alternating quantity is called a cycle. Figure 17.1 shows one cycle of an alternating voltage.

Time period (T): The time taken to complete one cycle of an alternating quantity is called its time period and is measured in seconds.

Frequency (f): The number of cycles that occur in one second is called its frequency. The unit of frequency is hertz (Hz); $f = 1/T$.

Average or Mean Value of AC Waveform

The average value of AC is obtained by averaging all instantaneous values over one half cycle.

If we calculate the average of a complete sine wave over a full cycle (0 to 2π), the result is zero (Fig. 17.2). This happens because the positive half-cycle and the negative half-cycle are mirror images of each other. They cancel each other. So, when we talk about the average value of an AC waveform, we usually mean the average over **one half-cycle** (0 to π).

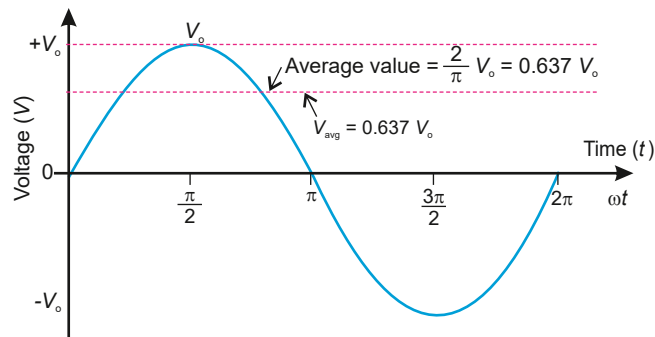


Fig.17.2: Average value of alternating voltage

The average value is used in situations involving rectification. When an AC signal passes through a rectifier (such as in a power supply), the output is a pulsating DC signal. The average of this pulsating DC is what matters for many applications. For a full-wave rectified sine wave, the average values of current and voltage are given by

$$I_{av} = 0.637 I_o \quad \text{and} \quad V_{av} = 0.637 V_o$$

where I_{av} and V_{av} represent average value of alternating current and alternating voltage respectively whereas I_o and V_o represents maximum or peak values of alternating current and voltage respectively (Fig. 17.3).

Instantaneous Value: The value of an AC quantity at any instant of time is known as instantaneous value. The value of instantaneous voltage is given as:

$$V = V_o \sin(\omega t)$$

$$\begin{aligned} \text{As } \omega &= \frac{2\pi}{T}, \text{ therefore} \\ &= 2\pi f \quad (\because f = \frac{1}{T}) \end{aligned}$$

$$\text{So } V = V_o \sin(2\pi ft) \dots\dots(17.2-a)$$

The value of instantaneous current is:

$$I = I_o \sin(\omega t)$$

$$\text{or } I = I_o \sin(2\pi ft) \dots\dots(17.2-b)$$

Peak or Maximum Value: It is the highest value that an AC waveform reaches during one cycle. It is the topmost point on the positive half-cycle or the bottommost point on the negative half-cycle (Figs. 17.1 and 17.2). It is written by V_o and I_o for voltage and current respectively.

Peak to Peak Value: This is the total distance from the positive peak to the negative peak of the waveform. For a symmetrical sine wave:

$$V_{p-p} = 2V_o$$

The p-p value of the voltage waveform shown in Fig. 17.3 is $2V_o$.

Root-mean-square (rms) or Effective Value of Current and Voltage:

Root-mean-square (rms) values of current, or voltage, are a useful way of comparing alternating current, or voltage, to its equivalent direct current, or voltage. The rms values

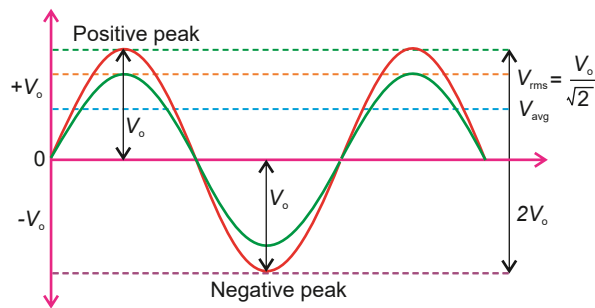


Fig.17.3: Graphically showing the rms value, average or mean value, peak value, and p-p value of alternating quantities-voltage.

Brain teaser

How can we detect the presence of AC under a carpet?

we can use a non-contact voltage detector, known as a pen tester or a multimeter with AC voltage setting. Non-contact detectors will sense the electromagnetic field generated by the current and indicate its presence with a light or sound. A multi-meter can provide a more precise measurement by directly detecting the voltage.

represent the direct current or voltage values, that will produce the same heating effect, or power dissipation as the alternating current or voltage. The rms value of an alternating current is defined as:

The value of a constant current that produces the same power in a resistor as the alternating current.

$$I_{\text{rms}} = \frac{I_0}{\sqrt{2}} \quad \text{or} \quad I_{\text{rms}} = 0.707I_0 \quad \dots\dots\dots (17.3)$$

The rms value of an alternating voltage is defined as:

The value of a constant voltage that produces the same power in a resistor as the alternating voltage.

$$V_{\text{rms}} = \frac{V_0}{\sqrt{2}} \quad \text{or} \quad V_{\text{rms}} = 0.707V_0 \quad \dots\dots\dots (17.4)$$

where I_0 = Peak or maximum current

V_0 = Peak or maximum voltage

Point to ponder!

Why do high voltage power lines crackle and hiss?

The hissing sound often heard near high voltage power lines is primarily caused by corona discharge, an electrical phenomenon where the air surrounding the conductor becomes ionized due to a strong electric field. This ionization creates a discharge that can produce a visible glow, radio noise, and audible hissing or crackling sounds, particularly when the voltage exceeds the breakdown strength of the air.

17.2 RELATION BETWEEN MEAN POWER AND MAXIMUM POWER FOR AN ALTERNATING CURRENT

An alternating current is that periodically reverses direction, causing the voltage and current to fluctuate sinusoidally. A resistive load is one that does not store energy (like a capacitor or inductor) and simply dissipates energy as heat.

The instantaneous power P in a circuit is the power at any given moment in time t .

$$P = VI \quad \dots\dots\dots (17.5)$$

where V is the instantaneous voltage and I is the instantaneous current. The maximum or peak P_{max} occurs when the current and voltage are at their peak values, I_0 and V_0 respectively. Also the mean power P_{mean} is the average power over one complete cycle of the AC waveform as shown in Fig.17.4.

For a sinusoidal AC, let I_0 be the maximum or peak value of the current while voltage V_0 is the maximum or peak value of the voltage, then the maximum or peak power P_{max} is given by

$$P_{\text{max}} = I_0^2 R \quad \dots\dots\dots (17.6)$$

Also we have expression for the mean power P_{mean} represented by

$$P_{\text{mean}} = I_{\text{rms}}^2 R \quad \dots\dots\dots (17.7)$$

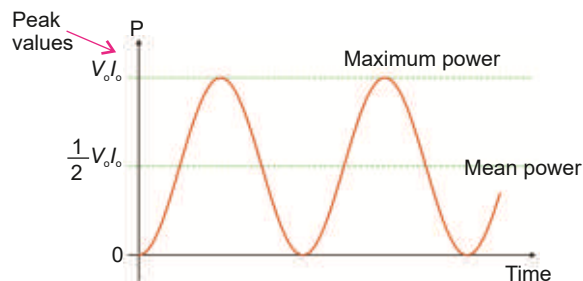


Fig.17.4: Mean power is exactly half the maximum power

As $I_o = \sqrt{2} I_{\text{rms}}$, so, using Eq. (17.6):

$$P_{\text{max}} = (\sqrt{2} I_{\text{rms}})^2 R = 2 I_{\text{rms}}^2 R$$

From Eq. 17.7 ; $I_{\text{rms}}^2 R = P_{\text{mean}}$, thus,

$$P_{\text{max}} = 2P_{\text{mean}}$$

or $P_{\text{mean}} = \frac{P_{\text{max}}}{2} \dots\dots\dots (17.8)$

Therefore, it can be concluded that the mean power in a resistive load is half the maximum power for a sinusoidal alternating current or voltage because the instantaneous power in an AC circuit varies over time. While the instantaneous power reaches at maximum value at certain points in the cycle, it spends significant time at zero or lower values, resulting in an average power that is half the peak value.

Example 17.1 An alternating voltage supplied across a resistor of 50Ω has a voltage of 220 V. Calculate the mean power in kilowatt of the system.

Solution $R = 50 \Omega$ and $V_{\text{rms}} = 220 \text{ V}$

We know that; $P_{\text{max}} = \frac{V_{\text{rms}}^2}{R}$

Putting the values $P_{\text{max}} = \frac{(220 \text{ V})^2}{50 \Omega} = 968 \text{ W}$

As the mean power is half of the maximum (peak) power,

$$P_{\text{mean}} = \frac{P_{\text{max}}}{2} = \frac{968 \text{ W}}{2} = 484 \text{ W}$$

Interesting Information



Digital Clamp Meter For Measuring AC and DC

Fascinating Fact



Scientists have converted human blood sugar into electricity

17.3 PHASE OF A.C.

The angle θ which specifies the instantaneous value of the alternating voltage or current is called the phase.

The instantaneous values of voltage and current are given by

$$V = V_o \sin \omega t \quad \text{and}$$

$$I = I_o \sin \omega t$$

where $\omega t = \theta$.

The phase of points A, B, C, D and E are 0, $\pi/2$, π , $3\pi/2$ and 2π respectively as shown in Fig.17.5(a). Thus, each point of A.C cycle corresponds to a phase.

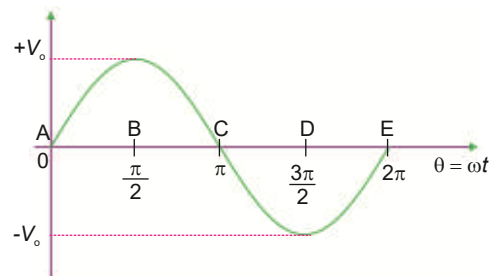


Fig.17.5(a): AC cycle corresponds to a phase.

Phase Lag and Phase Lead

Referring to the Figs. 17.5 (b) and 17.5(c), at $t=0$, the angle θ is also zero, the value of the phase angle θ at t equals to zero is called the initial phase of AC quantity. There are situations when current I and voltage V are not in phase i.e; they differ in phase. For example, the initial phase of current I may be positive or negative as compared to the initial phase of voltage V , which is zero. Consider a situation in which the initial phase of voltage V is zero and initial phase of current I is ϕ as shown in Fig. 17.5 (b). This situation can be represented by the following equations:

$$V = V_o \sin \theta \quad \text{and} \quad I = I_o \sin(\theta + \phi) \quad \dots\dots\dots (17.9-a)$$

At $t = 0$; voltage $V = 0$, but the current is positive and $I = I_o \sin \phi$. It means that the current had its zero value earlier by an angle ϕ than voltage. Thus, we can say that the current is leading the voltage by an angle ϕ in this situation (Fig. 17.5-b). The angle ϕ is the phase difference between the voltage and the current.

Similarly, the initial phase of current I is negative as compared to the initial phase of voltage V , which is zero as shown in Fig. 17.5(c). This situation can be represented by the following equations:

$$V = V_o \sin \theta \quad \text{and} \quad I = I_o \sin(\theta - \phi) \quad \dots\dots\dots (17.9-b)$$

It means the value of V is zero, but the current I has some negative value. It means that the current I will reach its zero value later on by an angle ϕ than the voltage V . Thus, we can say that current I is lagging behind the voltage V by an angle ϕ . The angle ϕ is called the phase difference between V and I .

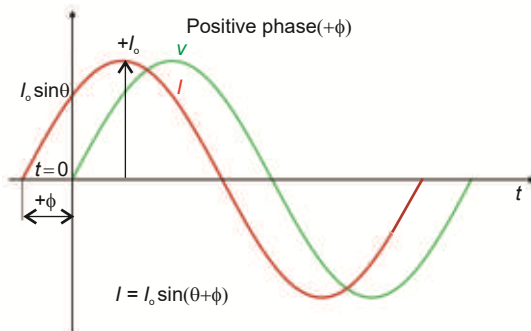


Fig.17.5(b): Graphical representation showing current leads the voltage.

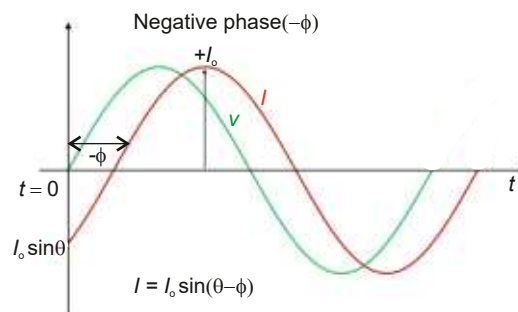


Fig.17.5(c): Graphical representation showing current lags behind the voltage.

Vector Representation of an Alternating Quantity

An alternating quantity can be represented by an anticlockwise rotating vector if it satisfies the conditions: (i) Its length on a certain scale represents the peak value or rms value of alternating quantity. (ii) In horizontal position at the instant, the alternating quantity is zero and is increasing positively (iii) The angular frequency of the rotating vector is same as the angular frequency of alternating quantity (Fig.17.5-d). shows an alternating voltage waveform leading the current wave form by 90° or $\pi / 2$ rad. In

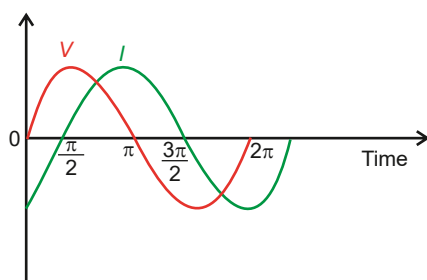


Fig.17.5 (d): Graphical representation showing voltage V leads current I by $\pi/2$

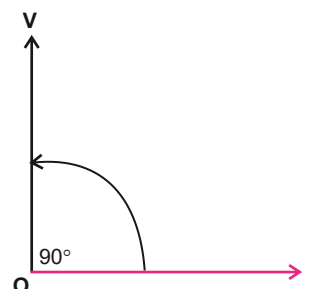


Fig. 17.5 (e): Vector diagram $\mathbf{OI}$ and $\mathbf{OV}$ showing V leads I by 90° or $\pi/2$ rad

Fig.17.5(e), vector $\mathbf{OI}$ represents peak or rms value of current which is taken as reference quantity. Similarly, $\mathbf{OV}$ represents peak or rms value of alternating voltage, which is leading the current by $\pi/2$ rad or 90° . Both vectors are rotating in anticlockwise direction with angular frequency.

17.4 AC CIRCUITS

The basic circuit element in a DC circuit is a resistor R which controls the current or voltage and the relation between them is given by Ohm's law ($V = IR$). But the basic circuit elements in AC circuit are resistor R , inductor L and capacitor C . These elements control the current and voltage through the circuit. The AC circuits with these components are discussed below:

17.5 AC THROUGH A RESISTOR

A resistor R connected with an AC source is shown in Fig.17.6 (a). The instantaneous voltage V is given as:

$$V = V_o \sin \omega t \quad \dots\dots\dots 17.10$$

The instantaneous current I through the circuit is:

$$I = \frac{V}{R}$$

Using Eq.17.10; $I = \frac{V_o \sin \omega t}{R}$

$$I = I_o \sin \omega t \quad \dots\dots\dots 17.11$$

where $I_o = \frac{V_o}{R}$ is the peak or maximum value of current.

It follows from Eqs.(17.10) and (17.11) that the instantaneous values of both voltage and current are sine functions which vary with time. Figure 17.6(b) shows that both voltage and current pass their minimum and maximum values at the same time and

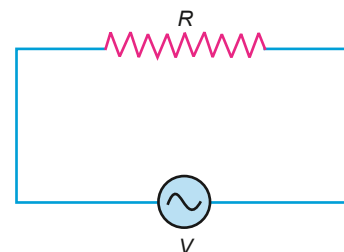


Fig.17.6(a): Showing a resistor R connected with an AC source

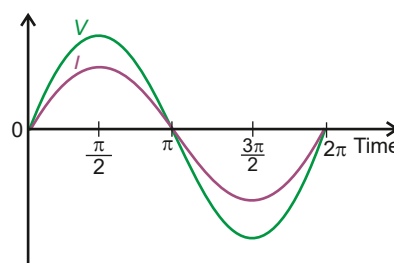


Fig. 17.6(b): Graphical representation for purely resistive circuit showing V and I in phase.

thus their instantaneous values are said to be in phase with each other.

Also Fig.17.6(c) shows V_R and I_R vectors for resistance. V and I are drawn parallel because there is no phase difference between them. The opposition to AC which the circuit presents is the resistance given by

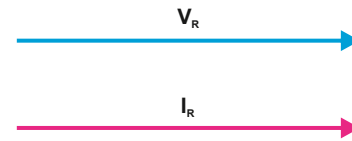
$$R = \frac{V}{I}$$


Fig.17.6(c): V_R and I_R vectors representation for resistance.

The power dissipation is proportional to the square of the current and makes no difference whether the current is direct or alternating i.e., whether the sign associated with the current is positive or negative. However, the power dissipation produced by AC having maximum value I_o is not the same as that produced by a direct current of maximum value I_o , because the alternating current is at this maximum value only for an instant during each half-cycle. Instantaneous power P dissipated across a resistor in AC circuit is:

$$P = VI \text{ or } P = I^2 R \text{ or } P = \frac{V^2}{R} \dots\dots\dots (17.12)$$

and average power $\langle P \rangle = \langle I^2 R \rangle$

$$\text{As } \langle I^2 \rangle = \frac{I_o^2}{2} = I_{\text{rms}}^2, \text{ so}$$

$$\text{Average power } \langle P \rangle = \frac{I_o^2}{2} R = I_{\text{rms}}^2 R$$

Also, it can be proved that;

$$\text{Average power } \langle P \rangle = I_{\text{rms}} V_{\text{rms}}$$

It is to be noted that here P is measured in watts V in volts, I in amperes and R in ohms respectively, and the Eq.(17.12) for power holds good only when V and I are in phase.

For your information

12 V AC Doorbell



Low voltage AC is safer to handle and reduces the risk of electrical shock.

Example 17.2 A 1 kW heating element is connected to a 250 V AC supply voltage. Calculate the amount of current taken from the supply and the resistance of the element when it is hot.

Solution Power $P = 1 \text{ kW} = 1000 \text{ W}$ and Applied voltage $V = 250 \text{ V}$

We know that; $P = VI$ or $I = \frac{P}{V}$

$$I = \frac{1000 \text{ W}}{250 \text{ V}} = 4 \text{ A}$$

As $R = \frac{V}{I}$, then

$$R = \frac{250 \text{ V}}{4 \text{ A}}$$

$$R = 62.5 \Omega$$

17.6 AC THROUGH INDUCTOR

An inductor, also called a coil, or choke is a passive two terminal electrical component having a large value of self-inductance and negligible resistance that stores energy in a magnetic field when electric current flows through it. The inductor is used to slow down current surges or spikes by temporarily storing energy in an electromagnetic field and then releasing it back into the circuit.

Suppose an inductor of inductance L is connected with an AC source of frequency f with a negligible resistance as shown in Fig.17.7 (a). Suppose the current is:

$$I = I_0 \sin \omega t \quad \dots\dots\dots (17.13)$$

If L is the inductance of the coil, then changing current sets up a back emf in the coil of magnitude:

$$V = L \frac{\Delta I}{\Delta t}$$

To maintain a constant current, the applied voltage must be equal to the back emf. The magnitude of applied voltage across the coil must have value given by

As $I = I_0 \sin \omega t$, so

$$V = L \frac{\Delta}{\Delta t} (I_0 \sin \omega t)$$

$$V = L I_0 \frac{\Delta}{\Delta t} \sin \omega t \quad \dots\dots\dots (17.14)$$

As $\frac{\Delta}{\Delta t} \sin(\omega t) = \omega \cos(\omega t)$, therefore,

$$V = \omega L I_0 \cos(\omega t)$$

As $\omega L I_0 = V_0$, so

$$V = V_0 \cos(\omega t)$$

As $\cos \omega t = \sin(\omega t + \frac{\pi}{2})$, so

$$V = V_0 \sin(\omega t + \frac{\pi}{2}) \quad \dots\dots (17.15)$$

It is seen from Eq.(17.15), the voltage across the inductor L is leading the current, or in other words, current is lagging behind the voltage in AC circuit containing an inductor as shown graphically in Fig.17.7(b) and by vector diagram in Fig. 17.7 (c). The current in an inductor always lags behind the

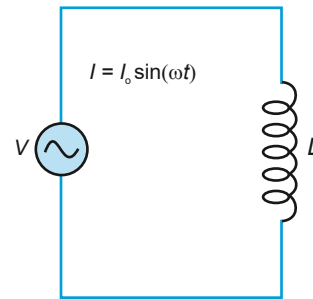


Fig.17.7 (a): Inductor L is connected with an AC source.

For your information

$\Delta I / \Delta t$ is rate of change of current with time. This also represents slope of I - t curve.

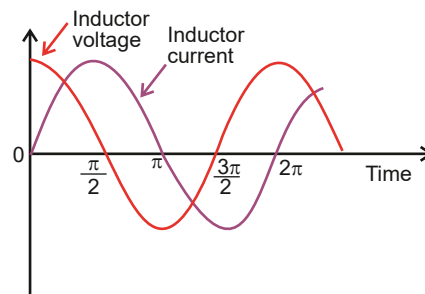


Fig.17.7(b): Graphically showing voltage leading current in an inductor

voltage by 90° or $\pi/2$ rad. The resistance offered by an inductor is called inductive reactance denoted by X_L and is given as

$$X_L = \frac{V_{\text{rms}}}{I_{\text{rms}}} \dots\dots\dots (17.16)$$

where V_{rms} is the rms value of alternating voltage in the inductor and I_{rms} is the rms current passing through it.

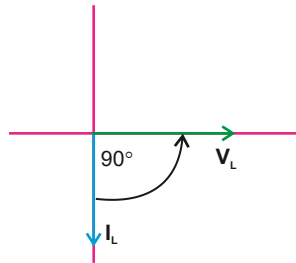


Fig.17.7(c): vector representation of $\mathbf{V}$ and $\mathbf{I}$ in an inductor.

$$\text{Also} \quad X_L = \omega L = 2\pi f L \dots\dots\dots (17.17)$$

It shows that inductive reactance X_L is directly proportional to both, frequency of current and the inductance L of the inductor. In case of DC, $f = 0$, so $X_L = 0$, while in case of large AC frequency, X_L is also large. Thus, we conclude that an inductor allows DC but blocks the AC. The unit of X_L is ohm.

Power dissipation in an inductor

The average power dissipated in a pure inductor is zero. This is because the inductor does not consume power in the traditional sense, as it stores energy in the magnetic field and returns it to the source during demagnetisation. The instantaneous power in the inductive circuit is non-zero and the average power per cycle is zero. This behaviour is due to the phase difference between the voltage and current in an AC circuit, which is 90° . Therefore, the average power dissipated in a pure inductor is zero. Since inductor does not consume energy, it is used for controlling AC without consuming energy.

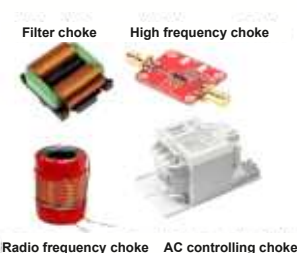
Example 17.3 A 400 mH coil of negligible resistance is connected to an AC circuit in which a current having its rms value of 6 mA is flowing. Find out the value of rms voltage across the coil if its frequency is 1 kHz.

Solution

Inductance of the coil $L = 400 \text{ mH} = 400 \times 10^{-3} \text{ H}$
 and $I_{\text{rms}} = 6 \text{ mA} = 6 \times 10^{-3} \text{ A}$
 Frequency $f = 1 \text{ kHz} = 1000 \text{ Hz}$
 Voltage across coil $V_{\text{rms}} = ?$
 We know that; $V_{\text{rms}} = I_{\text{rms}} X_L$
 As $X_L = 2\pi f L$, therefore,

For your information

Types of chokes



So $V_{\text{rms}} = I_{\text{rms}} (2\pi f L)$
 $V_{\text{rms}} = 6 \times 10^{-3} \text{ A} (2 \times 3.14 \times 1000 \text{ Hz} \times 400 \times 10^{-3} \text{ H}) = 15 \text{ V}$

17.7 CHOKE

It is a coil of thick copper wire wound closely in a large number of turns over a soft iron laminated core.

A choke is often modeled as a series RL-circuit, consisting of a resistor R , having very small value of resistance, in series with an inductor L , of quite large inductance $X_L = 2\pi f L$, as shown in Fig. 17.8. The inductance L of the choke coil is also very high due to the high permeability of iron core on which choke coil is wound. As the resistance R of a choke coil is negligibly small, therefore, the power factor ($\cos\theta$) of the choke coil is almost zero. Thus, the phase difference θ between the current and voltage for a choke coil is nearly equal to 90° . So practically, no power is dissipated as heat by a choke coil. Also choke is used to block high frequency AC while allowing DC and low frequency AC to pass. Choke coils are used to filter out high frequency AC noise from electronic circuits, ensuring a cleaner DC output. They are essential in switch-mode power supplies, helping to regulate voltage and filter out switching noise.



Fig. 17.8: A choke

Choke coils prevent unwanted radio frequency signals from leaking out of circuits, protecting other sensitive components. Chokes in fluorescent lights generate transient voltages across the tube, making it conducive to the breakdown voltage of the gas inside. Choke coils can limit the rate of current changes in circuits, preventing damage to insulation from sudden surges.

17.8 AC THROUGH A CAPACITOR

A capacitor does not allow direct current to pass through it because of the presence of an insulating medium between its plates. But alternating current can pass through a capacitor. In electric circuits, a capacitor is a reactive component. Unlike a resistor, a capacitor behaves differently in AC and DC circuits. It is because a capacitor can store energy in the form of electric field, whereas a resistor cannot store electrical energy in any form. Consider an alternating voltage source V is applied to a

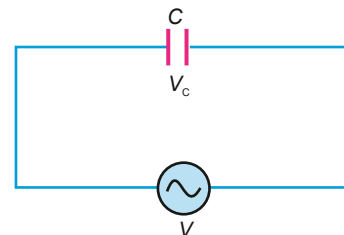


Fig.17.9(a):
Showing capacitor connected with ac voltage source

capacitor of capacitance C as shown in Fig. 17.9 (a).

When an alternating voltage is applied across the plates of a capacitor, it is charged in one direction and then in the other as the voltage reverses. The result is that electrons move to and fro around the circuit, connecting the plates, thus, constituting alternating current. The basic relationship between charge q and the voltage V across the plates: $q = CV$, holds good at every instant. Let V be the applied alternating voltage given by

$$V = V_o \sin \omega t \dots\dots\dots (17.18)$$

The charge on the capacitor at any instant is given by

$$q = CV = CV_o \sin \omega t$$

Then
$$I = \frac{\Delta q}{\Delta t} = \frac{\Delta}{\Delta t} (CV_o \sin \omega t)$$

$$I = CV_o \frac{\Delta}{\Delta t} (\sin \omega t)$$

As $\frac{\Delta}{\Delta t} (\sin \omega t) = \omega \cos \omega t$, thus

$$I = CV_o \omega \cos \omega t$$

As $\cos(\omega t) = \sin\left(\omega t + \frac{\pi}{2}\right)$, so

$$I = \omega CV_o \sin\left(\omega t + \frac{\pi}{2}\right)$$

$$I = \omega CV_o \sin\left(\omega t + \frac{\pi}{2}\right)$$

Here $\omega C V_o = I_o$, thus

$$I = I_o \sin\left(\omega t + \frac{\pi}{2}\right) \dots\dots\dots (17.19)$$

Equations 17.18 and 17.19 show that capacitor current I leads the voltage by 90° or $\pi/2$ rad or voltage lags behind the current I by 90° or $\pi/2$ rad as shown graphically in Fig. 17.9(b) and by vector diagram in Fig. 17.9(c).

Resistance offered by a capacitor is known as capacitive reactance denoted by X_c and is given by

$$X_c = \frac{V_{rms}}{I_{rms}}$$

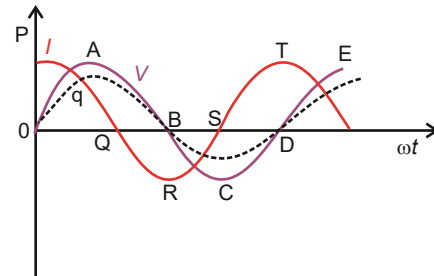


Fig.17.9(b): Graphical representation of current and voltage for capacitor.

Do you know?

Father of ALTERNATING CURRENTS?
Nikola Tesla was born in 1856 in Austria-Hungary and emigrated to the U.S.A. in 1884 as a physicist. He pioneered the generation, transmission, and use of alternating current which can be transmitted over much greater distances than direct current.

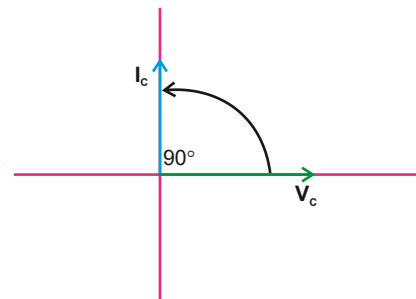


Fig. 17.9(c): Vector representation of current and voltage for a capacitor.

where V_{rms} is the rms value of alternating voltage in the capacitor and I_{rms} is the rms current passes through it.

$$X_C = \frac{1}{\omega C} = \frac{1}{2\pi f C}$$

The unit of capacitive reactance X_C is ohm.

Which shows that for low frequencies, capacitor will have a large reactance X_C and current I will be small whereas at high frequencies reactance X_C will be small and current I through the same capacitor will be large.

17.9 IMPEDANCE

In AC circuits, the opposition to current flow is called impedance denoted by Z , which includes both resistance R and reactance X . The resistance R is the opposition to current flow due to resistors. Reactance is the opposition to current flow due to capacitors (capacitive reactance X_C) and inductors (inductive reactance X_L). A pure resistive circuit has only resistance R and no reactance.

The combined effect of resistance and reactances in such circuits is known as impedance.

It is the ratio of the rms value of the voltage to the rms value of the current. Thus,

$$Z = \frac{V_{\text{rms}}}{I_{\text{rms}}} \quad \dots\dots\dots(17.20)$$

This SI unit of impedance is ohm.

Example 17.4 An AC circuit operates by a peak voltage of 200 V and 10 A as peak input current. Find the impedance of the circuit.

Solution

Peak voltage $V_o = 200 \text{ V}$

Peak current $I_o = 10 \text{ A}$

We know that; $V_{\text{rms}} = V_o \times 0.707$
 $= 200 \text{ V} \times 0.707 = 141.4 \text{ V}$

and $I_{\text{rms}} = I_o \times 0.707$
 $= 10 \text{ A} \times 0.707 = 7.07 \text{ A}$

As Impedance $Z = \frac{V_{\text{rms}}}{I_{\text{rms}}}$, therefore,
 $Z = \frac{141.4 \text{ V}}{7.07 \text{ A}} = 20 \Omega$

17.10 AC THROUGH RC-SERIES CIRCUIT

Consider a network of resistance R and a capacitor C connected in series by an alternating voltage source V as shown in Fig.17.10 (a). As R and C are in series, so same current would flow through each of them. The potential difference V across the resistance R would be $I_{\text{rms}} R$ and it would be in phase with the current. The vector diagram of the voltage and current is shown in Fig.17.10 (b). Taking the current as reference, potential difference $V_R = I_{\text{rms}} R$ across the resistance is represented by a line along the current line because potential drop $I_{\text{rms}} R$ is in phase with current. The potential difference across the capacitor will be:

$$V_c = I_{\text{rms}} X_c = \frac{I_{\text{rms}}}{\omega C}$$

As this, voltage lags behind the current by $\pi/2$ rad or 90° , so the line representing the vector $I_{\text{rms}} X_c$ is drawn at right angles to the current line Fig.17.10 (b).

The applied voltage V_{rms} that will send the current I in the circuit is obtained by resultant of the vectors

$I_{\text{rms}} R$ and $\frac{I_{\text{rms}}}{\omega C}$ i.e;

$$V_{\text{rms}} = \sqrt{(V_R)^2 + (V_c)^2}$$

$$V_{\text{rms}} = \sqrt{(I_{\text{rms}} R)^2 + \left(\frac{I_{\text{rms}}}{\omega C}\right)^2}$$

$$V_{\text{rms}} = I_{\text{rms}} \sqrt{R^2 + \left(\frac{1}{\omega C}\right)^2}$$

$$V_{\text{rms}} = I_{\text{rms}} \sqrt{R^2 + \left(\frac{1}{\omega C}\right)^2}$$

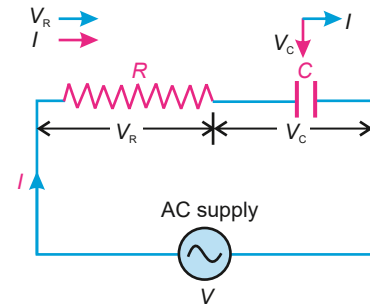


Fig. 17.10(a): Resistor and capacitor connected in series with AC voltage source

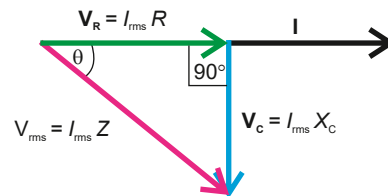


Fig. 17.10(b): Vector diagram showing relationship of I with V_c and V_R .

Impedance
$$Z = \frac{V_{\text{rms}}}{I_{\text{rms}}} = \sqrt{R^2 + \left(\frac{1}{\omega C}\right)^2} \dots\dots\dots(17.21)$$

Equation 17.21 suggests that we can find the impedance of a series AC circuit by vector addition. The resistance R is represented by a horizontal line in the direction of current which is taken as reference. The reactance $X_c = 1/\omega C$ is shown by a line lagging the R -line by 90° as shown in Fig.17.10(c). The impedance Z of the circuit is obtained by the vector summation of resistance and reactance. Figure 17.10(c) is known as impedance diagram of the circuit. The angle which the line representing the impedance Z makes with R line gives the phase difference between the voltage and current I in Fig. 17.10(c),

the current is leading the voltage applied by an angle $\pi/2$ or 90° as shown in Fig. 17.10(d) and is given by

$$\theta = \tan^{-1}\left(\frac{X_c}{R}\right) = \tan^{-1}\left(\frac{1}{\omega CR}\right) = \tan^{-1}\left(\frac{1}{2\pi fC}\right)$$

The power consumed in RC-series circuit is primarily due to the resistor, as the capacitor only stores energy and does not dissipate it. The power consumed by the resistor can be calculated by using the formula:

$$P = I^2 R$$

or $P = VI \cos \theta \dots\dots\dots (17.22)$

where I is the current, R is the resistance, V is the voltage, and θ is the phase angle between voltage and current.

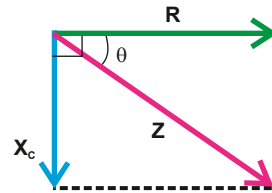


Fig. 17.10(c): Vector diagram of RC-series circuit.

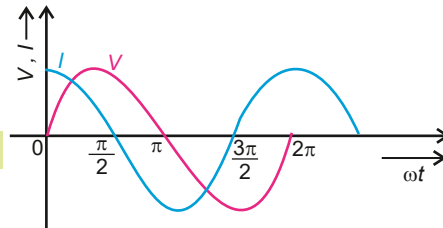


Fig. 17.10(d): Voltage and current waveform in RC-series circuit.

Example 17.5 A resistor of resistance $5 \text{ k}\Omega$ is connected in series with capacitor of capacitance $5 \mu\text{F}$ across the AC source of 220 V that has frequency 50 Hz . Calculate the phase angle.

Solution Resistance $R = 5 \text{ k}\Omega = 5000 \Omega$,
 Capacitance $C = 5 \mu\text{F} = 5 \times 10^{-6} \text{ F}$
 and Frequency $f = 50 \text{ Hz}$

We know that;

Phase angle $\theta = \tan^{-1}\left(\frac{1}{\omega CR}\right)$

As $\omega = 2\pi f$, so

$$\theta = \tan^{-1}\left(\frac{1}{2\pi f CR}\right)$$

Putting the values

$$\theta = \tan^{-1}\left(\frac{1}{2 \times 3.14 \times 50 \text{ Hz} \times 0.000005 \text{ F} \times 5000 \Omega}\right)$$

$$\theta = \tan^{-1}(0.1273) = 7.25^\circ$$

17.11 AC THROUGH RL-SERIES CIRCUIT

Consider an AC circuit containing a resistor R and inductor L connected in series with an AC source V as shown in Fig. 17.11(a). As R and L are in series, therefore, same current I flows through both R and L .

$$I = I_0 \sin \theta$$

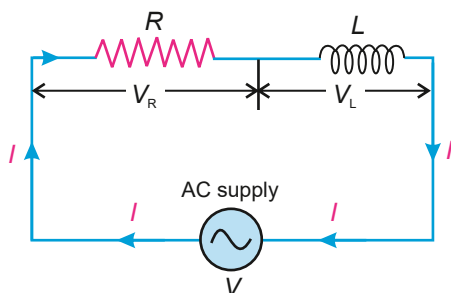


Fig.17.11 (a): RL-series circuit connected with AC supply

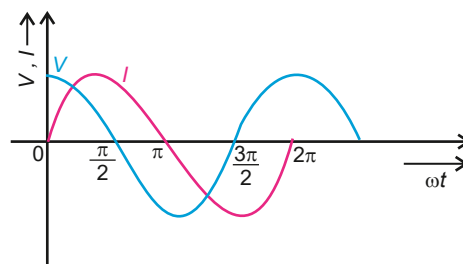


Fig. 17.11(b): Voltage and current waveform in RL-series circuit.

Let rms value of current through R is I . Then rms value of voltage V_R across R will be:

$$V_R = I_{\text{rms}} R$$

V will be in phase with the current I in the circuit, as R is shown vectorially by line OA in Fig.17.11(c). The value of rms voltage V_L across L will be:

$$V_L = I_{\text{rms}} X_L$$

Voltage V across L is leading the current I through L by $\pi/2$ or 90° as shown in Fig.17.11 (b) and vectorially by line OB in Fig.17.11(c). The resultant voltage V will be the vector sum of V_R and V_L shown vectorially by line OP , thus in Fig.17.11(c):

$$V_{\text{rms}} = \sqrt{V_R^2 + V_L^2}$$

$$V_{\text{rms}} = \sqrt{(I_{\text{rms}} R)^2 + (I_{\text{rms}} X_L)^2}$$

As $X_L = \omega L$, so

$$V_{\text{rms}} = I_{\text{rms}} \sqrt{R^2 + (\omega L)^2}$$

$$\text{Impedance } Z = \frac{V_{\text{rms}}}{I_{\text{rms}}} = \sqrt{R^2 + (\omega L)^2} \dots (17.23)$$

$$\text{Angle } \theta = \tan^{-1} \left(\frac{\omega L}{R} \right)$$

The angle θ which Z makes with R line, gives the phase difference between the applied voltage and current.

Power Dissipation in RL-Series Circuit

Power consumed in RL- series circuit is primarily due to the resistor, as the inductor stores energy but ideally does not dissipate it. The power consumed by the resistor is given by the formula:

$$P = I^2 R = VI \cos \theta$$

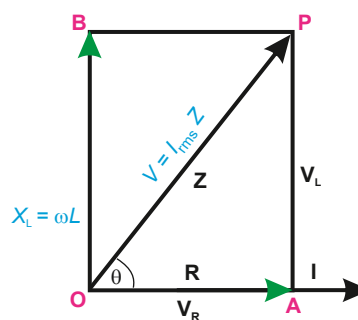


Fig. 17.11(c): Vector diagram of V_R and V_L in RL-series circuit.

For your information



Electric Eels create powerful shocks by its 3 main organs consisting of 80% of Eels body length through specialized cells called electrolytes.

where I , R and V represents current, resistance, and voltage respectively whereas θ is the phase angle between voltage and current. The factor $\cos \theta$ is known as power factor. It is to be noted that when we convert DC power P_{dc} , into AC power P_{ac} , we have to take an account of a quantity known as inverter efficiency. It usually may be 85 % to 90%. So, we can find AC power from DC power by using the following conversion as:

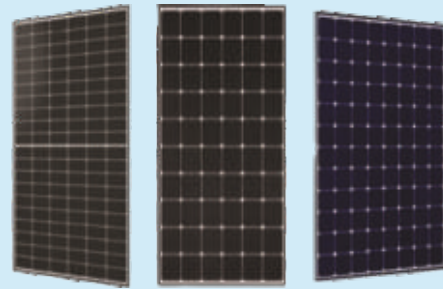
$$P_{ac} = (\text{Inverter efficiency}) \times P_{dc}$$

Do you know?



The electricity consumption of a 1.5 ton AC is approximately 1.2 to 1.5 units per hour.

Brain teaser



How much AC power will be received from a DC solar panel rating 575 watt, 75% inverting efficiency?

17.12 RECTIFICATION

Rectification is the process of converting alternating current into direct current.

Diodes allow current to flow in one direction (forward bias) but block it in the opposite direction (reverse bias). Rectification is of two types, half-wave rectification and full-wave rectification. Rectification is commonly used in electronic devices that require a stable DC power source, such as computers and smartphones, as most household and commercial power is AC.

Diode as a Rectifier

A diode allows large current to flow when forward biased. However, the current through a reverse biased diode is practically zero. It is due to this important property of the diode that it can be used for rectification.

Half-Wave Rectifier

In a half-wave rectification, an AC signal is converted into pulsating direct current DC by passing one half-cycle of waveform and blocking

Brain teaser

Cell phones store AC or DC!

Cell phone batteries store DC in them, since it is easier to store DC than AC. DC is also safer compared to AC. The electric grid provides AC only. Therefore, the AC is converted to DC using a rectifier before charging the cell phones or any other portable devices such as laptops, flashlights, etc.

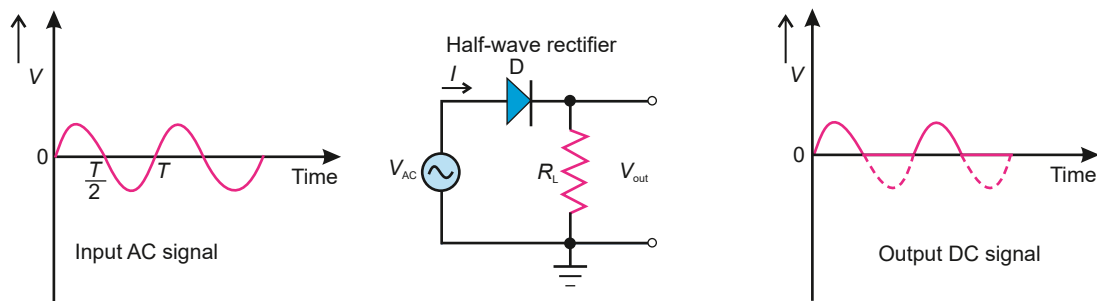


Figure 17.12: Working of a half-wave rectifier, showing graphically its input and output waveforms.

the other half. In Fig.17.12, alternating voltage of period T called input voltage is supplied to a diode D which is connected in series with a load resistance R_L . The working of half-wave rectifier is based on the fact that the diode allows the current flow only in one direction. Thus, it converts the AC signal into DC signal. During positive half-cycle of the input alternating voltage i.e., during the interval $0 \rightarrow T/2$, the diode D is forward bias, so it offers a very low resistance and current flows through R_L .

During negative half cycle, during the interval $T/2 \rightarrow T$, the diode becomes reverse bias and offers a very high resistance, so practically no current flows through R_L . Same events repeat during the next cycle and so on. The current across R_L flows in only one direction which means it is direct current. However, this current flows in pulses.

Full-Wave Rectification

Rectification during which both halves of alternating input voltage are converted into unidirectional current through a resistance is called full-wave rectification. Full-wave rectification is achieved by a bridge rectifier in which four individual rectifying diodes are connected in a closed loop “bridge” configuration to produce the desired output. The single secondary winding is connected to one side of the diode bridge network and the load to the other side. The four diodes labelled diodes D_1 to D_4 are arranged in “series pairs” with only two diodes conducting current during each half-cycle as shown in Fig.17.13(a)

The Positive Half-cycle: We know that diode conducts only when it is forward bias. During the positive half-cycle of the supply i.e., during the time $0 \rightarrow T/2$, the terminal A of the bridge is positive with respect to its other terminal B. Now the diodes D_1 and D_2 become forward biased and conduct in series while diodes D_3 and D_4 are reverse

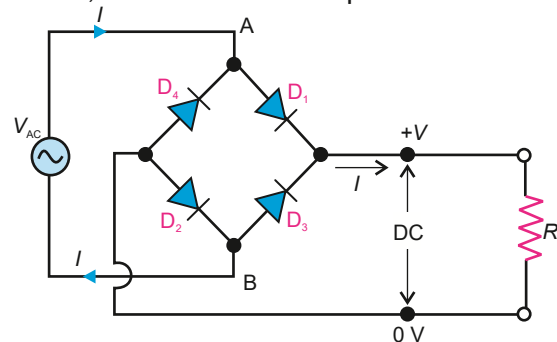


Figure 17.13(a): Full-wave bridge rectifier circuit showing 4 diode network with load resistance.

Do you know?

Uses of the PN junction diode:

1. Rectification
2. Signal clipping
3. Regulation
4. Clamping voltage
5. Signal detection
6. Photovoltaic cells and
7. Switching light emission

biased and the current flows through the load resistance R as shown in Fig. 17.13 (b).

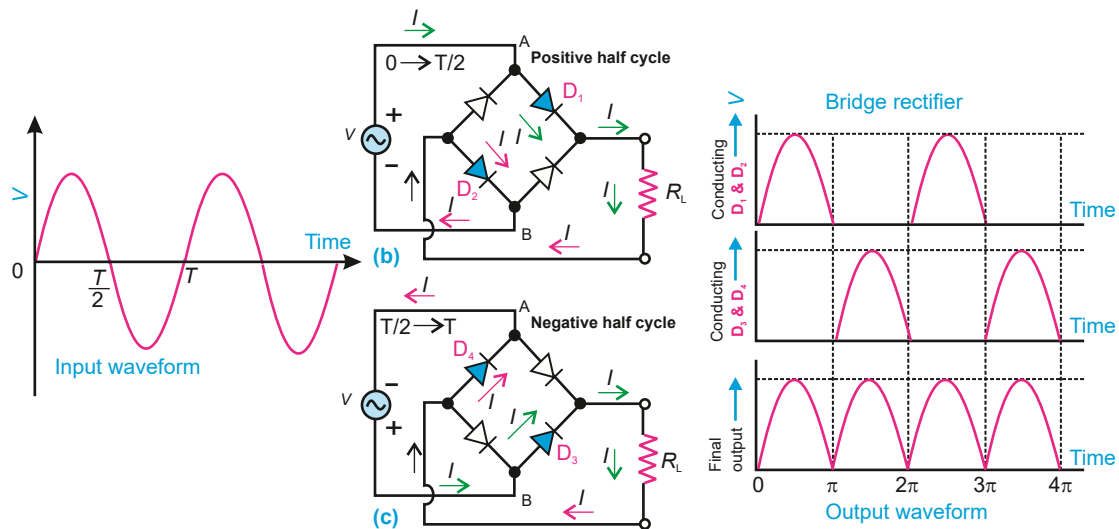


Figure 17.13(b), Figure 17.13 (c): Full-wave bridge rectifier circuit showing its working during its positive and negative half cycles and also representing graphically its input and output waveforms.

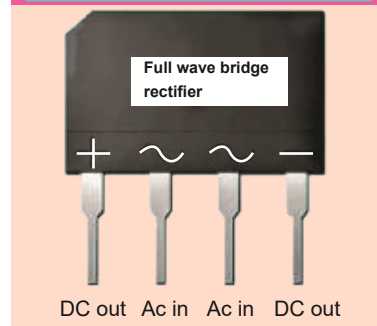
The Negative Half-cycle: During the negative half-cycle of the supply, during the time $T/2 \rightarrow T$, terminal A is negative and B is positive. Now the diodes D_3 and D_4 conduct in series, but diodes D_1 and D_2 switch “OFF” as they are now reverse biased. The current flowing through the load is the same direction as before as shown in Fig. 17.13 (c). If we take a comparison of Figs. 17.13(b) and 17.13(c), it can be seen that direction of current flow through the load resistance is same in both half of the cycle. Thus, both halves of the alternating input voltage send a unidirectional current through load resistance R . The output voltage is not smooth but pulsating. It can be made smooth using a circuit called filter which may consist of capacitors, inductors and resistors and their combinations.

Full-Wave Bridge Rectifier with Single Smoothing Capacitor

In AC to DC conversion, smoothing filters are used to reduce the ripple voltage after rectification. When AC is rectified, the output is pulsating DC, which is not suitable for most electronic devices. Smoothing filters help to produce a more stable DC voltage by filtering the AC components (ripples). They are of different types. In capacitor filter, a capacitor is connected in parallel with the load. It charges when the rectified voltage increases and discharges when the voltage decreases, thereby smoothing the output.

Smoothing or reservoir capacitors connected in parallel with the load across the output of the full-wave bridge rectifier circuit, increase the average DC output level

For your information



even higher as the capacitor acts like a storage device as shown in Fig.17.14. The smoothing capacitor converts the full-wave rippled output of the rectifier into a more smooth DC output voltage. We can see the effect it has on the rectified output waveform with different values of smoothing capacitor installed.

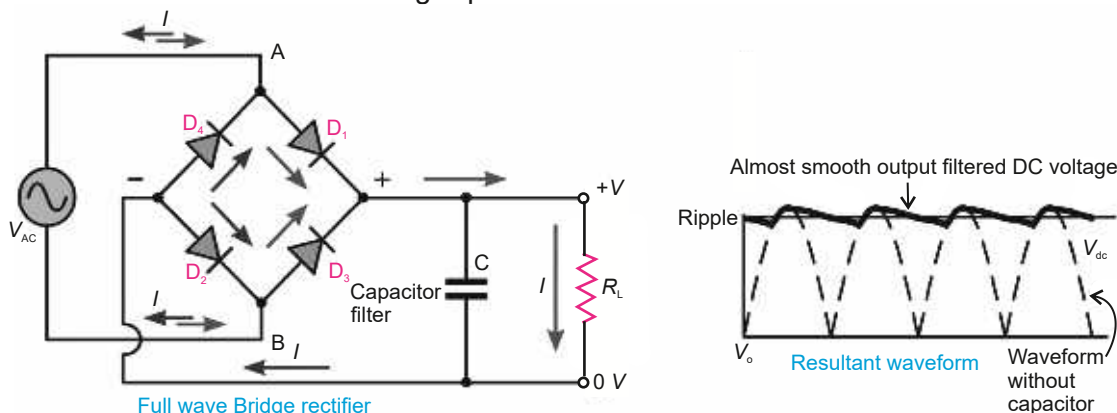


Fig.17.14: The smoothing capacitor converts the full-wave rippled output of the rectifier into a more smooth DC output voltage.

The maximum ripple voltage present in a full-wave rectifier bridge circuit is not determined by the value of the smoothing capacitor but by the frequency and load current, and is calculated as:

$$V_{\text{ripple}} = \frac{I_{\text{DC}}}{f_{\text{ripple}} C} \quad \text{or} \quad V_{\text{ripple}} = \frac{I_{\text{DC}}}{2fC}$$

where I_{DC} is the load current in ampere, f is input frequency and C is capacitance in farad (F). It is to be noted that ripple frequency f_{ripple} is twice that of the input frequency,

i.e., $f_{\text{ripple}} = 2f$

Brain teaser

Our heart is driven by electric pulses; the high electric frequency of AC current can affect the frequency of the heart and can cause fibrillation.

17.13 SELF-INDUCTION

Consider a circuit in which a coil is connected in series with a battery, a switch and a rheostat as shown in Fig.17.15. If we move the rheostat quickly, the primary current through the coil will change. The magnetic flux through the coil will also change, which finally induces an emf in the coil itself.

The phenomenon, in which changing current induces an emf inside the coil itself, is called self-induction.

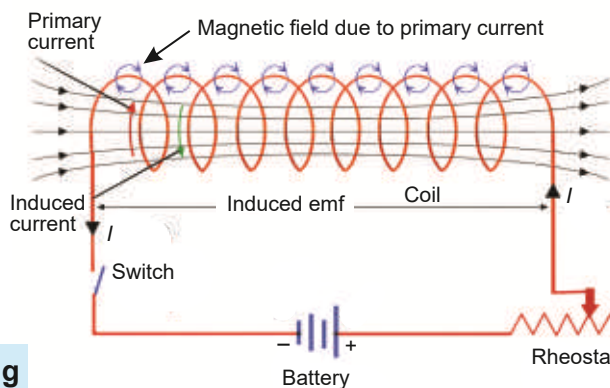


Fig.17.15: A coil showing production of induced current and induced emf in itself.

Let ϕ be the flux passing through one loop of the coil. The flux passing through the coil of N turns would be $N\phi$. As this flux ϕ is proportional to the magnetic field produced which is in turn proportional to the current I , therefore,

$$N\phi \propto I$$

or $N\phi = LI \dots\dots\dots (17.24)$

where L is proportionality constant and is called self-inductance of the coil.

The self-inductance L depends upon the following factors:

1. Number of turns of the coil.
2. Area of cross-section of the coil.
3. Core material, by winding the coil around ferromagnetic iron core, the magnetic flux and hence inductance can be increased significantly relative to that for an air core.

By Faraday's law, the emf induced in the coil is given by

$$\varepsilon_L = -N \frac{\Delta\phi}{\Delta t} \dots\dots\dots (17.25)$$

$$\varepsilon_L = - \frac{\Delta(N\phi)}{\Delta t}$$

Using Eq. (17.24); $\varepsilon_L = - \frac{\Delta(LI)}{\Delta t}$

$$\varepsilon_L = -L \frac{\Delta I}{\Delta t} \dots\dots\dots (17.26)$$

The magnitude of self-inductance is:

or $L = \frac{\varepsilon_L}{\left(\frac{\Delta I}{\Delta t}\right)} \dots\dots\dots (17.27)$

Here L is self-inductance of the coil and minus sign indicates that the induced emf opposes the applied voltage. In the above expression, If we take $\varepsilon_L = 1$ volt and $\Delta I/\Delta t = 1 \text{ A s}^{-1}$, then $L = 1$ henry (H), defined as:

If the emf induced in the coil is 1 V when the current flowing through it changes at the rate of 1 A s^{-1} , the coil has self-inductance of 1 henry (H).

The self-inductance of the coil can also be defined as ratio of the emf to the rate of change of current in the same coil. The SI unit of self-inductance L is Vs A^{-1} .

By equating Eqs. 17.25 and 17.26, we have

$$-N \frac{\Delta\phi}{\Delta t} = -L \frac{\Delta I}{\Delta t}$$

$$L = N \frac{\Delta \phi}{\Delta I} \dots\dots\dots (17.28)$$

which gives another form of equation for self-inductance of the coil.

Energy Stored in an Inductor

Consider a coil connected to a battery and a switch in series as shown in Fig.17.16. When the switch is turned ON, voltage is applied across the ends of coil and current increases from zero to maximum value. Due to change of current, an emf is induced, which is opposite to that of battery. Work is done by battery to move charges against the induced emf.

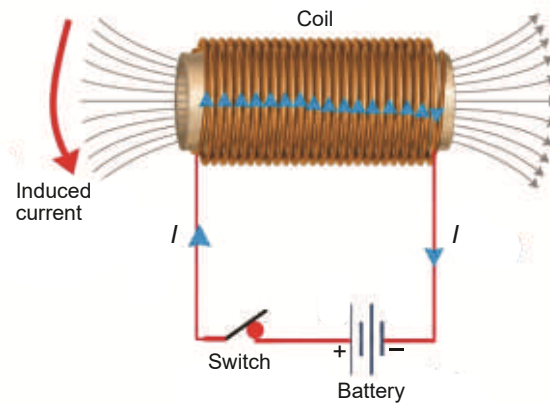


Figure 17.16: Circuit designed to calculate energy stored in an inductor.

Work done by battery in moving a small charge Δq is given by

$$W = \Delta q \mathcal{E}_L \dots\dots\dots (17.29)$$

where $\mathcal{E}_L$ is the magnitude of self-induced emf and is given by

$$\mathcal{E}_L = L \frac{\Delta I}{\Delta t}$$

So, putting values in Eq.17.29, we have

$$W = \Delta q L \frac{\Delta I}{\Delta t} = \frac{\Delta q}{\Delta t} L \Delta I \dots\dots\dots (17.30)$$

where $\frac{\Delta q}{\Delta t}$ represents average current and is given by

$$\frac{\Delta q}{\Delta t} = \left(\frac{0 + I}{2} \right) = \frac{1}{2} I$$

And $\Delta I = I - 0 = I$

Putting the values in Eq.17.30, we have

$$W = \left(\frac{1}{2} I \right) L I = \frac{1}{2} L I^2 \dots\dots\dots (17.31)$$

This work is stored as potential energy in the inductor. Hence, energy stored in an inductor is:

$$U_m = \frac{1}{2} LI^2 \dots\dots\dots (17.32)$$

It is to be noted that in an inductor, energy is stored in the magnetic field and can never be negative because it is proportional to square of current. The inductor stores energy while absorbing power, returns the previous stored energy when delivering power, so the net energy transfer can never be negative. Using Eq.17.32, expressed in terms of magnetic field **B** of solenoid can be found having:

n = number of turns per unit length

A = Area of cross-section of coil

B = Magnetic field strength = $\mu_o n I$ (Magnetic field due to solenoid)

I = current in solenoid = $\frac{B}{\mu_o n}$

L = Self Inductance = $\mu_o n^2 A \ell$ where ℓ is the length of solenoid

Putting the values in Eq.(17.32), we have

$$U_m = \frac{1}{2} \mu_o n^2 A \ell \left(\frac{B}{\mu_o n} \right)^2$$

$$U_m = \frac{B^2}{2\mu_o} (A\ell)$$

As Al = volume, so $U_m = B^2 \frac{1}{2\mu_o} (\text{volume})$

or $\frac{U_m}{\text{volume}} = B^2 \frac{1}{2\mu_o}$

The energy stored U_m , per unit volume inside the solenoid is called energy density, and for an inductor it is given by

$$\text{Energy density} = B^2 \frac{1}{2\mu_o} \dots\dots\dots (17.33)$$

17.14 MUTUAL INDUCTION

Consider two coils close to each other as shown in Fig. (17.17). One coil connected with a battery through a switch S and a rheostat, is called primary coil and the other one connected to the galvanometer is called secondary coil. If the current in the primary is changed by varying the resistance of the rheostat, the magnetic flux in the surrounding region changes. Since the secondary coil is in magnetic field of the primary, the changing flux also links with the secondary. This causes an induced emf in the secondary.

The phenomenon in which a changing current in one coil induces an emf in another coil is called mutual induction.

According to Faraday's law, the emf induced in the secondary coil $\mathcal{E}_s$ is proportional to the rate of change of flux $\Delta\phi_s/\Delta t$ passing through it and is given by

$$\mathcal{E}_s = -N_s \frac{\Delta\phi_s}{\Delta t} \dots\dots\dots (17.34)$$

where N_s is the number of turns in the secondary coil.

Let ϕ_s be the flux passing through the secondary coil. The net flux passing through the secondary coil of N_s loops is $N_s \phi_s$. As this net flux is proportional to the magnetic field produced by the current I_p in the primary and the magnetic field itself is proportional to I_p , therefore

$$N_s \phi_s \propto I_p$$

$$N_s \phi_s = M I_p \dots\dots\dots (17.35)$$

where M is proportionality constant and is called mutual inductance of the two coils, given by

$$M = N_s \frac{\phi_s}{I_p}$$

The mutual inductance M depends upon the following factors,

1. Number of turns of both the primary and secondary coils.
2. Area of cross-section of the two coils.
3. Magnetic permeability of medium between the coils.
4. Nature of material on which two coils are wound.
5. Distance between two coils.
6. Orientation between the primary and secondary coils.

By Faraday's law, the emf in the secondary coil is given by the rate of change of flux through the secondary coil as:

$$\mathcal{E}_s = -N \frac{\Delta\phi_s}{\Delta t} \quad \text{or} \quad \mathcal{E}_s = -\frac{\Delta(N_s \phi_s)}{\Delta t}$$

Using Eq. (17.35), we have: $\mathcal{E}_s = -\frac{\Delta(M I_p)}{\Delta t}$

$$\mathcal{E}_s = -M \frac{\Delta I_p}{\Delta t} \dots\dots\dots (17.36)$$

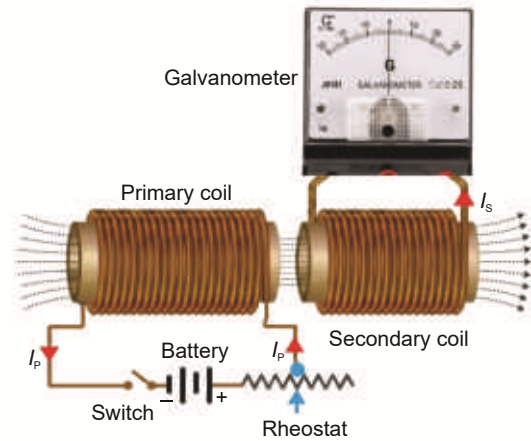


Figure (17.17): Circuit diagram to find Mutual induction between two coils.

Which shows that emf induced in the secondary coil is proportional to time rate of change of current in the primary. The negative sign indicates the fact that the induced emf is in such a direction that it opposes the change of current in the primary coil. The magnitude of mutual inductance is:

$$M = \frac{\varepsilon_s}{\left(\frac{\Delta I_p}{\Delta t} \right)} \dots\dots\dots(17.37)$$

Mutual inductance can be defined as ratio of average emf induced in the secondary to the time rate of change of current in the primary.

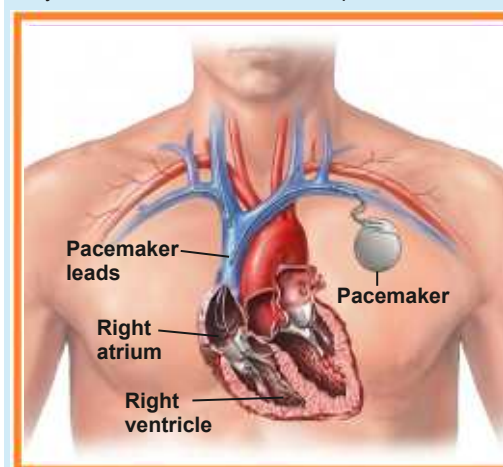
The SI unit of mutual inductance M is V s A^{-1} , which is called henry after Joseph Henry.

One henry is the mutual inductance of the pair of coils in which the rate of change of current of one ampere per second in the primary coil causes an induced emf of one volt in the secondary coil .

Mutual inductance finds many applications in various devices such as transformers, electric motors, and generators etc. It also plays a key role in digital signal processing and is utilized in devices like pacemakers and metal detectors.

Brain teaser

Do you have an idea how does a pacemaker work?



QUESTIONS

Multiple Choice Questions

Choose the correct answer.

17.1 Name of device that is used as a part of charging circuit for mobile phones is:

- (a) full-wave bridge rectifier
- (b) half-wave rectifier
- (c) AC generator
- (d) battery

17.2 A $10 \mu\text{F}$ capacitor is plugged into a $110 \text{ V}_{\text{rms}}$, 60 Hz voltage source, with an ammeter in series. The value of I_{rms} that passes through the capacitor is:

- (a) 0.0415 A
- (b) 0.415 A
- (c) 415 A
- (d) 0.0415 mA

17.3 In an RC-circuit, a $10 \text{ k}\Omega$ resistance is connected in series to a capacitor of $0.05 \mu\text{F}$. The applied voltage for charging is 36 V , the RC-time constant to charge a

capacitor is:

- (a) 0.0578 ms (b) 0.009 ms (c) 0.500 ms (d) 5.49 ms

17.4 In a purely inductive or capacitive circuit, the average power consumed is:

- (a) equal to apparent power (b) minimum
(c) maximum (d) zero

17.5 Through which of the AC-circuit elements both emf and current are in phase?

- (a) resistance (b) inductor (c) capacitor (d) LED

17.6 If magnetic field is doubled in an inductor, then, magnetic energy density becomes:

- (a) 6 times (b) half (c) 4 times (d) constant

17.7 RC-time constant is measured in:

- (a) ohm (b) henry (c) volt (d) second

17.8 The magnitude of mutual inductance M between two coils when current changes at 20 A s^{-1} in one coil induces an emf of 50 mV in the other is:

- (a) 5 mH (b) 2.5 mH (c) 0.0025 μH (d) 25 mH

Short Answer Questions

- 17.1 What impacts does resistance have on capacitance?
17.2 What happens when an AC line touches a DC line?
17.3 What is the difference between peak to peak and amplitude in case of sine wave?
17.4 Why are RC-circuits used in timing circuits?
17.5 Why are choke coils preferred over resistors for limiting current in AC circuits?
17.6 Define rms values of voltage and current with their phasor diagrams.

Constructed Response Questions

- 17.1 What do you mean by rating of electrical appliances? Explain briefly.
17.2 The energy stored in an inductor is analogous to the kinetic energy of a moving mass? Justify.
17.3 What are safety measures to take when working with AC or DC?
17.4 How does self-induction manifest when a current is switched on or off in a circuit?
17.5 Is it possible to achieve mutual inductance using a combination of four coils? If yes, justify your answer.
17.6 The currents of the order of 0.1 A through human body is fatal what causes death; heating due to current or something else?

Comprehensive Questions

- 17.1 Explain the behaviour of a capacitor in an AC circuit. Derive the expression for capacitive reactance and discuss why current leads the voltage by 90° or $\pi/2$ rad.
- 17.2 Explain the behaviour of an inductor in an AC circuit. Derive the expression for inductive reactance and explain why current lags behind the voltage by 90° or $\pi/2$ rad.
- 17.3 Find expression for impedance in case of RC and RL-series circuits.
- 17.4 Define full-wave rectification? Draw circuit, input and output signal phasor diagrams to explain how full-wave rectification is achieved by use of full wave bridge rectifier.
- 17.5 Define self-induction and derive the expression for self-inductance.
- 17.6 What is meant by mutual inductance? Derive a relationship for it and hence define henry.

Numerical Problems

- 17.1 AC circuit configured by a resistance of $20\ \Omega$, capacitor of capacitance $40\ \mu\text{F}$ is connected to an AC supply of $110\ \text{V}$ having frequency $50\ \text{Hz}$. Calculate:
(a) capacitive reactance. (b) current in the circuit (c) suggest phase between the voltage and current in a capacitive circuit.
(Ans: (a) $79.58\ \Omega$ (b) $1.34\ \text{A}$ (c) current leads voltage by 75.9°)
- 17.2 An alternating current i is represented by the equation; $i = 420 \sin(100\pi t)$. Find:
(a) the rms current (b) the frequency
(Ans: (a) $296.94\ \text{A}$ (b) $50\ \text{Hz}$)
- 17.3 A $50\ \text{Hz}$ AC of peak value $1\ \text{A}$ flows through primary coil of a transformer. If the mutual inductance between the primary and secondary be $1.5\ \text{H}$, then find the mean value of induced voltage.
(Ans: $300\ \text{V}$)
- 17.4 A $220\ \text{V}$ AC source of frequency $2\ \text{kHz}$ is connected across a capacitor of $10\ \mu\text{F}$. Find out: (a) the reactance X_c of capacitor (b) the current in the circuit.
(Ans: (a) $7.96\ \Omega$ (b) $27.63\ \text{A}$)
- 17.5 In a coil of inductance $150\ \text{mH}$, a current changes steadily from $50\ \text{mA}$ to $30\ \text{mA}$ in $164\ \text{ms}$, find: (a) the magnitude (b) direction of induced emf.
(Ans: (a) $18.29\ \text{mV}$ (b) same direction as the original current)
- 17.6 A choke coil is needed to operate a lamp at $160\ \text{V}_{\text{rms}}$ and $50\ \text{Hz}$. The lamp has a resistance of $5\ \Omega$ when running at $10\ \text{A}_{\text{rms}}$. Calculate the inductance of the choke coil.
(Ans: $48.4\ \text{mH}$)

Quantum Physics

Students Learning Outcomes

After studying this chapter, the students will be able to:

- ◆ state that electromagnetic radiation has a particulate nature.
- ◆ explain and apply the photonic model of light to solve problems.
- ◆ [use $E = hf$ to solve problems, and use the electron volt (eV) as a unit of energy].
- ◆ explain that a photon has momentum [including that the momentum is given by $p = E/c$ (connect with the idea that light can exert a force)].
- ◆ describe that photoelectrons may be emitted from a metal surface when it is illuminated by electromagnetic radiations.
- ◆ describe and use the terms threshold frequency and threshold wavelength.
- ◆ explain photoelectric emission in terms of photon energy and work function energy.
- ◆ state and apply; $hf = \phi + \frac{1}{2}mv_{\max}^2$
- ◆ explain why the maximum kinetic energy of photoelectrons is independent of intensity, whereas the photoelectric current is proportional to intensity.
- ◆ evidence for light as a wave and as a particle [Explain that the photoelectric effect provides evidence for a particulate nature of electromagnetic radiation while phenomena such as interference and diffraction provide evidence for a wave nature].
- ◆ Discuss qualitatively the evidence provided by electron diffraction for the wave nature of particles.
- ◆ explain and apply the de Broglie wavelength to solve problems [use to solve problems]; $\lambda = h/p$.
- ◆ state that there are discrete electron energy levels in isolated atoms (e.g. atomic hydrogen).
- ◆ explain the appearance and formation of emission and absorption line spectra.
- ◆ use of; $E = hf$ to solve problems.
- ◆ describe the Compton effect qualitatively.
- ◆ Recall the phenomena of pair production and pair annihilation.
- ◆ explain how electron microscopes achieve very high resolution.
- ◆ state and explain Heisenberg's uncertainty principle qualitatively.
- ◆ use the uncertainty principle to explain why empirical measurements must necessarily have uncertainty in them.

Quantum Theory assumes the behaviour of electromagnetic radiations as discrete packets of energy and particles on a very small scale to behave as waves. It is probably the most successful theory in physics at providing explanations for the unresolved issues and accurate predictions.

However, the classical physics is still valid in ordinary processes of everyday life. We shall discuss various aspects of quantum theory in this chapter.

18.1 QUANTUM THEORY OF RADIATION

In 1901, Max Planck suggested that energy is radiated or absorbed in discrete packets of energy called quanta rather than as a continuous wave. Quanta is plural of quantum, a discrete packet of energy. Each quantum is associated with radiations of a single frequency. The energy E of each quantum is proportional to its frequency f related as

For your information

The one important constant h is a constant of quantum theory and another important constant c is the speed of light which is a constant of special theory of relativity.

$$E = hf \text{(18.1)}$$

where h is Planck's constant, its value is $6.63 \times 10^{-34} \text{ J s}$.

Max Planck received Nobel Prize in physics in 1918 for his discovery of energy quanta.

The Photon

Max Planck suggested that as matter is not continuous but consists of a large number of tiny particles, so is the radiation energy from a source. He assumed that granular or particle nature of radiation from hot bodies was due to some property of the atoms producing it. Einstein extended his idea and postulated that packets or tiny bundles of energy are integral part of all electromagnetic radiations and that they could not be subdivided. These indivisible tiny bundles of energy he called "photons". A monochromatic beam of light with wavelength λ consists of a stream of photons travelling at speed c and carries energy; $E = hf$.

From the theory of relativity; $E = mc^2$, the relativistic momentum of photon is; $p = mc$,

hence, $E = pc \text{(18.2)}$

Thus $pc = hf$ or $p = hf/c$ Since $c = f\lambda$, therefore,

Momentum of photon $p = h/\lambda \text{(18.3)}$

The quantum theory may be extended to include any system such as a mass oscillating on a spring. However, energy steps are far too small to be detected, so any particle nature is invisible. Quantum effects are only important when observing atom sized objects, where h is a significant factor in any detectable energy change.

18.2 PHOTOELECTRIC EFFECT

The process of emission of electrons from a metal surface when exposed to light of suitable frequency is called the photoelectric effect. The emitted electrons are known as photoelectrons.

The photoelectric effect is demonstrated by the apparatus shown in Fig. 18.1. An evacuated glass tube X contains two electrodes. The electrode A connected to the positive terminal of the battery acts as anode. The metal electrode C connected to negative terminal acts as cathode. When monochromatic light of suitable frequency is

allowed to shine on cathode, it begins to emit electrons. These photoelectrons are attracted by the positive anode and the resulting current is measured by a micro-ammeter. The current stops when light is cut off, which proves that the current flows because of incident light. This current is, hence, called photoelectric current. The maximum energy of the photoelectrons can be determined by reversing the connections of the battery in the circuit i.e., now the anode A is negative and cathode C is at positive potential. In this condition, the photoelectrons are repelled by the anode and the photoelectric current decreases. If this potential is made more and more negative, at a certain value, called stopping potential V_o , the current becomes zero. Even the electrons of maximum energy are not able to reach collector plate. The maximum kinetic energy of photoelectrons is, thus

$$\frac{1}{2}mv_{\max}^2 = V_o e \dots\dots\dots(18.4)$$

where m is mass, v is velocity and e is the charge on electron. If the experiment is repeated with light beam of higher intensity, the amount of current increases but the current stops for the same value of V_o . Figure 18.2 shows two curves of photoelectric current as a function of potential V where $I_2 > I_1$. However, if the intensity is kept constant and experiment is performed with different frequencies of incident light, we obtain the curves as shown in Fig.18.3. The current is same but stopping potential is different for each frequency of incident light, which indicates the proportionality of maximum kinetic energy with frequency of light f .

The important results as indicated by the graph between $(K.E.)_{\max}$ and frequency of incident light of the experiments are:

1. The electrons are emitted with different energies. The maximum energy of photoelectrons depends on the particular metal surface and the frequency of incident light.
2. There is a minimum frequency called threshold frequency f_o below which no electrons are emitted, however intense the light may be. For example, blue light

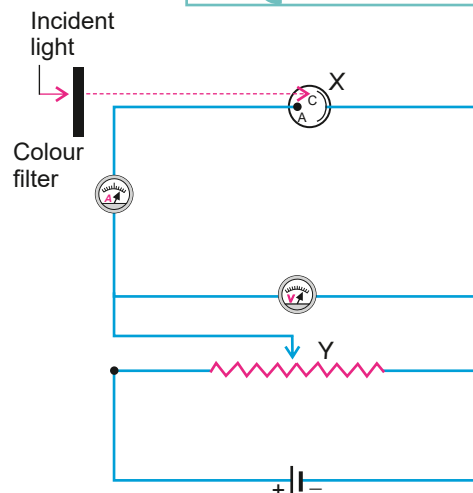


Fig. 18.1: Experimental arrangement to observe photoelectric effect

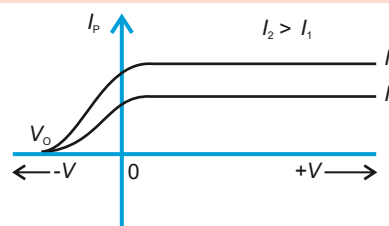


Fig. 18.2: Characteristic curves of photocurrent vs applied voltage for two intensities of monochromatic light.

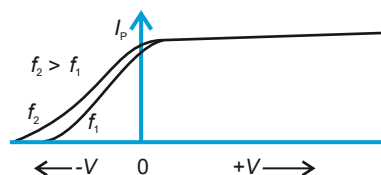


Fig. 18.3: Characteristic curves of photocurrent vs applied voltage for light of different frequencies.

produces photoelectric emission in sodium metal but red light does not (Fig. 18.4).

This threshold frequency f_0 varies from metal to metal. The corresponding wavelength of incident light is called threshold wavelength λ_0 .

By equation; $c = f\lambda$ or $c = f_0\lambda_0$.

- Electrons are emitted instantaneously; the intensity of light determines only their number.

These results could not be explained on the basis of electromagnetic wave theory of light. According to this theory, increasing the intensity of incident light should increase the *K.E.* of emitted electrons which contradicts the experimental result. The classical theory cannot also explain the threshold frequency of light.

According to this theory, even the light of lesser energy should eventually transfer enough energy to liberate electrons. But this does not happen.

Explanation on the Basis of Quantum Theory

Einstein extended the idea of quantization of energy proposed by Max Planck that light is emitted or absorbed in quanta, a tiny packets of energy, known as photons. The energy of each photon of frequency f as given by quantum theory is:

$$E = hf$$

A photon could be absorbed by a single electron in the metal surface. The electron needs a certain minimum energy called the work function ϕ to knockout from the metal surface. If the energy of incident photon is sufficient, the electron is ejected instantaneously from the metal surface. Apart of the photon energy is used by the electron to break

away from the metal and the rest appears as the kinetic energy of the electron. That is;

$$\text{Incident photon energy} - \text{Work function} = (K.E.)_{\max} \text{ of photoelectron}$$

$$\text{or} \quad hf - \phi = \frac{1}{2}mv_{\max}^2 \quad \dots\dots\dots (18.5)$$

This is known as Einstein's photoelectric equation.

When $(K.E.)_{\max}$ of the photoelectron is zero, the frequency f is equal to threshold frequency f_0 , hence, Eq. 18.5 becomes:

$$hf - \phi = 0 \quad \text{or} \quad \phi = hf \quad \dots\dots\dots (18.6)$$

Hence, we can also write Einstein's photoelectric equation as:

$$(K.E.)_{\max} = hf - hf_0 \quad \dots\dots\dots (18.7)$$

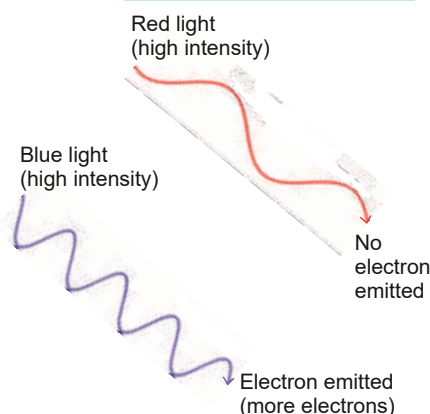
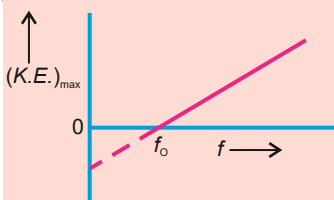


Fig. 18.4: Red and blue light incident on sodium metal

For your information



A graph of the maximum kinetic energy of photoelectrons vs frequency. Below a certain frequency f , no photoelectrons are emitted.

It is to be noted that all the emitted electrons do not possess the maximum kinetic energy, some electrons come straight out of the metal surface and some lose energy in atomic collisions before coming out. Equation 18.7 holds good only for those electrons which come out with full surplus energy.

Albert Einstein was awarded Nobel Prize in physics in 1921 for his explanation of photoelectric effect.

The phenomenon of photoelectric effect cannot be explained if we assume that light consists of waves and energy is uniformly distributed over its wavefront. It can only be explained by assuming light consists of small packet of energy known as photons. Thus, it shows the particle nature of light.

Example 18.1 Find the energy of a photon in eV of the blue light of wavelength 450 nm.

Solution

$$\text{Wavelength } \lambda = 450 \text{ nm} = 450 \times 10^{-9} \text{ m}$$

$$h = 6.63 \times 10^{-34} \text{ J s}$$

Using the formula; $E = hf = \frac{hc}{\lambda}$

$$\begin{aligned} E &= \frac{6.63 \times 10^{-34} \text{ J s} \times 3 \times 10^8 \text{ m s}^{-1}}{450 \times 10^{-9} \text{ m}} \\ &= 4.42 \times 10^{-19} \text{ J} \end{aligned}$$

$$\text{As } 1.6 \times 10^{-19} \text{ J} = 1 \text{ eV} \quad \text{or } 1 \text{ J} = \frac{1}{1.6 \times 10^{-19}} \text{ eV, therefore,}$$

$$E = \frac{4.42 \times 10^{-19}}{1.6 \times 10^{-19}} \text{ eV}$$

$$E = 2.76 \text{ eV}$$

Example 18.2 Yellow light of 577 nm wavelength is incident on a cesium metal surface. The stopping voltage is found to be 0.25 V. Find:

- the maximum *K.E.* of photoelectrons
- the work function of cesium

Solution

- Given data is

$$\text{Wavelength } \lambda = 577 \text{ nm} = 577 \times 10^{-9} \text{ m}$$

$$\text{Stopping voltage } V_o = 0.25 \text{ V}$$

$$\begin{aligned} \text{As } (K.E.)_{\max} &= V_o e \\ &= 0.25 \text{ V} \times 1.6 \times 10^{-19} \text{ C} \\ &= 4 \times 10^{-20} \text{ J} \end{aligned}$$

$$\begin{aligned} \text{Hence } (K.E.)_{\max} &= \frac{4 \times 10^{-20}}{1.6 \times 10^{-19}} \text{ eV} \\ &= 0.25 \text{ eV} \end{aligned}$$

(b) Using photoelectric equation:

$$hf - \phi = \frac{1}{2}mv_{\max}^2$$

$$\text{or } \phi = \frac{hc}{\lambda} - \frac{1}{2}mv_{\max}^2$$

Putting the values

$$\begin{aligned}\phi &= \frac{6.63 \times 10^{-34} \text{ J s} \times 3 \times 10^8 \text{ m s}^{-1}}{577 \times 10^{-9} \text{ m}} - 4 \times 10^{-20} \text{ J} \\ &= 3.45 \times 10^{-19} \text{ J} - 4 \times 10^{-20} \text{ J} = 3.05 \times 10^{-19} \text{ J}\end{aligned}$$

$$\text{As } 1.6 \times 10^{-19} \text{ J} = 1 \text{ eV, therefore, } 1 \text{ J} = \frac{1}{1.6 \times 10^{-19}} \text{ eV}$$

$$\text{Hence } \phi = \frac{3.05 \times 10^{-19}}{1.6 \times 10^{-19}} \text{ eV}$$

$$\phi = 1.9 \text{ eV}$$

For your information

Interaction of electromagnetic radiation with matter

- (i) At low energies (less than 0.51 MeV), the dominant process is photoelectric effect.
- (ii) At intermediate energies, the dominant process is Compton effect.
- (iii) At higher energies (more than 1.02 MeV), the dominant process is pair production.

18.3 COMPTON EFFECT

Arthur Holly Compton at Washington university in 1923 studied the scattering of X-rays by loosely bound electrons from a graphite target (Fig. 18.5). He measured the wavelength of X-rays scattered at an angle θ with the original direction. He found that wavelength λ_s , of the scattered X-rays is greater than the wavelength λ_i of the incident X-rays. This is known as Compton effect. The increase in wavelength of scattered X-rays could not be explained on the basis of classical wave theory. In order to explain this effect, Compton suggested that X-rays consist of photons with energy hc/λ and momentum h/λ .

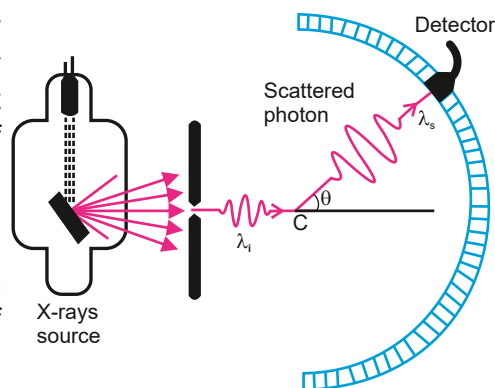


Fig. 18.5: Compton's scattering experiment

Compton Shift

Let us derive an expression for change in wavelength $\Delta\lambda$ known as Compton shift in wavelength.

Consider scattering of a photon by an electron in which the photon suffers collision with the electron like a billiard ball. Figure 18.6 shows the collision between an X-ray photon and stationary electron.

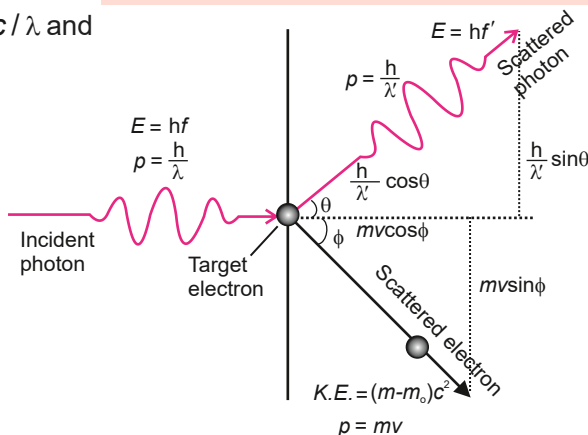


Fig. 18.6: Collision between X-ray photon and stationary electron

The incident X-ray photon of energy ($hf = hc/\lambda$) strikes an electron, initially at rest. On collision, the photon loses some energy which is taken up by the electron. The photon is scattered at an angle θ with the original direction with a smaller energy ($hf' = hc/\lambda'$) and the electron recoils at an angle ϕ , with the original direction of the photon, with *K.E.* equal to $(m - m_0)c^2$, where m is the mass of electron when it is moving with velocity v and m_0 is its rest mass. The difference of the total energy of the electron after collision (mc^2) and before collision, (m_0c^2) is equal to the *K.E.* of the electron.

Actually, this is simple problem of the elastic collision of two balls. Both energy and momentum are conserved. So, by law of conservation of energy:

Change in photon energy = *K.E.* of electron

$$hf - hf' = mc^2 - m_0c^2$$

After solving the above, referring to energy and also momentum conservation, we have the relation for the change in wavelength of the electron. i.e., the expression for Compton shift ($\Delta\lambda$) is:

$$\Delta\lambda = \frac{h}{m_0c} (1 - \cos\theta) \quad \dots\dots\dots (18.8)$$

The factor h/m_0c has dimensions of wavelength and is called Compton wavelength and has the numerical value:

$$\begin{aligned} \frac{h}{m_0c} &= \frac{6.63 \times 10^{-34} \text{ J s}}{9.1 \times 10^{-31} \text{ kg} \times 3 \times 10^8 \text{ m s}^{-1}} \\ &= 2.43 \times 10^{-12} \text{ m} \end{aligned}$$

If the scattered X-ray photons are observed at $\theta = 90^\circ$, the Compton shift $\Delta\lambda$ equals the Compton wavelength. The Eq. 18.8 was found to be in complete agreement with Compton's experimental result, which again is a striking confirmation of particle like interaction of electromagnetic waves with matter.

Arthur Holly Compton was awarded Nobel Prize in physics in 1927 for his discovery of the effect named after him.

We have already discussed in previous class, another kind of very high energy photon such that of a γ -rays with matter. It is pair production in which photon energy is changed into electron-positron pair given by Einstein energy equation.

Energy of photon = Energy need for pair-production + *K.E.* of the particles

$$hf = 2m_0c^2 + K.E. (e^-) + K.E. (e^+) \quad \dots\dots\dots (18.9)$$

The converse of pair-production is known as annihilation of matter in which a particle and its antiparticle interact and annihilate into high energy photons in the γ -rays range along with some other elementary particles.

18.4 WAVE NATURE OF PARTICLES

It has been observed that light displays a dual nature, it acts as a wave and also as a particle. Assuming symmetry in nature, the French physicist, Louis de Broglie proposed in 1924 that particles should also possess wave-like properties. As momentum p of photon is given by Eq. (18.3) as:

$$p = \frac{h}{\lambda}$$

de Broglie suggested that momentum of a material particle of mass m moving with velocity v should be given by the same expression. Thus,

$$p = \frac{h}{\lambda} = mv$$

$$\text{or} \quad \lambda = \frac{h}{p} = \frac{h}{mv} \quad \dots\dots\dots(18.10)$$

where λ is the wavelength associated with particle waves. Hence, an electron can be considered to be a particle and a wave. Equation 18.10 is called de Broglie relation.

An object of large mass and ordinary speed has such a small wavelength that its wave effects such as interference and diffraction are negligible. For example, a rifle bullet of mass 20 g flying with speed 330 m s⁻¹ will have a wavelength λ given by

$$\begin{aligned} \lambda &= \frac{h}{mv} = \frac{6.63 \times 10^{-34} \text{ J s}}{2 \times 10^{-2} \text{ kg} \times 330 \text{ m s}^{-1}} \\ &= 1 \times 10^{-34} \text{ m} \end{aligned}$$

This wavelength is so small that it is not measurable or detectable by any of its effects.

On the other hand, for an electron moving with a speed of $1 \times 10^6 \text{ m s}^{-1}$:

$$\begin{aligned} \lambda &= \frac{6.63 \times 10^{-34} \text{ J s}}{9.1 \times 10^{-31} \text{ kg} \times 1 \times 10^6 \text{ m s}^{-1}} \\ &= 7 \times 10^{-10} \text{ m} \end{aligned}$$

This wavelength is in the X-rays range. Thus, diffraction effects for electrons are measurable whereas interference effects for bullets are not.

Example 18.3 Find the wavelength of two photons produced when a positron-electron pair annihilates. The rest mass energy of each photon is 0.51 MeV.

Solution

$$E = 0.51 \text{ MeV}$$

$$\begin{aligned} E &= 0.51 \times 10^6 \text{ eV} \times 1.6 \times 10^{-19} \text{ J/eV} \\ &= 8.16 \times 10^{-14} \text{ J} \end{aligned}$$

Wavelength $\lambda = ?$

As $E = hf = \frac{hc}{\lambda}$

or $\lambda = \frac{hc}{E}$

Putting the values

$$\lambda = \frac{6.63 \times 10^{-34} \text{ J s} \times 3 \times 10^8 \text{ m s}^{-1}}{8.16 \times 10^{-14} \text{ J}}$$

$$\lambda = 2.43 \times 10^{-12} \text{ m} = 2.43 \text{ pm}$$

Davisson and Germer Experiment

A convincing evidence of the wave nature of electrons was provided by Clinton J. Davisson and Lester H. Germer. They showed that electrons are diffracted from metal crystals in exactly the same manner as X-rays or any other wave. The apparatus used by them is shown in Fig. 18.7, in which electrons from heated filament are accelerated by an adjustable applied voltage. The electron beam is then made incident on a nickel crystal. The scattered electrons from metal surface came off in regular peaks. They interpreted these peak pattern as a result of diffraction just like X-rays diffraction by NaCl crystal. The wavelength of the waves associated with electrons was found to be just that predicted by de Broglie.

The electron beam of energy Ve is made incident on a nickel crystal. The beam diffracted from crystal surface enters a detector and is recorded as a current I . The gain in $K.E.$ of the electron as it is accelerated by a potential V in the electron gun is given by

$$\frac{1}{2}mv^2 = Ve$$

or $mv^2 = 2Ve$

or $m^2v^2 = 2mVe$

or $mv = \sqrt{2mVe}$

From de Broglie equation: $\lambda = \frac{h}{mv}$

Thus $\lambda = \frac{h}{\sqrt{2mVe}} \dots\dots\dots (18.11)$

In one of the experiments, the accelerating voltage V was 54 volts, hence,

$$\lambda = \frac{h}{\sqrt{2mVe}}$$

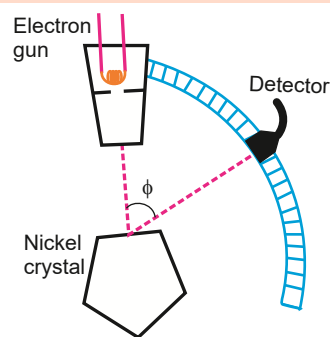


Fig.18.7: Experimental arrangement of Davisson and Germer for electron diffraction

$$= \frac{6.63 \times 10^{-34} \text{ J s}}{\sqrt{2 \times 9.1 \times 10^{-31} \text{ kg} \times 54 \text{ J C}^{-1} \times 1.6 \times 10^{-19} \text{ C}}}$$

$$\lambda = 1.66 \times 10^{-10} \text{ m}$$

This beam of electrons diffracted from crystal surface was obtained for a glancing angle of 65° . According to Bragg's equation;

$$2d \sin \theta = m\lambda$$

For 1st order diffraction $m = 1$

For nickel $d = 0.91 \times 10^{-10} \text{ m}$

Thus $2 \times 0.91 \times 10^{-10} \text{ m} \times \sin 65^\circ = \lambda$, which gives

$$\lambda = 1.65 \times 10^{-10} \text{ m}$$

Thus, experimentally observed wavelength is in excellent agreement with theoretically predicted wavelength.

Diffraction patterns have also been observed with protons, neutrons, hydrogen atoms and helium atoms, thereby giving substantial evidence for the wave nature of particles.

For his work on the dual nature of particles, Prince Louis Victor de Broglie received the 1929 Nobel Prize in physics. Clinton Joseph Davisson and George Paget Thomson shared the Nobel Prize in 1937 for their experimental confirmation of the wave nature of particles.

18.5 WAVE-PARTICLE DUALITY

Interference and diffraction of light provide evidence for its wave nature, while photoelectric effect and Compton effect prove the particle nature of light. Similarly, the experiments of Davisson and Germer and G. P. Thomson reveal wave-like nature of electrons and in the experiment of J. J. Thomson to find e/m , we had to assume particle-like nature of the electron. In the same way, we have to assume both wave-like and particle-like properties for all matter: electrons, protons, neutrons, molecules, etc. and also light, X-rays, γ -rays, etc. have to be included in this. In other words, matter and radiation have a dual 'wave-particle' nature and this new concept is known as wave-particle duality. Niels Bohr pointed out in stating his principle of complementarity that both wave and particle aspects are required for the complete description of both radiation and matter. Both aspects are always present and either may be revealed by an experiment. However, both aspects cannot be revealed simultaneously in a single experiment. The two aspects of light are different in nature that light shows to experimenters. A particular aspect is determined by the nature of the experiment being done. If we put a diffraction grating in the path of a light beam, we reveal it as a wave. If we allow the light beam to hit a metal surface, we need to regard the beam as a stream of particles to explain our observations. There is no simple experiment that we can carry

out with the beam that will require us to interpret it as a wave and as a particle at the same time. Light behaves as a stream of photons when it interacts with matter and behaves as a wave in traveling from a source to the place where it is detected. Thus, the duality of light had to be accepted as a fact of life. In effect, all micro-particles (electrons, protons, photons, atoms, etc.) propagate as if they were waves and exchange energies as if they were particles - that is referred to as the wave-particle duality.

18.6 ELECTRON MICROSCOPE

Electron microscope uses the wave nature of electrons whose de Broglie wavelength is thousand times shorter than the wavelength of visible light. Which enables the electron microscope to see fine details not visible to the optical microscope. The electrons from a source are passed through the electromagnetic lens called condenser which focuses as the beam on the specimen. The beam is accelerated on applying the voltages of several megavolts. The higher is the speed of electrons, the shorter is the wavelength and hence, higher is the resolution. An electromagnetic lens (objective lens) form the image of specimen. The image is further magnified by the electromagnetic projector lens to photograph the image on a film, called micrograph.

A three dimensional image of remarkable quality can be achieved by modern versions called scanning electron microscope.

As stated above, electron microscope can see much smaller details down to nanometre i.e., about the atomic levels of various materials. The uses include study of cell structure, viruses and bacteria in microbiology, investigation of metal fractures, materials and crystal structure.

18.7 ATOMIC SPECTRA

The evidence of particle nature of quantized photons using spectrometer is the basis of atomic spectra. The experimental arrangement consists of a discharge tube, spectrometer and diffraction grating as shown in Fig. 18.9.

Do you know?

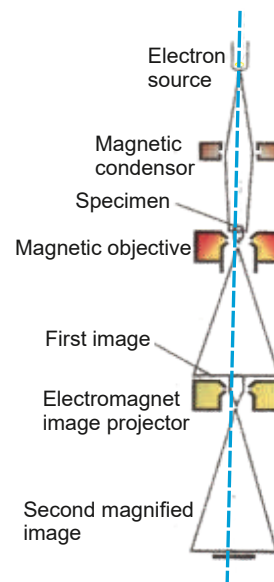
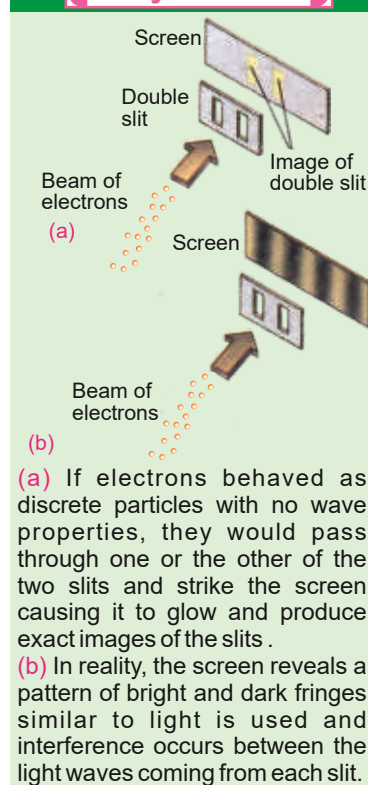


Fig. 18.8: A schematic diagram of the transmission electron microscope.

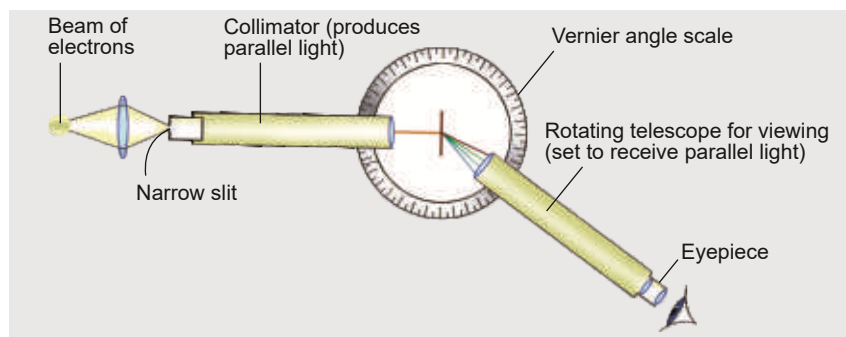


Fig. 18.9: Spectrum produced with a diffraction grating

A series of lines are viewed in dark background through the eyepiece of the spectrometer. The vapours of different elements exhibit different patterns of such lines called spectral series which can be used to identify different elements. One such series was identified by J. J. Balmer in 1885 in the spectrum of atomic hydrogen shown in Fig. 18.10. It is in the visible region of the electromagnetic spectrum consisted of four spectral lines of wavelength. Its results were expressed by J.R. Rydberg in the following mathematical formula:

$$\frac{1}{\lambda} = R_H \left(\frac{1}{2^2} - \frac{1}{n^2} \right)$$

where n is the discrete excited energy level of the electrons with values 3, 4, 5, ... and R_H is the Rydberg constant. Its value is $1.0974 \times 10^7 \text{ m}^{-1}$. In fact, spectral lines of atomic hydrogen extends in the invisible ultraviolet and infrared regions. In the ultraviolet region, the Lyman series contains the wavelengths given by the formula:

$$\frac{1}{\lambda} = R_H \left(\frac{1}{1^2} - \frac{1}{n^2} \right) \dots\dots\dots (18.12)$$

where $n = 2, 3, 4, \dots\dots\dots$

In the infrared region, the wavelengths of spectral lines are given by

$$\frac{1}{\lambda} = R_H \left(\frac{1}{3^2} - \frac{1}{n^2} \right) \dots\dots\dots (18.13)$$

where $n = 4, 5, 6, \dots\dots\dots$

This series is known as Paschen series. Energy level diagram is given in Fig. 18.11 and different transitions are shown in Fig. 18.12.

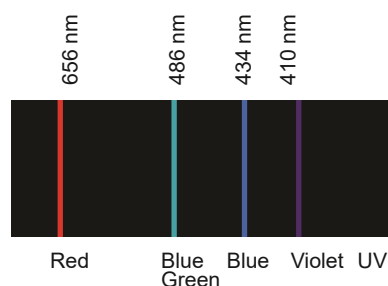
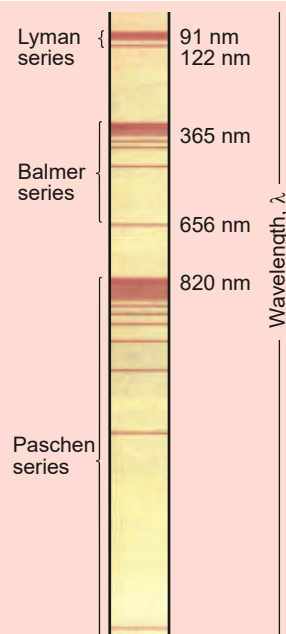


Fig. 18.10

For your information



Line spectrum of atomic hydrogen. Only the Balmer series lies in the visible region of the electromagnetic spectrum.

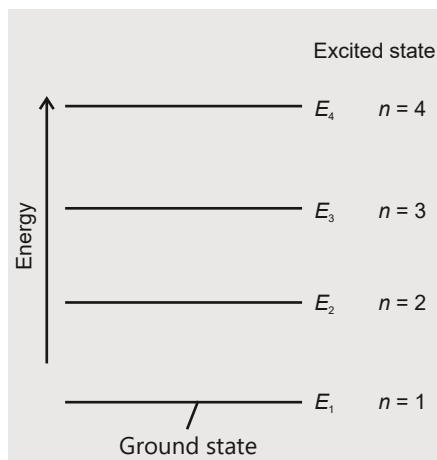


Fig. 18.11: Energy levels

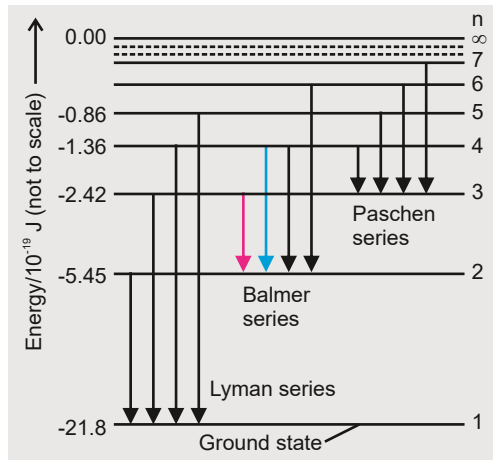


Fig. 18.12: Different transitions in the hydrogen atom

The above mentioned equations give the allowed energy value of the electron in the isolated atoms of hydrogen. The state for $n = 1$ is said to be the ground state whereas the state $n = 2, 3, 4, \dots$ are called excited states. An atom absorbs energy in discrete amounts only and is raised to one of its excited states. The excited states have very short life. The electron soon returns to lower energy level by emitting photons observed as emission spectrum consisting of several spectral lines.

Absorption and Emission Spectra

When light with a continuous spectrum such as white light is passed through a gas at low pressure, and spectrum of light is then analysed, it is found that light of certain wavelengths is missing. In their places dark lines are seen. This type of spectrum is called absorption spectrum. As the light passes through the gas, electrons absorb discrete amount of energy and make transition to higher energy levels. Only photons of certain energy or frequency are absorbed. The wavelengths of light absorbed correspond exactly to energy needed to make such upward transitions. When these excited electrons return to lower levels, the photons are emitted of specific wavelengths given out in emission

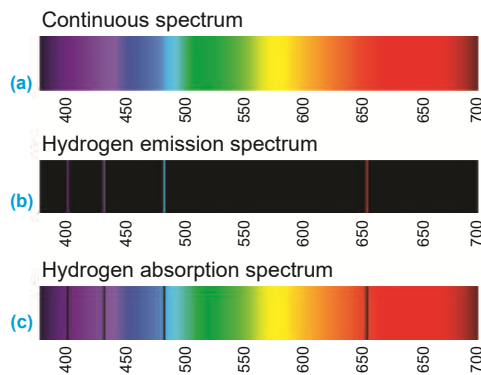


Fig. 18.13: Relation between an absorption spectrum and the emission spectrum of the same element: (a) spectrum of white light, (b) absorption spectrum of element (c) emission spectrum of the same element.

Do you know?

Photon must have energy exactly equal to the energy difference between the two shells for excitation of an atom but an electron with $K.E.$ greater or equal can excite the gas atoms. This transition of energy is equal to; $hf = E_n - E_p$.

For your information

The orbital electrons have specific amount of energies whereas free electrons may have any amount of energy.

spectrum. It means the wavelengths missing from an absorption line spectrum are those present in the emission line spectrum as shown in Fig. 18.13.

18.8 UNCERTAINTY PRINCIPLE

Position and momentum of a particle cannot be measured simultaneously with perfect accuracy. There is always a fundamental uncertainty associated with any measurement. This uncertainty is not associated with the measuring instrument. It is a consequence of the wave-particle duality of matter and radiation. This was first proposed by Werner Heisenberg in 1927 and hence, is known as Heisenberg uncertainty principle. This fundamental uncertainty is completely negligible for measurements of position and momentum of macroscopic objects in daily life but is a predominant fact of life in the atomic domain. For example, a stream of light photons striking a flying tennis ball hardly affects its path, but one photon striking an electron drastically alters its motion. Since light has also wave properties, we would expect to be able to determine the position of the electron only to within one wavelength of the light being used. Hence, in order to observe the position of an electron with less uncertainty, we must use light of short wavelength. But it will alter the motion drastically making momentum measurement less precise. If light of wavelength λ is used to locate a micro-particle moving along x-axis, the uncertainty in its position measurement is: $\Delta x \approx \lambda$

At most, the photon of light can transfer all its momentum h/λ to the micro-particle whose own momentum will then be uncertain by an amount:

$$\Delta p \approx \frac{h}{\lambda}$$

Multiplying these two uncertainties gives:

$$\Delta x \cdot \Delta p \approx \lambda \left(\frac{h}{\lambda} \right) \approx h \dots \dots \dots (18.14)$$

Equation 18.14 is the mathematical form of position-momentum uncertainty principle. It states that:

The product of the uncertainty Δx in the position of a particle at some instant and the uncertainty Δp in the x-component of its momentum at the same instant approximately equals Planck's constant h .

Using uncertainty principle, we can prove that electrons are not present inside atomic nucleus. As the diameter of a nucleus is of the order of 10^{-15} m for an electron to be in the nucleus, the uncertainty in its position would be 10^{-15} m. Using uncertainty principle, the uncertainty in the momentum will be:

$$\Delta p \approx \frac{h}{\Delta x} \approx \frac{6.63 \times 10^{-34} \text{ N s}}{1 \times 10^{-15} \text{ m}} \approx 6.63 \times 10^{-19} \text{ N s}$$

The corresponding energy uncertainty should be of the order of GeV. No such electron in the atom has been found. Thus, electrons do not exist inside the nucleus.

There is another form of uncertainty principle which relates the energy of a particle and

Do you know?

In the sub-atomic world, few things can be predicted with 100% precision.

the time at which it had that energy. i.e.,

$$\Delta E \cdot \Delta t \approx h \text{ (18.15)}$$

Thus, more accurately we determined the energy of a particle, the more uncertainty we will be of the time during which it has that energy.

According to Heisenberg's more careful calculations, he found that at the very best:

$$\Delta x \cdot \Delta p \geq \hbar$$

where $\hbar = \frac{h}{2\pi} = 1.05 \times 10^{-34} \text{ J s}$

and $\Delta E \cdot \Delta t \geq \hbar$

For your information

We can never accurately describe all aspects of subatomic particles simultaneously.

Werner Heisenberg was awarded Nobel Prize in 1932 for his contribution towards quantum physics.

Example 18.4 What can be the velocity of an electron enclosed in a box about the size of an atom of the order of $1.0 \times 10^{-10} \text{ m}$?

Solution

Size of atom $\Delta x = 1.0 \times 10^{-10} \text{ m}$

Mass of electron $m = 9.1 \times 10^{-31} \text{ kg}$

Velocity of electron $\Delta v = ?$

Using the formula: $\Delta p \cdot \Delta x \approx h$

As $\Delta p \approx m \Delta v$

Hence $m \Delta v \cdot \Delta x = h$

$$\Delta v = \frac{h}{m \times \Delta x}$$

Putting the values

$$\Delta v = \frac{6.63 \times 10^{-34} \text{ J s}}{9.1 \times 10^{-31} \text{ kg} \times 1.0 \times 10^{-10} \text{ m}}$$

$$\Delta v = 7.3 \times 10^6 \text{ m s}^{-1}$$

Example 18.5 Find the uncertainty in the energy of a photon which is emitted from an atom radiating for about 10^{-8} second .

Solution

$\Delta t = 10^{-8} \text{ s}, \quad h = 6.63 \times 10^{-34} \text{ J s} \quad \text{and} \quad \Delta E = ?$

Using uncertainty principle:

$$\Delta E \cdot \Delta t \approx h$$

or $\Delta E \approx \frac{h}{\Delta t}$

$$\Delta E \approx \frac{6.63 \times 10^{-34} \text{ J s}}{10^{-8} \text{ s}} \approx 6.63 \times 10^{-26} \text{ J}$$

As $1.6 \times 10^{-19} \text{ J} = 1 \text{ eV}$, therefore,

$$\Delta E \approx \frac{6.63 \times 10^{-26}}{1.6 \times 10^{-19}} \text{ eV}$$

$$\Delta E \approx 4 \times 10^{-7} \text{ eV}$$

QUESTIONS

Multiple Choice Questions

Choose the correct answer.

18.1 Which particle is emitted when UV light is made incident on a zinc surface?

- (a) Photon (b) Positron (c) Electron (d) Alpha particle

18.2 The wave nature of electrons is supported by experiments on:

- (a) line spectrum of atoms (b) the production of X-rays
(c) photoelectric effect (d) electron diffraction by crystalline material

18.3 The dimensions of Planck's constant are same as that of:

- (a) momentum (b) torque
(c) gravitational constant (d) angular momentum

18.4 Which of the following has the most energetic photons?

- (a) Light waves (b) Microwaves (c) X-rays (d) Gamma rays

18.5 de-Broglie waves are associated with:

- (a) moving charged particles only (b) moving neutral particles only
(c) all moving particles (d) all particles whether in motion or at rest

18.6 An electron microscope employs the principle:

- (a) electron have a wave nature
(b) electron can be focused by an electric field
(c) electron can be focused by a magnetic field
(d) all of the above

18.7 Electron, proton, neutron and alpha particle, all have the same speed, which particle will have the shortest wavelength?

- (a) Electron (b) Proton (c) Neutron (d) Alpha particle

18.8 The Balmer series is obtained when all the transitions of electrons terminate at:

- (a) $n = 1$ (b) $n = 2$ (c) $n = 3$ (d) $n = 4$

18.9 The spectrum in which different colours are not diffused into each other and are separated by dark spaces is usually known as:

- (a) line spectrum (b) continuous spectrum
- (c) band spectrum (d) absorption spectrum

18.10 Using uncertainty principle, it can be proved that:

- (a) light has particle nature
- (b) light has wave nature
- (c) electron lies out of the nucleus
- (d) there is always uncertainty in measuring accurately energy and momentum for atomic particles

18.11 According to uncertainty principle, in order to observe the position of an electron with greater accuracy, we must use light of:

- (a) longer wavelength (b) shorter wavelength
- (c) single wavelength (d) any wavelength

Short Answer Questions

- 18.1 How energy of a packet is related to its frequency according to quantum theory? Describe briefly.
- 18.2 How does the photoelectric current vary with intensity of incident light on a metal surface?
- 18.3 Which has more energy, a photon of UV radiation or a photon of yellow light? Describe briefly.
- 18.4 Is work function of a metal related to threshold frequency in a photoelectric emission? State briefly.
- 18.5 Explain briefly, why we can observe the wave-like properties of electrons but not for a flying tennis ball.
- 18.6 Radiation with a certain frequency causes electrons to be emitted from the surface of one metal and not from the surface of another metal. Why?
- 18.7 Will high frequency light has more number of electrons than a low frequency light? Describe briefly.
- 18.8 When does light behave as a wave and when does it behaves as consisted of particles?
- 18.9 Name two possible spectral line series in the spectrum of atomic hydrogen. In which region of electromagnetic spectrum each lies?
- 18.10 What is meant by shift, and on which factors does it depend?

Constructed Response Questions

- 18.1 Why do solids give rise to a continuous spectrum while hot vapours emit line spectrum?

- 18.2 Light can emit electrons from a metal surface and light can also be diffracted. Comment on the statement.
- 18.3 Why X-rays have different properties from light even though both originate from the transition of electrons between different energy levels in excited atoms? Describe briefly.
- 18.4 Is energy conserved when an atom emits a photon of light? Describe how.
- 18.5 How can spectrum of hydrogen atom contains many spectral lines when hydrogen atom contains one electron only?
- 18.6 What should be the speed of an electron if it is to be confined in a box of the size of a nucleus?
- 18.7 An incident X-ray on a metal surface is scattered. What is the wavelength, frequency, energy and speed of the scattered X-ray as compared to incident X-ray?
- 18.8 Why does energy and momentum conservation play a key role in driving Compton's effect?
- 18.9 Describe how a line emission spectrum leads to understand the existence of discrete electron energy levels in atoms.
- 18.10 Mass of a photon is considered zero, then how does the photon possess momentum? Describe.

Comprehensive Questions

- 18.1 Explain photoelectric effect, its experimental arrangements and observations. Deduce photoelectric equation; $hf = \phi + \frac{1}{2}mv_{\max}^2$.
- 18.2 Describe de-Broglie waves associated with material particles and discuss wave particle duality.
- 18.3 Explain the appearance and formation of emission and absorption of line spectra from excited atoms. How can they be used to identify the presence of various elements in a mixture of vapours?
- 18.4 What are the similarities and differences between electron microscope and an optical microscope?
- 18.5 State and explain uncertainty principle. What is its significance in particle physics?

Numerical Problems

- 18.1 A photon has wavelength 350 nm. What is its energy in joules and also in eV?
(Ans: 5.68×10^{-19} J, 3.55 eV)
- 18.2 A certain atom has a separation in energy of 3×10^{-18} J between two of its energy levels. What frequency photon is required to make the atom change from the lower state to higher state?
(Ans: 4.5×10^{15} Hz)

- 18.3 What is the wavelength of photon with an energy of 10×10^{-19} J?
(Ans: 1.99×10^{-7} m)
- 18.4 Calculate the threshold frequency for sodium metal whose work function is 2.28 eV and the value of h is 6.63×10^{-34} J s.
(Ans: 5.5×10^{14} Hz)
- 18.5 What can be the longest wavelength for photoelectric emission from tungsten if its work function is 4.54 eV.
(Ans: 273 nm)
- 18.6 An electron is accelerated in an evacuated tube from rest through a potential difference of 550 V. What will be its final momentum? Calculate the wavelength associate with this electron.
(Ans: 1.27×10^{-23} N s, 5.22×10^{-11} m)
- 18.7 The position of a free electron is determined with an uncertainty of 10^{-7} m. What is the uncertainty in its velocity? What would be the position of an electron after one second?
(Ans: 7×10^3 m s $^{-1}$, anywhere within a distance of 7×10^3 m)
- 18.8 A stationary nickel nucleus of mass 9.95×10^{-26} kg emits a photon of energy 1.17 MeV. Find:
(i) the wavelength of the photon and its momentum.
(ii) the speed of the nucleus after the emission of photon.
(Ans: (i) 1.06×10^{-12} m, 6.25×10^{-22} N s (ii) 6.29×10^3 m s $^{-1}$)
- 18.9 The work function of a metal is 3.5 eV. Light of wavelength 450 nm is incident on the surface. Find out whether electrons will be emitted by the photoelectric effect, from the surface.
(Ans: Electrons will not be emitted)
- 18.10 A 90 keV X-ray photon is fired at a carbon target and Compton scattering occurs. Find the wavelength of the incident photon and the wavelength of the scattered photon for scattering angle of (a) 30° (b) 60° .
(Ans: 13.8 pm (a) 14.1 pm (b) 15 pm)

Nuclear and Particle Physics

Students Learning Outcomes

After studying this chapter, the students will be able to:

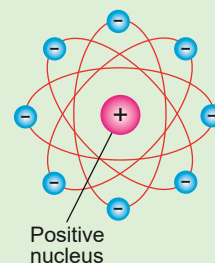
- ◆ recognize the equivalence between energy and mass as represented by $E = mc^2$ and state use of this equation
- ◆ define and use the terms mass defect and binding energy
- ◆ sketch the variation of binding energy per nucleon with nucleon number
- ◆ recall what is meant by nuclear fusion and nuclear fission
- ◆ explain the relevance of binding energy per nucleon to nuclear reactions, including nuclear fusion and nuclear fission
- ◆ explain how the neutrons produced in fission create a chain reaction and that this is controlled in a nuclear reactor [including the action of coolant, moderators and control rods]
- ◆ calculate the energy released in nuclear reactions using $E = \Delta m c^2$
- ◆ explain that fluctuations in count rate provide evidence for the random nature of radioactive decay
- ◆ explain that radioactive decay is both spontaneous and random
- ◆ define activity and decay constant, and state the use of $A = \lambda N$
- ◆ explain half-life with examples
- ◆ use it to solve numerical problems $\lambda = 0.693/T_{1/2}$
- ◆ state the exponential nature of radioactive decay
- ◆ use the relationship $N = N_0 e^{-\lambda t}$ [where N could represent activity, number of undecayed nuclei or received count rate) to solve problems analytically and graphically and N_0 is the initial number of Nuclides]
- ◆ describe the function of the principle components of a water moderated power reactor [core, fuel, rods, moderator, control rods, heat exchange, safety rods and shielding]
- ◆ explain why uranium fuel needs to be enriched before use
- ◆ compare the amount of energy released in a fission reaction with the (given) energy released in a chemical reaction.
- ◆ explain what is a medical tracer [a substance containing radioactive nuclei that can be introduced into the body and is then absorbed by the tissue being studied]
- ◆ explain annihilation reactions [they occur when a particle interacts with its antiparticle and that mass-energy and momentum are conserved in the process]
- ◆ calculate the energy of the gamma ray photons emitted during the annihilation of an electron-positron pair
- ◆ explain that the gamma-ray photons from an annihilation event travel outside the body and can be detected

Ernest Rutherford's experiments in 1911 indicated the existence of a dense, positively charged central part of atom (the nucleus) of very small size surrounded by electrons. In 1920, Rutherford further suggested the positive charge due to protons inside the nucleus and also predicted the presence of another particle having no charge. The prediction came true when neutron was discovered in 1932.

Nuclear physics is the study of various aspects of atomic nuclei and sub-atomic particles. We will particularly discuss here unstable nuclei giving off radiations, their decay process, nuclear reactions such as fission and fusion, benefitting mankind with the release of huge amount of energy. Radioactive isotopes of various elements are used as radioactive tracers for medical diagnostics and treatment, agriculture, scientific research and industry.

An atomic nucleus is represented by the symbol A_ZX often called a nuclide. 'X' represents the chemical symbol of the element, 'Z' the atomic or charge number which indicates the numbers of protons inside the nucleus of an atom. 'A' stands for mass number which indicates the total number of nucleons (protons and neutrons). The number of neutrons 'N' in the nucleus is given by $A - Z$. Thus, we write the elements of hydrogen, carbon and uranium as 1_1H , ${}^{12}_6C$, ${}^{235}_{92}U$ respectively. The nucleus is very dense with a radius of the order 10^{-14} m, surrounded by a cloud of electrons giving the atomic radius of the order 10^{-10} m.

Do you know?



From α -particles scattering experiments, Lord Rutherford concluded that most of the part of an atom is empty and that mass is concentrated in a very small region called nucleus.

19.1 MASS DEFECT AND BINDING ENERGY

The results of experiments on the masses of different nuclei show that the mass of the nucleus is always less than the total mass of all the protons and neutrons making up the nucleus. The lost or missing mass is called mass defect Δm given by the equation:

$$\Delta m = [Zm_p + (A - Z)m_n] - m_{\text{nucleus}} \quad \text{..... (19.1)}$$

where m_p is the mass of a proton, m_n is the mass of a neutron and m_{nucleus} is the mass of the entire nucleus. The missing mass is converted to energy, used in the formation of the nucleus. This energy can be calculated from Einstein's mass-energy relation:

$$\Delta E = (\Delta m)c^2 \quad \text{..... (19.2)}$$

This energy is the potential energy called the binding energy (B.E) of the nucleus. It is defined as:

The energy required to break the nucleus into its constituents (neutrons and protons) is called binding energy.

Binding energy is considered to have negative value similar to absolute gravitational potential value taken as zero on the surface of the Earth. A better measure is the binding energy per nucleon. It can be considered as the average energy needed to separate a

nucleus into its individual nucleons.

We usually use the masses of sub-atomic particles in atomic mass unit (amu or u), which is $\frac{1}{12}$ th of the mass of an unbound neutral atom of $^{12}_6\text{C}$.

$$1 \text{ u} = 1.66 \times 10^{-24} \text{ g} \quad \text{or} \quad 1 \text{ u} = 1.66 \times 10^{-27} \text{ kg}$$

Example 19.1 Find the mass defect and binding energy of the helium nucleus, given that;

$$\text{Mass of proton } m_p = 1.672623 \times 10^{-27} \text{ kg}$$

$$\text{Mass of neutron } m_n = 1.674929 \times 10^{-27} \text{ kg}$$

$$\text{Mass of nucleus of } ^4_2\text{He} = 6.646786 \times 10^{-27} \text{ kg}$$

Solution Using mass defect equation:

$$\Delta m = 2m_p + 2m_n - m_{\text{nucleus}}$$

$$= 2(1.672623) \times 10^{-27} \text{ kg} + 2(1.674929) \times 10^{-27} \text{ kg} - 6.646786 \times 10^{-27} \text{ kg}$$

$$= 6.695104 \times 10^{-27} \text{ kg} - 6.646786 \times 10^{-27} \text{ kg}$$

$$\Delta m = 0.048318 \times 10^{-27} \text{ kg}$$

Using Einstein's Equation:

$$\text{B.E} = \Delta m c^2$$

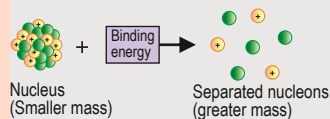
$$= 0.048318 \times 10^{-27} \text{ kg} \times (3.0 \times 10^8 \text{ m s}^{-1})^2$$

$$\text{B.E} = 4.34862 \times 10^{-12} \text{ J}$$

Changing into eV:

$$\text{B.E} = \frac{4.34862 \times 10^{-12}}{1.6 \times 10^{-19}} \text{ eV} = 27 \text{ MeV}$$

For your information



Energy must be supplied to break the nucleus apart into its constituent protons and neutrons.

For your information

Some Atomic Masses

Particle	Mass (u)
e	0.00055
n	1.008665
^1_1H	1.007825
^2_1H	2.014102
^3_1H	3.01605
^3_2He	3.01603
^4_2He	4.002603
^7_3Li	7.016004
$^{10}_5\text{B}$	10.013534
$^{14}_7\text{N}$	14.0031
$^{17}_8\text{O}$	16.9991

Binding Energy per Nucleon

It is useful to draw a graphical curve of binding energy per nucleon against the mass number 'A'. The position of a nucleus on this curve gives information about the stability of the nucleus and whether or not we can get energy by their fission or by fusion. Figure 19.1 shows that initially the hydrogen nucleus with just one proton has a binding energy zero. The fall from hydrogen to helium is large, and suggests a large energy release when hydrogen is converted to helium.

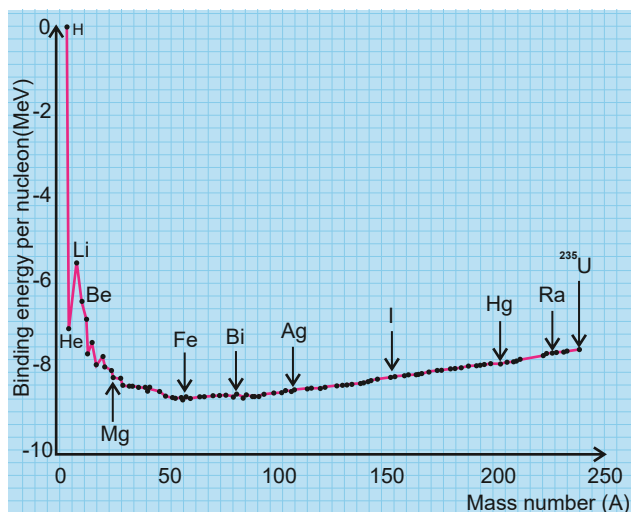


Fig. 19.1: Average binding energy per nucleon against mass number

The curve reaches a maximum value of B.E per nucleon about 8.7 MeV for iron (${}^{56}_{26}\text{Fe}$) which is the most stable nuclide. The region of most stable nuclei is between mass number 'A' equal to 50 and 80. Then, there is a steady decrease in binding energy per nucleon up to the end of the curve.

Thus binding energy per nucleon is less for very light and very heavy nuclei. This is very significant. If a heavy nucleus is split into two nuclei that lie near the lowest part of the curve would be more tightly bound (stable). Such a process we call nuclear fission. The mankind is already benefiting from this process getting huge amount of power production.

Similarly, the nucleons in any pair of neighbouring nuclei on the lighter side of the curve would be more tightly bound (stable) if the pair combine to form a single nucleus. If a deuterium and tritium fuse to form a more stable helium nucleus, huge amount of energy is released. But such a process is not easy to initiate and needs extremely high temperature such as occurring naturally in our Sun and other stars where tremendous amount of energy is released by fusion processes. In fact, energy is obtained from any nuclear reaction in which the binding energy per nucleon of the products increases.

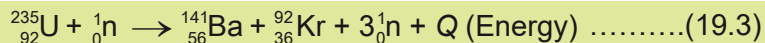
19.2 NUCLEAR FISSION

Otto Hahn and Fritz Strassmann of Germany while working upon the nuclear reactions made a startling discovery. They observed that when slow moving neutrons are bombarded on ${}^{235}_{92}\text{U}$, then as a result of the nuclear reaction ${}^{141}_{56}\text{Ba}$ and ${}^{92}_{36}\text{Kr}$, on average about two or three neutrons are emitted. It may be remembered that the mass of both krypton and barium is less than that of the mass of uranium. This nuclear reaction was different from other nuclear reactions, in two ways:

Firstly, as a result of the breakage of the uranium nucleus, two nuclei of almost equal size are obtained, whereas in the other nuclear reactions the difference between the masses of the reactants and the products was not large. Secondly, a very large amount of energy is given out in this reaction.

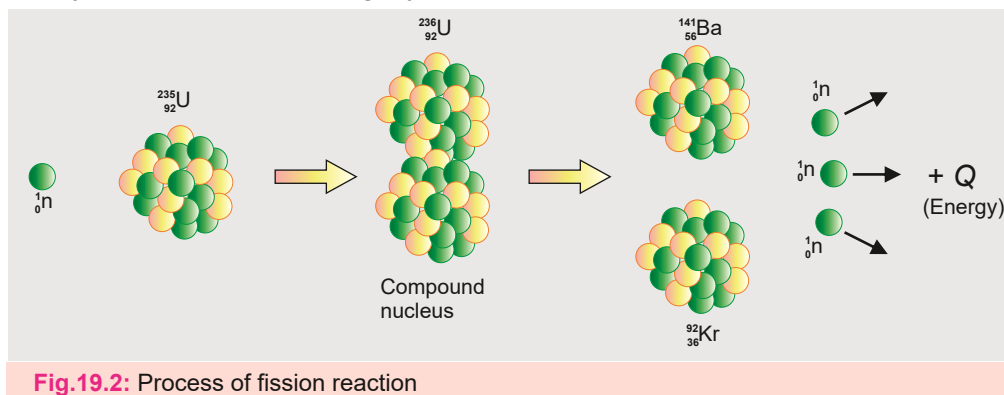
A reaction in which a heavy nucleus splits up into two or more lighter nuclei of roughly equal size along with the emission of energy is called fission reaction.

Fission reaction of ${}^{235}_{92}\text{U}$ can be represented by the equation:



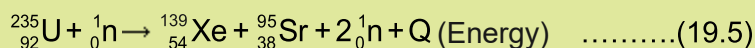
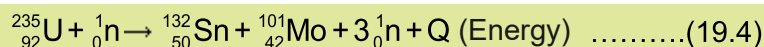
Here Q is the energy given out in this reaction. By comparing the total energy on the left side of the equation with total energy on the right side, we find that in the fission of one uranium nucleus, about 200 MeV energy is given out. It may be kept in mind that there is no difference between the sum of the mass and the charge numbers on both sides of the equation. Fission reaction is shown in Fig. 19.2. Fission reaction can be easily explained with the help of graph of Fig. 19.1. This graph shows that the binding energy per nucleon is greatest for the middle elements of the periodic table and this binding energy per

nucleon is a little less for the light or very heavy elements i.e., the nucleons in the light or very heavy elements are not so rigidly bound.



For example, the binding energy per nucleon for uranium is about 7.6 MeV and the products of the fission reaction of uranium, namely barium and krypton, have binding energy of about 8.5 MeV per nucleon. Thus, when a uranium nucleus breaks up into barium and krypton, as a result of fission reaction, an energy at the rate of $(8.5-7.6) \text{ MeV} = 0.9 \text{ MeV}$ per nucleon is given out. This means that an energy $235 \times 0.9 = 211.5 \text{ MeV}$ is given out in the fission of one uranium nucleus.

The fission process of uranium does not always produce the same fragments (Ba, Kr). In fact, any of the two nuclei present in the upper horizontal part of binding energy could be produced. Two possible fission reactions of uranium are given below as an example:



Hence, in the uranium fission reaction, several products may be produced. All of these products (fission fragments) are radioactive. Fission reaction is not confined to uranium alone; it is possible in many other heavy elements. However, it has been observed that fission takes place very easily with the slow neutrons in uranium-235 and plutonium-239, and mostly these two are used for fission purposes.

Fission Chain Reaction

We have observed that during fission reaction, a nucleus of uranium-235 absorbs a neutron and breaks into two nuclei of almost equal masses besides emitting two or three neutrons. By properly using these neutrons, fission reaction can be produced in more uranium atoms such that a fission reaction can continuously maintain itself.

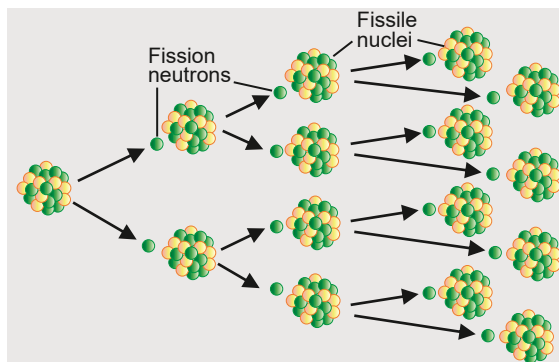


Fig. 19.3: Fission chain reaction

This process is called fission chain reaction. Suppose that we have a definite amount of ${}^{235}_{92}\text{U}$ and a slow neutron originating from any source which produces fission reaction in one atom of uranium. Out of this reaction about three neutrons are emitted. If conditions are appropriate, these neutrons produce fission in some more atoms of uranium. In this way, this process rapidly proceeds and in an infinitesimally small time, a large amount of energy along with huge explosion is produced. Figure 19.3 is the representation of fission chain reaction.

It is possible to produce such conditions in which only one neutron, out of all the neutrons created in one fission reaction, becomes the cause of further fission reaction. The other neutrons either escape out or are absorbed in any other medium except uranium. In this case, the fission chain reaction proceeds with its initial speed. To understand these conditions carefully, look at Fig. 19.4. In Fig. 19.4 (a), a fission reaction in a thin sheet of fissile material (such as enriched uranium-235). The resulting neutrons escape in the air and so they cannot produce any fission chain reaction. Figure 19.4 (b) shows some favourable conditions for chain reaction. Some of the neutrons produced in the first fission reaction produce only one more fission reaction but here also no chain reaction is produced. In Fig. 19.4 (c), a spherical lump of highly fissile heavy material mostly uranium-235 is shown. If the lump is sufficiently big, then most of the neutrons produced by the fission reaction get absorbed in before they escape out of the sphere and produce chain reaction.

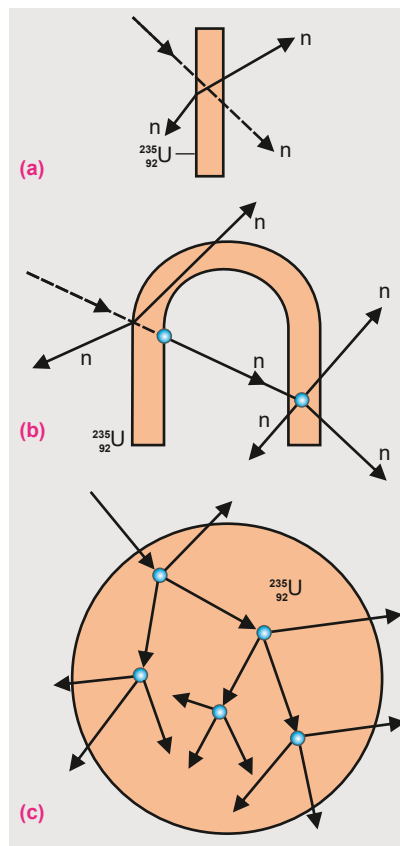
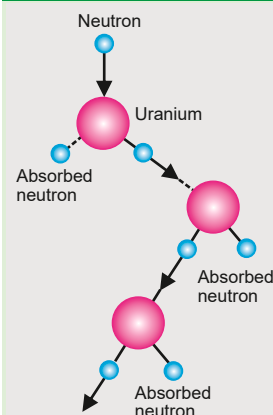


Fig. 19.4

Do you know?



In a controlled chain reaction, only one neutron, on the average, from each fission event causes another nucleus to fission. As a result, energy is released at a steady rate.

The minimum mass of fissile material like (Uranium-235 or Plutonium-239) required to sustain a nuclear chain reaction is called critical mass.

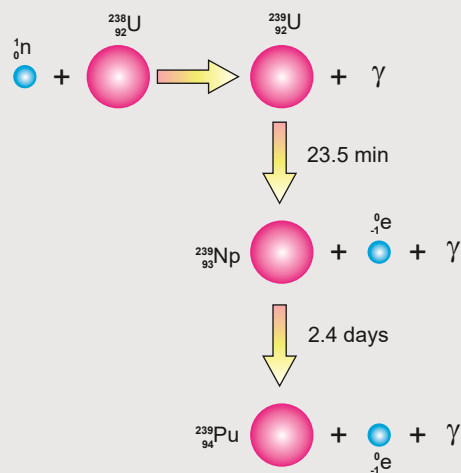
The volume of this mass of uranium is called critical volume. If the mass of uranium is much greater than the critical mass, then the chain reaction proceeds at a rapid speed and a huge explosion or blast is produced. Atom bomb works at this principle. If the mass of uranium is less than the critical mass, the chain reaction does not proceed. If the mass of uranium is equal to the critical mass, the chain reaction proceeds at its initial speed and in this way, we get a source of energy.

Energy in a nuclear reactor is obtained according to this principle. Nuclear reactor works by controlling chain reaction and so in this way, energy can be produced for the human needs like in nuclear power plants. It is to be mentioned that energy released in nuclear fission reaction is much more than chemical reactions in burning of fossil fuels.

19.3 NUCLEAR REACTOR

In a nuclear power station, the reactor plays the same part as does furnace in a thermal power station. In a furnace, coal or oil is burnt to produce heat, while in a reactor fission reaction produces heat. When fission takes place in the atom of uranium or any other heavy atom, then an energy at the rate of about 200 MeV per nucleon is produced. This energy appears in the form of kinetic energy of the fission fragments. These fast moving fragments besides colliding with one another also collide with the uranium atoms. In this way, their kinetic energy gets transformed in heat energy. This heat is used to produce steam which in turn rotates the turbine. Turbine rotates the generator which produces electricity. A sketch of a nuclear power plant is shown in Fig.19.5.

Do you know?



An induced nuclear reaction in which ${}^{238}_{92}\text{U}$ is transmuted into the transuranium element plutonium ${}^{239}_{94}\text{Pu}$ with half-life 24,100 years.

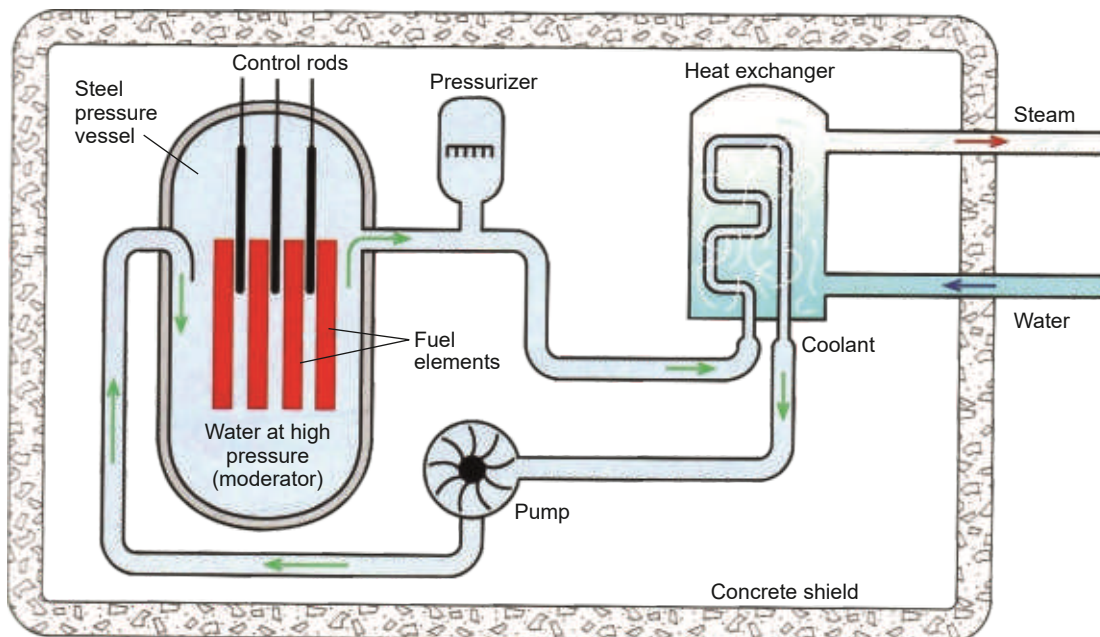


Fig. 19.5: Nuclear power plant

A reactor usually has four important parts. These are:

1. The most important and vital part of a reactor is called core. Here the fuel is kept in the shape of cylindrical tubes. Reactor fuels are of various types. Uranium was used as fuel in first generation reactors. In this fuel, the quantity of $^{235}_{92}\text{U}$ is enriched from 2.4 to 3.0 percent. It may be remembered that the quantity of $^{235}_{92}\text{U}$ in the naturally occurring uranium is only 0.7 percent. Now a days plutonium-239 and uranium-233 are also being used as fuel.
2. The fuel rods are placed in a substance of small atomic weight, such as water or heavy water. They are called moderators. The function of these moderators is to slow down the speed of the neutrons produced during the fission process and to direct them towards the fuel. Heavy water, it may be remembered, is made of ^2_1H , a heavy isotope of hydrogen instead of ^1_1H . The neutrons produced in the fission reaction are very fast and energetic and are not suitable for producing fission in reactor fuel like $^{235}_{92}\text{U}$ or $^{239}_{94}\text{Pu}$ etc. For this purpose, slow neutrons are more useful. To achieve this moderators are used.
3. Besides moderator, there is an arrangement for the control of number of neutrons, so that of all the neutrons produced in fission, only one neutron produces further fission reaction. The purpose is achieved either by cadmium or by boron rods because they have the property of absorbing fast neutrons. The control rods made of cadmium or boron are moved in or out of the reactor core to control the neutrons that can initiate further fission reaction. In this way, the speed of the chain reaction is kept under control. In case of emergency or for repair purposes, control rods are allowed to fall back into the reactor and thus stop the chain reaction and shut down the reactor.
4. Heat is produced due to chain reaction taking place in the core of the reactor. The temperature of the core, therefore, rises to about 500°C . To produce steam from this heat, it is transported to heat exchanger with the help of water, heavy water or any other liquid under high pressure. In the heat exchanger, this heat is used to produce high temperature steam from ordinary water. The steam is then used to run the turbine which in turn rotates the generator to produce electricity. The temperature of the steam coming out of the turbine is up to 50°C . This is further cooled to convert it into water again. To cool this steam, water from some river or sea is generally used.

Do you know?



This symbol is universally used to indicate an area where radioactivity is being handled or artificial radiations are being produced.

In Karachi Nuclear Power Plant (KANUPP), heavy water is being used as a moderator and for the transportation of heat also from the reactor core to heat exchanger, heavy water is used. However, to cool steam coming out of the turbine, sea water is being used.

5. At present, three nuclear power plants are operational in Karachi whereas four at Chashma (Mianwali) in Pakistan with a total output capacity of more than 3,000 MW. Their cooling source is the water from nearby Indus river.

The nuclear fuel once loaded into the reactor can operate it for a few months. There after the fissile material begins to decrease. Now the used fuel is removed and fresh fuel is fed instead. In the used up fuel, intense radioactive substances still remain present. The half-life of these radioactive remnant materials is many thousand years. The radiations and the particles emitted out of this nuclear waste are very injurious and harmful to the living things. Unfortunately, there is no proper arrangement of the disposal of the nuclear waste. This cannot be dumped into oceans or left in any place where they will contaminate the environment, such as through the soil or the air. They must not be allowed to get into the drinking water. The best place so far found to store these wastes is in the bottom of old salt mines, which are very dry and are thousands of metres below the Earth surface. Here they can remain and decay without polluting the environment on the surface of the Earth.

Do you know?

Radioactive wastes are of three types i.e., high level, medium and low level. All these wastes are dangerous for ground water and land environment.

For your information

It is very difficult to dispose off radioactive waste safely due to their long half-lives. For example, 'Pu' half-life is 24,100 years, therefore, it remains dangerous for about 1,82,000 years.

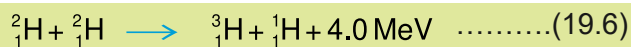
19.4 NUCLEAR FUSION

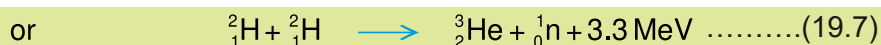
The curve of the graph in Fig. 19.1 also shows that the binding energy per nucleon increases up to $A = 50$. Hence, when two light nuclei merge together to form a heavy nucleus whose mass number A is less than 50, then energy is given out. In the topic on mass defect and binding energy, we have observed that when two protons and two neutrons merge to form a helium nucleus, then about 27 MeV energy is given out.

A nuclear reaction in which two light nuclei merge to form a heavy nucleus is called fusion reaction.

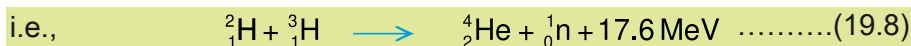
During a fusion reaction some mass is lost and its equivalent energy is given out. In a fusion reaction, more energy per nucleon can be obtained as compared to the fission reaction. But unfortunately, it is more difficult to produce fusion. Two positively charged nuclei must be brought very close to one another. To do so, work has to be done against the electrostatic force of repulsion between the positively charged nuclei. Thus, a very large amount of energy is required to produce fusion reaction. It is true that a greater amount of energy can be obtained during a fusion reaction compared to that produced during a fission reaction, but in order to start this reaction, a very large amount of energy has to be spent.

The probability of occurring fusion of two lighter nuclei is great where one proton or one neutron is produced as given below:





In both of these reactions, about 1.0 MeV energy per nucleon is produced which is equal to the energy produced during fission. If ${}^2_1\text{H}$ and ${}^3_1\text{H}$ are forced to fuse, then 17.6 MeV energy is obtained.



In above reactions ${}^1_1\text{H} + {}^3_1\text{H}$ are termed as fusion fuels. We know that for fusion of two light nuclei, the work has to be done to overcome the repulsive force which exists between them. For this, the two nuclei are accelerated towards one another at a very high speed (Fig. 19.6).

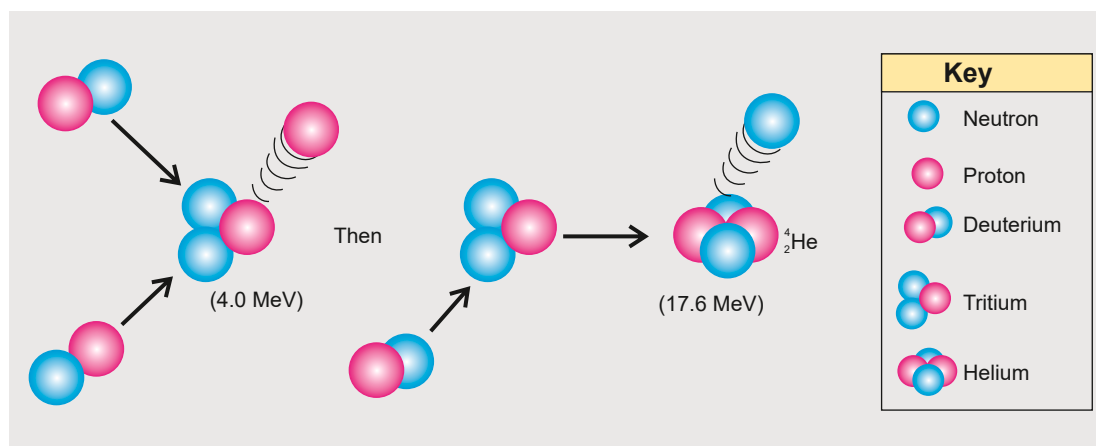


Fig. 19.6: Fusion reaction

This method has been used in the research study of nuclear fusion of ${}^2_1\text{H}$ and ${}^3_1\text{H}$, but cannot be used on a large scale. There is another method to produce fusion reaction. It is based upon the principle that the speed of atoms of a substance increases with the increase in the temperature of that substance. To start a fusion reaction, the temperature at which the required speed of the light nuclei can be obtained is about 15 million kelvin. At such extraordinarily high temperature, the reaction that takes place is called thermonuclear reaction. Ordinarily such a high temperature cannot be achieved. However, during the explosion of an atom bomb this temperature can be had for a very short time.

Until now the fusion reaction has been observed in the experimental testing of hydrogen bomb by triggering it by an atomic bomb. A very large amount of energy can be achieved from a fusion reaction, but till now this reaction has not been brought under control like a fission reaction and so is not being used to produce electricity. Efforts are in full swing in this field and it is hoped that in near future, some method would be found to control this reaction as well. However, the fusion process is taking place in the core of stars including our Sun.

19.5 ACTIVITY AND HALF-LIFE

We have seen that whenever an α or β -particle is emitted from a radioactive element, it is transformed into some other element. This radioactive decay process is quite random and does not follow any regular pattern. Fluctuation in decay curves is also an evidence of its random nature. This means that we cannot foretell about any particular atom as to when will it decay. It could decay immediately or it may remain unchanged for thousands of years. Thus, we cannot say anything about the life of any particular atom of a radioactive element.

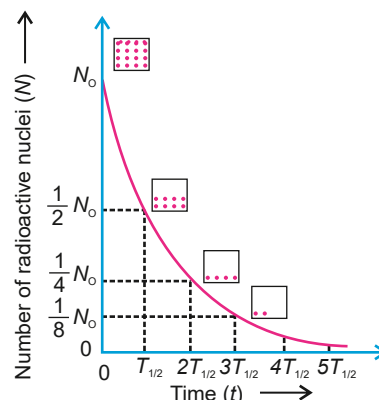


Fig. 19.7: Half-life

While we cannot predict when a single nucleus will decay but in a very large sample the decay follows a pattern or consistency when we draw the number of radioactive nuclei against time (Fig. 19.7). This graph shows that the curve never touches the zero. It means that the radioactivity never ends.

The activity of a radioactive substance is the number of particles it emits per second.

The number of such emissions per second is called decay rate (activity). Its SI unit is Becquerel (Bq).

$$1 \text{ Bq} = 1 \text{ decay per second}$$

The count rate by a nuclear radiation detector is the number of counts recorded per second. It has been found directly proportional to the activity. As time goes, the activity of a source decreases in a consistent manner.

The time taken for the activity of a sample to decrease to half of any starting value is called half-life.

Besides getting the definition of half-life, we can deduce two other conclusions from this example. These are, firstly no radioactive element can completely decay. It is due to the reason that in any half-life period, only half of the nuclei decay and in this way, an infinite time is required for all the atoms to decay.

Secondly, the number of atoms decaying in a particular period is proportional to the number of atoms present in the beginning of the period. If the number of atoms to start with is large then a large number of atoms will decay in this period and if the number of atoms present in the beginning is small then less atoms will decay.

We can represent these results with an equation. If at any particular time, the number of radioactive atoms be N , then in an interval Δt , the number of decaying atom, ΔN is proportional to the time interval Δt and the initial number of atoms N , i.e.,

$$\Delta N \propto -N\Delta t$$

or
$$\Delta N = -\lambda N\Delta t \quad \dots\dots\dots(19.9)$$

where λ is the constant of proportionality and is called decay constant. Equation (19.9) shows that if the decay constant of any element is large, then in a particular interval more of its atoms will decay and if the constant λ is small, then in that very interval less number of atoms will decay. Thus, decay constant of any element is equal to the fraction of the decaying atoms per unit time. The unit of the decay constant is s^{-1} . From Eq. (19.9) we can find activity: $\frac{\Delta N}{\Delta t} = -\lambda N$ As $\frac{\Delta N}{\Delta t} = \text{Activity (A)}$, hence,

$$A = \lambda N \quad \text{..... (19.10)}$$

Eq. (19.10) gives the absolute activity, where number of atoms N decreases with time. The decay ability of any radioactive element is shown by the graphical curve (Fig. 19.7).

We know that every radioactive element decays at a particular rate with time. If we draw a graph between number of atoms in the sample of the radioactive element present at different times, then a curve as shown in Fig. 19.7 will be obtained. This graph shows that in the beginning, the number of atoms present in the sample of the radioactive element were N_0 , with the passage of time the number of these atoms decreased due to their decay. This graph is called decay curve.

After a period of one half-life; $N_0 / 2$ number of atoms of this radioactive element are left behind. If we wait further for another half-life period, then half of the remaining $N_0 / 2$ atoms decay, and $1 / 2 \times N_0 / 2 = (1/2)^2 N_0$ atoms remain behind. After the expiry of further period of a half-life, half of the remaining $(1 / 2)^2 N_0$ atoms decay. The number of atoms that remain un-decayed is $1 / 2 \times (1 / 2)^2 N_0 = (1 / 2)^3 N_0$. We can conclude from this example that if we have N_0 number of any radioactive element, then after a period of n half-lives, the number of atoms left behind is $(1 / 2)^n N_0$.

In a similar way, we can use the general formula for decayed radioactive nuclides:

$$\text{After first half-life} = N_0 - \frac{N_0}{2} = \frac{N_0}{2}$$

$$\text{After second half-life} = N_0 - \frac{N_0}{4} = \frac{3N_0}{4}$$

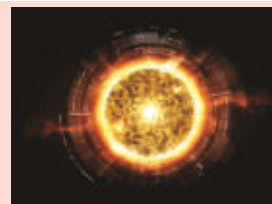
$$\text{After third half-life} = N_0 - \frac{N_0}{8} = \frac{7N_0}{8}$$

$$\begin{array}{ccc} \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot \end{array}$$

$$\text{After } n \text{ half-lives} \quad N = N_0 \left(1 - \frac{1}{2^n} \right) \quad \text{..... (19.11)}$$

It has been found that the estimate of decay of every radioactive element is according to the graph of Fig. 19.7, but the half-life of every radioactive element is different. For example, the half-life of uranium-238 is 4.5×10^9 years while the half-life of radium-226 is 1620 years. The half-life of some radioactive elements is very small, for example, the half-life of radon gas is 3.8 days and that of uranium-239 is 23.5 minutes.

For your information



Artificial Sun has been designed to mimic real Sun. It is an experimental advanced super conducting Tokamak (EAST) project in China. It uses powerful magnetic fields to confine superheated plasma reaching temperature over 100 million $^{\circ}\text{C}$ in an attempt to fuse hydrogen isotopes into helium nucleus releasing clean energy, just like our Sun.

From the above discussion, it is found that the estimate of any radioactive element can be made from its half-life or by determining its decay constant λ . It can be proved with the help of calculus that the following relations exist between the decay constant λ and the number of existing nucleons at any time t .

$$N = N_0 e^{-\lambda t} \quad \dots\dots\dots(19.12)$$

where N_0 is the initial number of nuclides and N is the undecayed nuclei after time t .

When N reduces to $N_0/2$, t is known as half-life $T_{1/2}$. The solution of Eq. (19.12) gives:

$$T_{1/2} = \frac{0.693}{\lambda} \quad \dots\dots\dots(19.13)$$

Thus, if the decay constant λ of any radioactive element is known, its half-life can be found.

Any stable element; besides the naturally occurring radioactive element, can be made radioactive. For this, very high energy particles are bombarded on the stable element. This bombardment excites the nuclei and the nuclei after becoming unstable become radioactive element. Such radioactive elements are called artificial radioactive elements.

Example 19.2 The half-life of Polonium ($^{210}_{84}\text{Po}$) is 140 days. If there are 1,000 atoms of Polonium initially present in a sample, how many of them will decay in 280 days?

Solution

Original atoms of Polonium	= 1000
Half-life of Polonium	= $T_{1/2}$ = 140 days
Number of atoms decaying in 140 days	= $\frac{1000}{2}$
Number of Polonium atoms left	= $1000 - 500 = 500$
Number of atoms decaying in next 140 days	= $\frac{500}{2} = 250$
Total atoms of Polonium decaying in 280 days	= $500 + 250 = 750$ atoms

Example 19.3 Americium-241 is an artificially produced radioactive element that emits α -particles. Its sample of mass 5.1 μg is found to have an activity 5.9×10^5 Bq. Determine, for this sample:

- (i) total number of nuclei (ii) decay constant (iii) half-life in years

Solution

(i) Using Avogadro's number $N_A = 6.02 \times 10^{23}$

$$\begin{aligned} \text{Total number of nuclei} &= \frac{m \times N_A}{M} \\ &= \frac{5.1 \times 10^{-6} \text{ g} \times 6.02 \times 10^{23}}{241} = 1.27 \times 10^{16} \end{aligned}$$

(ii) Using Eq. 19.10:

$$A = \lambda N$$

Putting the values $5.9 \times 10^5 \text{ s}^{-1} = \lambda \times 1.27 \times 10^{16}$

$$\lambda = 4.6 \times 10^{-11} \text{ s}^{-1}$$

(iii) Half-life $T_{1/2} = \frac{0.693}{\lambda}$

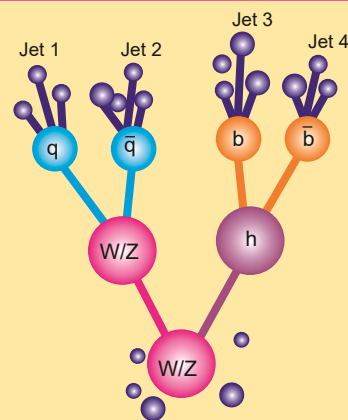
$$= \frac{0.693}{4.6 \times 10^{-11} \text{ s}^{-1}}$$

$$T_{1/2} = 1.49 \times 10^{10} \text{ s}$$

Converting into years; $T_{1/2} = \frac{1.49 \times 10^{10} \text{ s}}{365 \times 24 \times 60 \times 60} \text{ years}$

$$= 472 \text{ years}$$

For your information



Schematic diagram of p^+ and $\bar{p}$ collision

19.6 RADIOACTIVE TRACERS

A radioactive isotope behaves in just the same way as the normal isotope inside a living organism. But the location and concentration of a radioactive isotope can be determined easily by measuring the radiation it emits. Thus, a radioactive isotope acts as an indicator or tracer that makes it possible to follow the course of a chemical or biological process. The technique is to substitute radioactive atoms for stable atoms of the same kind in a substance and then to follow the 'tagged' atoms with the help of radiation detector in the process.

For example, some chemicals such as hydrogen and sodium present in water and food are distributed uniformly throughout the human body. Certain other chemicals are selectively absorbed by certain organs.

Radioisotopes of many elements can be made easily by bombardment with neutrons and other particles. As such isotopes have become available and are inexpensive, their use in medicine, agriculture, scientific research and industries has expanded tremendously.

Radioisotopes are used to find out what happens in many complex chemical reactions and how they proceed. Similarly, in biology, they have helped in investigating the chemical reactions that take place in plants and animals. By mixing a small amount of radioactive isotope with fertilizer, we can easily measure how much fertilizer is taken up by a plant using radiation detector. From such measurements, farmers can know the proper amount of fertilizer to use. Through the use of radiation-induced mutations, improved varieties of certain crops such as rice, chickpea, wheat and cotton have been developed. They have improved plant structure. The plants have shown more resistance to diseases and pest, and give better yield and grain quality. Nuclear Institute

for Agriculture and Biology (NIAB) at Faisalabad doing a great job in this regard.

Uses in Medical Diagnostic and Treatment of Diseases

Tracers are widely used in medicine to study the process of digestion and the way chemical substances move about in the human body.

Radio-iodine, for example, is absorbed mostly by the thyroid gland, phosphorus by bones and cobalt by liver. They can serve as tracers. Small quantity of low activity radioisotope mixed with stable isotope is administered by injection or otherwise to a patient and its location in diseased tissue can be ascertained by means of radiation detectors. For example, radioactive iodine can be used to check that a person's thyroid gland is working properly. A diseased or hyperactive gland absorbs more than twice the amount of normal thyroid gland. Radioactive Iodine-131 is also used to combat cancer of the thyroid gland. A similar method can be used to study the circulation of blood using radioactive isotope sodium-24.

Experiments on cancerous cells have shown that those cells that multiply rapidly absorb more radiation and are more easily destroyed than normal cells by ionizing radiation.

Radioactive tracers in imaging devices have helped in the understanding and diagnosis and treatment of many diseases.

In some cases, radioisotopes in capsules known as 'seed' are implanted in the malignant tissue for local and short ranged treatment. For skin cancer, phosphorous-32 or strontium-90 may be used but the dose has to be carefully controlled to avoid healthy tissues.

The patient undergoing radiation treatment often feel ill as the radiation also damage some healthy tissues.

Some Radioisotopes and their uses

Isotope	Half-life	Example of use
Sodium ^{24}Na	15 hours	Plasma volume
Iron ^{52}Fe	8.3 hours	Imaging bone marrow function
Technetium ^{99}Tc	6 hours	Thyroid uptake scans
Iodine ^{131}I	8.3 days	Kidney tests
Iodine ^{125}I	60 days	Plasma volume Vein flow

19.7 ANNIHILATION REACTIONS

We have already discussed annihilation of positron-electron pair in the previous class. Infact, annihilation is an event taking place whenever a particle and its antiparticle come close to each other. In elementary particles research, very high energies are required for investigations. Instead of collisions at their rest mass energies, they are accelerated to nearly the speed of light and then collided. At CERN's Large Hadron Collider (LHC), proton p^+ antiproton $\bar{p}$ and other heavy nuclei are collided at nearly with the speed of light to create conditions similar to early universe. p^+ , $\bar{p}$ and other heavy particles collision at energies of the order of TeV offer insight to early universe conditions. During such collisions electro-weak force mediation particles W^+ , W^- and Z bosons predicted along with top and bottom quarks and most significant Higgs boson in July, 2012.

Example 19.4 Calculate the energy of the gamma ray photons emitted during the annihilation of an electron-positron pair.

Solution

Mass of positron = 9.1×10^{-31} kg

Mass of electron = 9.1×10^{-31} kg

Using Einstein mass-energy relation:

$$E = 2mc^2$$

$$= 2 \times 9.1 \times 10^{-31} \text{ kg} \times (3 \times 10^8 \text{ m s}^{-1})^2$$

$$E = 1.64 \times 10^{-13} \text{ J}$$

In eV: $E = \frac{1.64 \times 10^{-13}}{1.6 \times 10^{-19}} \text{ eV}$

$$= 1.02 \text{ MeV}$$

Hence, 2γ -ray photons will be emitted each with 0.51 MeV energy travelling in opposite directions to conserve momentum.

For your information

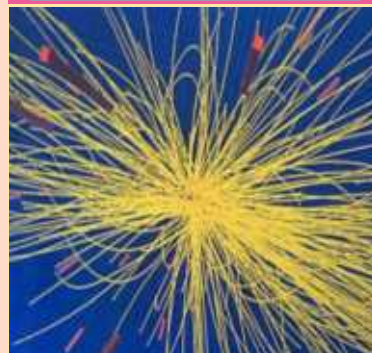


Image of a 7 TeV proton-proton collision in CMS producing more than 100 charged particles.

QUESTIONS

Multiple Choice Questions

Choose the correct answer.

19.1 What is the half-life of a radio nuclide if $1/16$ of its mass is present after 2h?

- (a) 15 min. (b) 30 min. (c) 45 min. (d) 60 min.

19.2 Gamma radiations are emitted due to:

- (a) de-excitation of atoms (b) excitation of atoms
(c) de-excitation of nuclei (d) excitation of nuclei

19.3 Cadmium rods are used in nuclear reactor for:

- (a) slowing down fast neutron (b) speeding up slow neutron
(c) absorbing neutrons (d) starting the nuclear reaction

19.4 The mass of fissionable material needed for self-sustaining chain reaction is called:

- (a) supercritical mass (b) fermi mass
(c) critical mass (d) sub-critical mass

19.5 Graphite and heavy water are two common moderators used in a nuclear reactor. The function of the moderator is:

- (a) to slow down the neutrons to thermal energies
(b) to absorb the neutrons and stop the chain reaction
(c) to cool the reactor
(d) to control the energy released in the reactor

19.6 The mass of a nucleus is always:

- (a) an integral number of proton masses
- (b) equal to the mass of all other nuclei of the same element
- (c) equal to the sum of the masses of proton and neutron
- (d) less than the sum of masses of the constituent particles

19.7 Binding energy per nucleon is:

- (a) greater for heavy nuclei
- (b) least for heavy nuclei
- (c) greater for light nuclei
- (d) greater for medium weight nuclei

19.8 What remains unchanged in a fusion event?

- (a) Energy
- (b) Mass of nucleons
- (c) Number of nucleons
- (d) Temperature

19.9 Radioactive tracers are also employed to follow the path that various chemicals or food constituents take in:

- (a) human bodies
- (b) animals
- (c) plants
- (d) all of these

19.10 The K.E. of the fission fragments is ultimately converted to:

- (a) potential energy
- (b) heat energy
- (c) magnetic energy
- (d) chemical energy

Short Answer Questions

- 19.1 Define binding energy and how is it calculated?
- 19.2 Comment on the statement. "In radioactive decay, if the probability of decay per unit time is doubled, its half-life will also be doubled."
- 19.3 After four half-lives, what percentage of a radioactive sample remains?
- 19.4 Why natural uranium is said to be low grade nuclear fuel? Describe briefly.
- 19.5 Give an application of a radioactive tracers in medicine.
- 19.6 How is energy generated in the Sun? Describe briefly.
- 19.7 What is meant by the half-life of a radioactive material?
- 19.8 What is meant by controlled and uncontrolled fission chain reaction? Describe briefly.

Constructed Response Questions

- 19.1 In what ways a nuclear power station is different to a fossil fuel burning power station also called thermal power plant?
- 19.2 If fusion reactors are developed, what advantage are they likely to have over fission nuclear reactors?
- 19.3 The activity of a radioactive sample is found to decay by 10 percent in a year. What will be the activity after another year? Explain.

- 19.4 A source has an activity of 10×10^4 Bq and a half-life of 20 minutes. What will be the activity of the source after one hour?
- 19.5 What factors make a fusion reaction difficult to achieve?
- 19.6 In what way a nuclear reaction is different from chemical reaction? Explain.
- 19.7 Suggest a reason why neutrons are absorbed in the boron rods and the rods become hot as a result of this reaction.

Comprehensive Questions

- 19.1 Define the terms mass defect and binding energy. What information is given by the graphical curve between binding energy per nucleon against nucleon number?
- 19.2 What is nuclear fission? Explain critical mass and nuclear fission chain reaction.
- 19.3 Describe a nuclear power reactor with its components and their functions.
- 19.4 Describe why the very lighter nuclei and very heavy nuclei are unstable. How are they useful as a source of huge amount of energy?
- 19.5 What are radioactive tracers? Give some applications in medical diagnostics and treatment.
- 19.6 What is half-life? Derive the relation between number of half-lives and time. How can we calculate the remaining and decayed quantity after several half-lives?

Numerical Problems

- 19.1 Find the mass defect and binding energy of the deuterium. The mass of deuteron is 3.3435×10^{-27} kg.
(mass of proton $m_p = 1.6726 \times 10^{-27}$ kg, mass of neutron $m_n = 1.6749 \times 10^{-27}$ kg)
(Ans: 2.23 MeV)
- 19.2 The decay constant of strontium is $1.99 \times 10^{-5} \text{ s}^{-1}$. Find its half-life. (Ans: 9.7 hours)
- 19.3 Half-life of krypton is 3.16 minutes. Out of 40 g of krypton, how much will be left after 12.64 minutes? (Ans: 2.5 g)
- 19.4 A 2048 g (≈ 2 kg) sample of a radioactive material decreases to 32 g in 18 years. What is its half-life? (Ans: 3 years)
- 19.5 An alpha particle has an energy of 5.56 MeV when released from the Radium atom. Determine its velocity. (Using $m_\alpha = 6.64 \times 10^{-27}$ kg). (Ans: $1.67 \times 10^7 \text{ m s}^{-1}$)
- 19.6 Find the energy released during the nuclear reaction: ${}^7_3\text{Li} + {}^1_1\text{H} \rightarrow 2 {}^4_2\text{He}$
Mass of ${}^4_2\text{He} = 4.00388 \text{ u}$, ${}^1_1\text{H} = 1.00814 \text{ u}$ and of ${}^7_3\text{Li} = 7.01823 \text{ u}$.
(Ans: 17.3 MeV)

Medical Physics

Students Learning Outcomes

After studying this chapter, the students will be able to:

- ◆ illustrate how PET-scanning works: Positrons emitted by the decay of the tracer annihilate when they interact with electrons in the tissues, producing a pair of gamma-ray photons travelling in opposite directions.
- ◆ explain the piezoelectric effect and its applications in medical science: Ultrasound waves are generated and detected by a piezoelectric transducer.
- ◆ explain how ultrasound can be used to obtain diagnostic information about internal body structures.
- ◆ explain that X-rays are produced by electron bombardment of a metal target and calculate the energy.
- ◆ explain the use of X-rays in imaging internal body structures, including a description of the term contrast in X-ray imaging.
- ◆ explain how computed tomography (CT) scanning works: It produces a 3D image of an internal structure by first combining multiple X-ray images taken in the same section from different angles to obtain a 2D image of the section, then repeating this process along an axis and combining 2D images of multiple sections.

Medical physics is a branch of physics that applies physical principles and techniques to medicine, especially in the diagnosis and treatment of diseases. It combines knowledge of physics with medical science for the development and use of technologies. This helps doctors to see inside the human body and treat patients safely and effectively. Medical physics plays an important role in modern healthcare technologies by improving the accuracy of medical imaging, reducing radiation risks, and enhancing treatment methods. In our daily life, medical physics is commonly used in hospitals and diagnostic centers through technologies such as X-ray imaging, ultrasound scans, CT-scans, and PET-scans. These techniques allow doctors to detect fractures, diagnose tumors, and to study internal organs without surgery. Thus, medical physics directly contributes to early diagnosis, effective treatment, and improved quality of life. In this chapter, we will learn how advanced imaging techniques are developed using principles of physics. We will study how positron emission tomography (PET) works through positron–electron annihilation and gamma-ray detection. The chapter will also explain the piezoelectric effect and its role in generating and detecting ultrasound waves for medical diagnosis. Furthermore, we will explore the production and use of X-rays.

Introduction to Radiations Used in Medical Physics

Radiation is the transfer of energy through space or matter in the form of waves or particles. In medical physics, radiations are mainly used for diagnosis (imaging) and treatment. Radiations are broadly classified into non-ionizing and ionizing radiations.

Ionizing radiations have enough energy to remove electrons from atoms, producing ions. Non-ionizing radiations do not have enough energy to remove electrons from atoms.

Examples of radiations used in medical physics are X-rays, gamma rays, and positrons. Interaction of X-rays and gamma rays with matter occurs through photoelectric effect and Compton scattering. The positrons are emitted from radioactive isotopes. On the other hand, examples of non-ionizing radiations are ultrasound waves and radiowaves. These waves are reflected, refracted, or absorbed when interact with the matter.

Gamma rays, X-rays, ultrasonic waves, and radiowaves are widely used in medical physics for diagnosis and treatment. X-rays are commonly used to image bones and internal organs, while gamma rays and positrons are used in nuclear medicine, such as Single Photon Emission Computer Tomography, Positron Emission Tomography scans and cancer treatment, due to their high penetrating power. Ultrasonic waves are used in ultrasound imaging to examine soft tissues, monitor fetal diseases, and study internal organs safely without ionizing radiation. Radiowaves are used in techniques such as MRI, where they help produce detailed images of internal body structures by interacting with atomic nuclei, making them especially useful for imaging soft tissues without harmful radiation exposure.

20.1 ULTRASOUND AND PIEZOELECTRIC EFFECT

Ultrasound is widely used non-invasive medical imaging technique that allows visualization of internal body structures such as tissues, organs, and the fetus during pregnancy. This imaging method operates on a physical principle known as the piezoelectric effect, which is central to the working of an ultrasound transducer. The transducer plays a dual role; it generates ultrasound waves and detects the reflected signals from inside the body.

The Piezoelectric Effect

The word “piezoelectric” comes from the Greek word “piezein,” which means “to press”, and “electric” refers to the production of electricity.

This effect was discovered in 1880 by two French brothers, Pierre Curie and Jacques Curie. They were experimenting with crystals such as quartz, tourmaline, and Rochelle salt (sodium potassium tartrate). They found that when mechanical pressure was applied to these crystals, they generated an electric charge. Their experiments proved a strong connection between mechanical force and electrical response in certain crystals. The piezoelectric effect is possible because of the unique internal structure of some crystals. There are two types of piezoelectric effect; direct piezoelectric effect and

Tidbit!

The piezoelectric effect makes it possible to turn pressure into voltage and vice versa. This is how ultrasound “talks” to the body.

converse piezoelectric effect. These crystals have no centre of symmetry, and inside them, there are small regions with positive and negative charges. Normally, these charges are balanced, but when pressure is applied, the balance is disturbed, and the charges move slightly, causing an electric field to appear on the surface. This is called the direct piezoelectric effect. On the other hand, if we apply an electric field to the crystal, it changes shape a little. This is called the converse piezoelectric effect. The working of this effect is shown in Fig. 20.1. A transducer generates ultrasound waves by the converse effect (an applied alternate voltage makes the crystal vibrate) and detects the returning echoes, the direct effect.

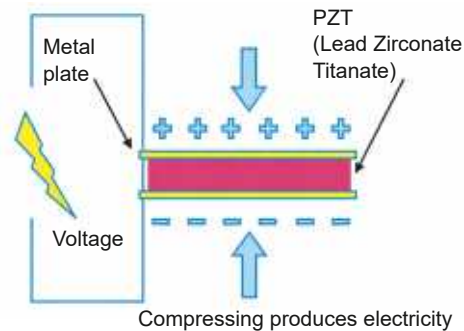


Fig. 20.1: Schematic diagram demonstrating piezoelectric effect

20.2 DIAGNOSTIC USE OF ULTRASOUND IN MEDICAL IMAGING

Ultrasound is a widely used medical imaging technique that allows doctors to view internal body structures in real time. The labelled diagram of ultrasound machine is shown in Fig. 20.2. It operates by sending high frequency sound waves, typically in the range of 1 MHz-15 MHz, into the body using a device called a transducer. These waves travel through the body and reflect off boundaries between different tissues. The reflected echoes are then collected and analyzed to create an image. One of the key advantages of ultrasound is that it does not use ionizing radiation, making it a safe, non-invasive, and repeatable imaging method. The functioning of ultrasound imaging is based on several physical principles. First, the piezoelectric transducer emits pulses of ultrasound waves, which travel at different speeds depending on the type of tissue they pass through. When these waves encounter a boundary between two tissues with different acoustic impedances (determined by tissue density and sound speed), part of

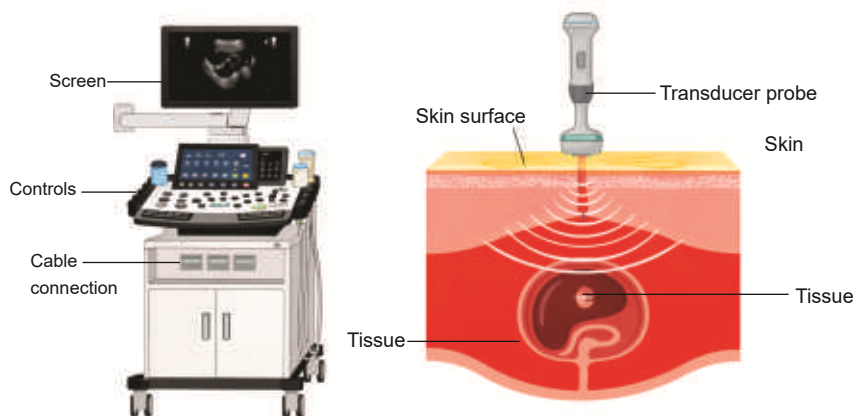


Fig. 20.2: Ultrasound machine

the wave is reflected. The transducer then switches to receive mode and detects these reflected waves using the direct piezoelectric effect. These echoes are converted into electrical signals and processed by the machine. Finally, a computer calculates the time delay and intensity of the returning echoes to construct a visual image. Areas with stronger reflections appear brighter in the image, helping to reveal the internal structures of organs and tissues.

Do you know?

The crystals in an ultrasound probe are smaller than a grain of rice, but they can both send and detect sound waves.

Factors Affecting Image Quality

- **Acoustic Impedance Difference:** A higher difference between two tissues (e.g., soft tissue and bone) causes stronger echoes, producing clear boundaries in the image.
- **Frequency of Ultrasound:** Higher frequency ultrasound offers better resolution but less penetration; lower frequency waves penetrate deeper but have lower image clarity.
- **Use of Gel:** A special gel is applied between the transducer and the skin to remove air gaps and improve sound waves transmission.

Tidbit

Ultrasound can even measure blood flow speed using the Doppler effect detecting shifts in sound frequency.

Applications

Ultrasound imaging is extensively used across various medical specialties for diagnostic purposes. In obstetrics, it helps monitor fetal growth, estimate gestational age, and identify abnormalities during pregnancy. In the abdominal region, ultrasound is used to examine vital organs such as the liver, kidneys, gallbladder, pancreas, and bladder. In cardiology, echo cardiography allows doctors to evaluate the heart's structure and function in real time. The field of urology also benefits from ultrasound for examining the prostate, testicles, and urinary system. Additionally, in the musculoskeletal system, it aids in detecting ligament injuries, joint inflammation, and other soft tissue conditions.

20.3 X-RAYS

The transitions of electrons in hydrogen or other light elements result in the emission of spectral lines in the infrared, visible or ultraviolet region of electromagnetic spectrum due to small energy differences in the transition levels.

In heavy atoms, the electrons are assumed to be arranged in concentric shells labelled as K, L, M, N, O, etc. The K-shell being closest to the nucleus, the L-shell next, and so on (Fig. 20.3). The inner shell electrons are tightly bound, and a large amount of energy is required to knock them out from their normal energy levels. After excitation,

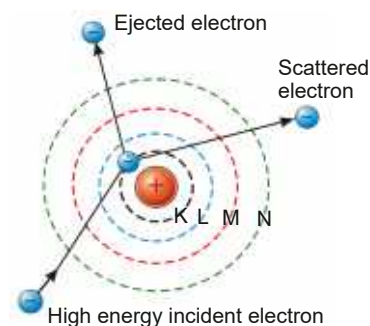


Fig. 20.3: Concentric energy shells

when an atom returns to its normal state, photons of larger energy are emitted. Thus, transition of inner shell electrons in heavy atoms gives rise to the emission of high energy photons or X-rays. These X-rays consist of series of specific wavelengths or frequencies and hence are called characteristic X-rays. The study of characteristic X-rays spectra has played a very important role in the study of atomic structure and the periodic table of elements (Fig. 20.4).

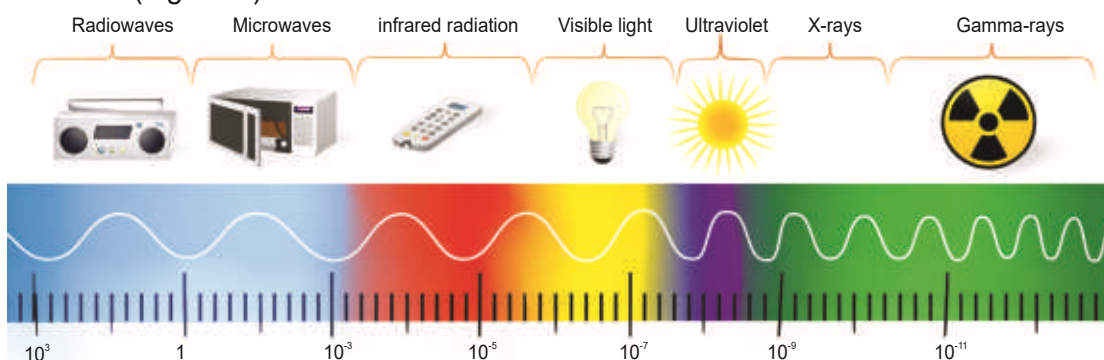


Fig. 20.4: Electromagnetic spectrum

Production of X-rays

Figure 20.5 shows an schematic arrangement of producing X-rays. It consists of a high vacuum tube called X-ray tube. When the cathode C is heated by the filament F, it emits electrons which are accelerated towards the anode T. If V is the potential difference between C and T, the $(K.E.)_{\max}$ with which the electron strikes the target is given by

$$(K.E.)_{\max} = Ve \dots\dots\dots (20.1)$$

Suppose that these fast moving electrons of energy Ve strike a target made of tungsten or any other heavy element. It is possible that in collision, the electrons in the innermost shells, such as K or L, will be knocked out. Suppose that one of the electrons in the K-shell is removed, thereby producing a vacancy or hole in that shell. The electron from the L-shell jumps to occupy the hole in the K-shell, thereby emitting a photon of energy

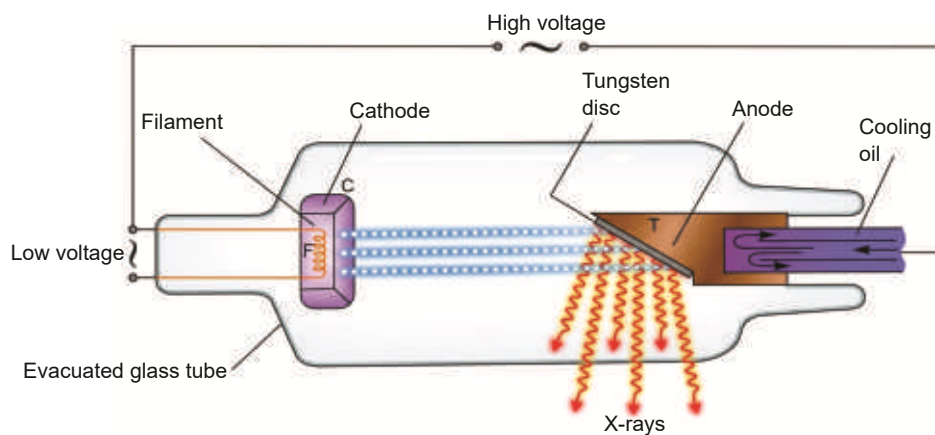


Fig. 20.5: Schematic arrangement for producing X-rays

hf , called the K_α X-ray given by

$$hf_{K_\alpha} = E_L - E_K \dots\dots\dots (20.2)$$

It is also possible that the electron from the M-shell might also jump to occupy the hole in the K-shell. The photons emitted are K_β X-ray with energies hf these photons give rise to K_β X-ray and so on.

$$hf_{K_\beta} = E_M - E_K \dots\dots\dots (20.3)$$

The photons emitted in such transitions i.e., inner shell transitions are called characteristic X-rays, because their energies depend upon the type of target material.

The holes created in the L and M-shells are occupied by transitions of electrons from higher energy states creating more X-rays. The characteristic X-rays appear as discrete lines on a continuous spectrum as shown in Fig. 20.6.

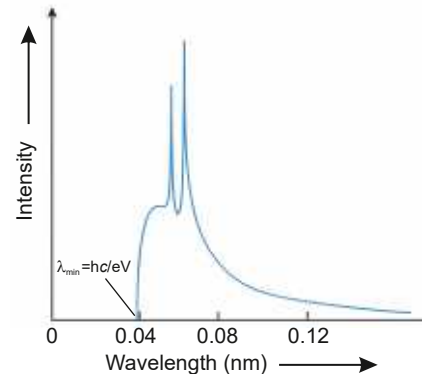


Fig. 20.6: Continuous X-ray spectrum

The Continuous X-ray Spectrum

The continuous spectrum is due to an effect known as bremsstrahlung or braking radiation. When the fast moving electrons bombard the target, they are suddenly slowed down on impact with the target. We know that an accelerating charge emits electromagnetic radiation. Hence, these impacting electrons emit radiation as they are strongly decelerated by the target. Since the rate of deceleration is so large, the emitted radiation correspond to short wavelength and so the bremsstrahlung is in the X-ray region. In the case, when the electrons lose all their kinetic energy in the first collision, the entire kinetic energy appears as a X-ray photon of energy $hf_{\max}$,

$$\text{i.e.;} \quad (K.E.)_{\max} = hf_{\max}$$

The wavelength $\lambda_{\min}$ in Fig. 20.6 corresponds to frequency $f_{\max}$. Other electrons do not lose all their energy in the first collision. They may suffer a number of collisions before coming to rest. This will give rise to photons of smaller energy or X-rays of longer wavelength. Thus, the continuous spectrum is obtained due to deceleration of impacting electrons.

X-Ray Imaging

As X-rays travel through the body, they lose energy because they interact with atoms. This process is called attenuation. X-rays can ionize atoms; it means that they knock out electrons which is how energy is transferred. The intensity of the X-ray beam decreases gradually as it passes through a material. This drop follows an exponential pattern.

Image intensifiers are used in real-time X-ray imaging. They help reduce radiation exposure by brightening the image, making it easier to see. Contrast in X-ray images means how clearly, we can tell different tissues apart. Contrast depends on X-rays energy. Hard X-rays (high energy) are used for bone imaging. Soft X-rays (low energy) are better for soft tissues like the breast. Contrast media to temporarily enhance the

visibility of internal structures, organs and blood vessels are basically special substances like iodine or barium that absorb X-rays well. They are injected or swallowed to highlight certain parts of the body.

20.4 COMPUTED TOMOGRAPHY SCANNING (CT-SCAN)

Computed Tomography, commonly known as CT-scan, is a medical imaging technique that combines X-ray technology with computer processing to generate high resolution cross-sectional and 3D images of internal body structures. It is especially useful for diagnosing diseases, detecting internal injuries, and guiding medical treatment. An ordinary X-ray shows a flat (2D) image of the body, where different parts like bones and organs overlap each other. This makes it hard to see individual structures clearly. To solve this problem, doctors use a special technique called CT-scan (Computed Tomography).

A CT-scanner takes many X-ray images of the same body section from different angles. A computer combines these images to make a detailed cross-sectional slice or image of the body. Then, the scanner moves slightly along the body and repeats the process for the next slice. In the end, all these slices are combined by the computer to create a 3D image of the internal structure.

In a CT-scanner, the X-ray tube rotates around the patient while detectors stay in place. This setup collects complete information from all sides, as shown in Fig. 20.7.



Fig. 20.7: A boy 3D head cross-section CAT-scan

Do you know?

- i. The CT-scanner was invented by Sir Godfrey Hounsfield (England) and Allan Cormack (South Africa) in the early 1970s. For this invention, both scientists were awarded the Nobel Prize in Physiology or Medicine in 1979.
- ii. CT-scans create hundreds of slices or images in seconds like peeling a digital apple layer by layer. CT-scans can detect strokes within minutes of onset crucial for lifesaving treatment.

The name computed tomography comes from:

- **Computed:** A computer processes the data.
- **Tomography:** From the Greek word *tomos*, meaning **slice**.
- **Axial:** The scan is taken along the body's axis.

This way, doctors can get a clear and detailed view of internal organs, bones, or tissues. Figure 20.7 shows a boy having a CT-scan, and the monitor displays a cross-sectional image of the head. This technique is called Computed Axial Tomography (CAT-scan) because it uses a computer to control the scanning process, collect data, and converts it into images. The X-ray tube rotates around the patient's body, capturing images from

different angles. These images are then used to create detailed cross-sectional slices of the body.

Advantages of CT-Scan

CT-scans offer several important advantages over traditional X-rays. While single X-ray images are still useful and quick to produce, they only give a flat, 2D view. In contrast, CAT-scans produce 3D images, allowing doctors to see the size, shape, and exact location of organs, bones, and other structures inside the body. This makes it easier to distinguish between tissues, even when they have very similar densities. For example, a CT-scan can clearly show the position and size of a tumor, which helps doctors accurately target it during treatment using high energy X-rays or gamma rays. However, it is important to remember that CT-scans involve exposure to X-rays, which are a form of ionizing radiation. Although the amount of radiation is relatively low, it still carries a small risk to the patient. On average, the radiation dose from a CT-scan typically delivers a dose comparable to, or a few times greater than one year of natural background radiation depending on the body part scanned, or roughly equal to the radiation dose received during four long-distance flights. Therefore, doctors take special precautions, especially when scanning pregnant women or patients with existing health concerns.

Tidbit!

CT-scanners take multiple X-ray images from different angles to reconstruct a 2D cross-section of the body. When these slices are combined, they create a 3D model, like stacking layers of a cake.

20.5 POSITRON EMISSION TOMOGRAPHY SCANNING (PET-SCAN)

Medical imaging methods that use ionizing radiation rely on how radiation interacts with matter; as photons pass through tissue they may be absorbed or scattered, and the pattern of energy deposition and attenuation is what ultimately produces measurable signals at a detector. In nuclear medicine, the radiation originates inside the patient from an administered radiopharmaceutical, and the detected signal reflects physiological function rather than only anatomy. A key functional technique is SPECT, in which a gamma-emitting tracer is introduced into the body and a rotating gamma camera records many projection views; these projections are then reconstructed to form cross-sectional images of tracer distribution, providing 3D information about perfusion and organ function. PET extends this idea using positron-emitting tracers; after emission, the positron travels a short distance in tissue and annihilates with an electron, producing two gamma photons emitted in nearly opposite directions; detecting these photons in coincidence allows the system to localize the event and reconstruct a quantitative map of tracer uptake. Consequently, PET is a non-invasive technique that images metabolic and biochemical activity at the cellular or molecular level, making it highly valuable for early cancer detection and staging, evaluating brain function, and identifying cardiac abnormalities before clear anatomical changes appear (Fig. 20.8).

Working of PET SCAN

A schematic diagram of PET scanning machine is shown in Fig. 20.8. In the PET-scan, a small amount of radioactive tracer is injected into a vein. The tracer is carried by the

blood and builds up more in malignant tissues that are very active, such as many cancers or working parts of the brain and heart. As the unstable nuclei of the tracer decay, they emit positrons. Each positron travels only a short distance before meeting an electron in nearby tissue. When they meet, the positron and electron annihilate and their mass is converted into energy in the form of two gamma-ray photons which move in opposite directions. Detectors arranged in a ring around the patient detect these photon pairs and send the signals to a computer, which constructs clear images showing where the tracer has collected. Brighter areas on the PET image correspond to regions of higher tracer concentration and therefore higher tissue activity.

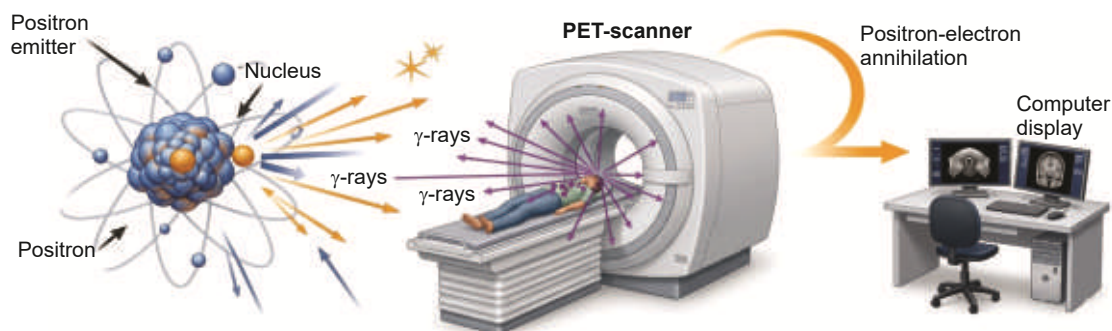


Fig. 20.8: Schematic diagram of PET-scanning

QUESTIONS

Multiple Choice Questions

Choose the correct answer.

- 20.1** What kind of radiation is detected in a PET scan after positron annihilation?
 (a) Beta particles (b) Gamma rays (c) X-rays (d) Alpha particles
- 20.2** The role of coincidence detection in PET-scanning is:
 (a) to control the temperature of the scanner
 (b) to detect gamma rays arriving simultaneously
 (c) to block unwanted radiation
 (d) to generate X-rays
- 20.3** Which principle allows ultrasound transducers to emit and detect sound waves?
 (a) Piezoelectric effect (b) Doppler effect
 (c) Photoelectric effect (d) Thermoelectric effect
- 20.4** The body part which commonly examined using ultrasound imaging is:
 (a) muscles (b) bones (c) abdomen (d) teeth
- 20.5** In X-ray imaging, what does "contrast" mainly depend on?
 (a) Density and thickness of tissues (b) Heartbeat rate
 (c) Distance from the detector (d) Surface temperature of the body
- 20.6** What advantage does a CT-scan have over a standard X-ray image?
 (a) Less radiation exposure (b) Simpler equipment
 (c) Produces 3D images from 2D slices (d) Uses sound waves instead of X-rays

20.7 The wave used in ultrasound imaging is:

- (a) longitudinal sound wave
- (b) transverse sound wave
- (c) electromagnetic wave
- (d) ionizing radiation wave

20.8 Characteristic X-rays are produced due to:

- (a) nuclear reactions
- (b) inner shell electronic transitions
- (c) outer shell transitions
- (d) ion collisions

20.9 What does “CT” stand for in medical imaging?

- (a) Combined Tomography
- (b) Computed Telemetry
- (c) Computed Tomography
- (d) Centralized Therapy

Short Answer Questions

- 20.1 Which type of radiation is produced during a PET-scan?
- 20.2 Which physical effect enables ultrasound transducers to produce sound waves?
- 20.3 Mention one area of the body commonly examined using ultrasound.
- 20.4 What causes contrast in X-ray images of different body tissues?
- 20.5 What occurs when ultrasound hits the boundary between two tissues?
- 20.6 How does CT-scan improve imaging? Explain briefly.
- 20.7 What do you know about the terms k_α and k_β X-rays.

Constructed Response Questions

- 20.1 What does a time difference in detecting gamma rays at two detectors indicate in a PET scan?
- 20.2 Why is ultrasound not effective for imaging air-filled organs like the lungs? Explain.
- 20.3 What happens when the X-ray contrast between two tissues is too low? Suggest a solution.
- 20.4 How does Computed Tomography scanning reduce the overlapping of organs seen in plain X-rays? Explain.
- 20.5 What could be a possible medical cause of uneven tracer concentration seen in a PET brain scan?
- 20.6 What is the purpose of lead collimators in PET gamma-ray detectors? Explain.

Comprehensive Questions

- 20.1 What is the piezoelectric effect, and how does it allow transducers to both emit and detect ultrasound waves?
- 20.2 Describe how an ultrasound image is created using the concepts of echoes and reflected sound waves.
- 20.3 Explain how attenuation of X-rays by different tissues creates contrast in the resulting image. Include an example.
- 20.4 Compare CT and conventional X-ray imaging with respect to image clarity and diagnostic value.
- 20.5 Write a brief note on PET. How is it used to scan malignant tissues?
- 20.6 How would you explain the inner shell transitions of X-rays? Also explain characteristic and continuous X-rays.

Space and Environment

Students Learning Outcomes

After studying this chapter, the students will be able to:

- ◆ explain the term luminosity [as the total power of radiation emitted by a star].
- ◆ apply the inverse-square law for radiant flux intensity F in terms of the luminosity L of the source $= \frac{L}{4\pi r^2}$.
- ◆ define and apply standard candles [Explain the use of standard candles to determine distances to galaxies].
- ◆ explain blackbody radiation and apply Wien's displacement law to solve problems [$\lambda_{\text{max}} T = \text{constant}$ to estimate the peak surface temperature of a star].
- ◆ apply the Stefan-Boltzmann's law [$L = 4\pi r^2 \times \sigma T^4$ to solve problems].
- ◆ estimate the radius of a star [applying Wien's displacement law and the Stefan-Boltzmann's law].
- ◆ explain that the lines in the emission and absorption spectra from distant objects show an increase in wavelength from their known values.
- ◆ explain why redshift leads to the idea that the universe is expanding [include using $\frac{\Delta \lambda}{\lambda} \approx \frac{\Delta f}{f} \approx \frac{v}{c}$ for the redshift of electromagnetic radiation from a source moving relative to an observer to solve problems relating to the expanding universe].
- ◆ state and explain Hubble's law and how it leads to the Big Bang Theory.
- ◆ describe Earth's climate system as a complex system having five interacting components [the atmosphere (air), the hydrosphere (water), the cryosphere (ice and permafrost), the lithosphere (Earth's upper rocky layer) and the biosphere (living things)].
- ◆ relate ocean currents and wind patterns to the climate system.
- ◆ explain how global climate is determined by energy transfer from the Sun [with specific reference to the factors and terms: state and use the term Earth energy budget. Explain how the energy imbalance between the poles and the equator can affect atmospheric circulation].
- ◆ explain that due to the conservation of angular momentum, the Earth's rotation diverts the air to the right in the northern hemisphere and to the left in the southern hemisphere, thus, forming distinct atmospheric cells.
- ◆ explain that ocean water that has more salt has a higher density and differences in density play an important role in ocean circulation.
- ◆ explain how the thermohaline circulation transports heat from the tropics to the polar regions.

Space and environment are two important and closely related fields that help us to understand both the universe and our planet. Space refers to a region beyond Earth's atmosphere in which all celestial bodies, such as planets, stars, and galaxies, exist and interact through gravity and electromagnetic radiations. Environment is the surrounding physical, chemical, and biological conditions in which living and non-living systems exist and interact. It includes air, water, land, climate, ecosystems, and all

natural processes that support life on the earth. In this chapter, we will explore how stars emit energy (luminosity), how we use tools like standard candles and redshift to measure cosmic distances, and how radiation laws help in estimating the temperature and size of stars. We will also learn how Earth's climate system functions through complex interactions among air, water, ice, land, and life, and how these elements respond to solar energy, wind, and ocean circulation.

21.1 LUMINOSITY OF A STAR

Luminosity is the total amount of energy emitted by a star per unit time in the form of electromagnetic radiations.

It is a fundamental property that defines the intrinsic brightness of a star. Luminosity is represented by the symbol L and is measured in watts, indicating the total power of radiation emitted by a star. Luminosity of a star primarily depends on its radius and surface temperature. Larger stars have more surface area to emit radiations and therefore, tend to be more luminous. Similarly, for stars of equal size, the one with the higher surface temperature will radiate more energy and hence, exhibits greater luminosity. The luminosity of a star can be represented as the product of its surface area ($A=4\pi r^2$) and radiant intensity ($I=\sigma T^4$), so

$$L = 4\pi r^2 \sigma T^4 \dots\dots\dots (21.1)$$

Equation (21.1) is called Stefan-Boltzmann's law. It states that:

The total output power of a star is directly proportional to the product of its surface area and the fourth power of its temperature.

where σ is Stefan-Boltzmann constant. Its value is $5.67 \times 10^{-8} \text{ W m}^{-2} \text{ K}^{-4}$. T is the surface temperature and r is the radius of the star.

Example 21.1 The luminosity of a star is $2.5 \times 10^{26} \text{ W}$, and its surface temperature is 6000 K. Estimate the radius of the star.

Solution Luminosity of the star $L = 2.5 \times 10^{26} \text{ W}$
 Surface temperature of the star $T = 6000 \text{ K}$
 Radius of the star $r = ?$
 Using Stefan-Boltzmann's law; $L = 4\pi r^2 \sigma T^4$

$$\text{or } r = \sqrt{\frac{L}{4\pi\sigma T^4}} = \sqrt{\frac{2.5 \times 10^{26} \text{ W}}{4 \times 3.14 \times 5.67 \times 10^{-8} \text{ W m}^{-2} \text{ K}^{-4} \times (6000 \text{ K})^4}} = 5.2 \times 10^8 \text{ m}$$

Tidbit!



The Sun shines with an immense luminosity of about 4×10^{26} watts that is equivalent to billions of nuclear explosions happening every second.

Do you know?

Stefan-Boltzmann's law helps astronomers estimate the size and brightness of stars just by knowing their temperature.

Brain Teaser

If a star becomes twice as hot, how many times more energy will it radiate?

21.2 RADIANT FLUX INTENSITY

Radiant flux intensity is the amount of electromagnetic energy per unit area per unit time emitted from a star that reaches the Earth's surface.

It is a small fraction of the total value of luminosity. In other words, radiant flux intensity is the luminosity per unit area measured on the surface of Earth. Radiant flux intensity from a star at a distance r is expressed as:

$$\text{Radiant flux intensity } F = \frac{L}{4\pi r^2} \dots\dots\dots (21.2)$$

Its SI unit is watt per square metre (W m^{-2}).

Where L is the luminosity and $4\pi r^2$ is the surface area of star. As the luminosity of a star remains constant, therefore, radiant flux intensity follows inverse-square law with distance as demonstrated in Fig. 21.1. According to this law, when the distance from centre of a star is doubled (from r to $2r$), intensity of the radiant flux reduces to one-fourth of its original value. Similarly, increasing the distance to three times its initial length ($3r$), decreases the flux by a factor of nine, and this pattern continues accordingly.

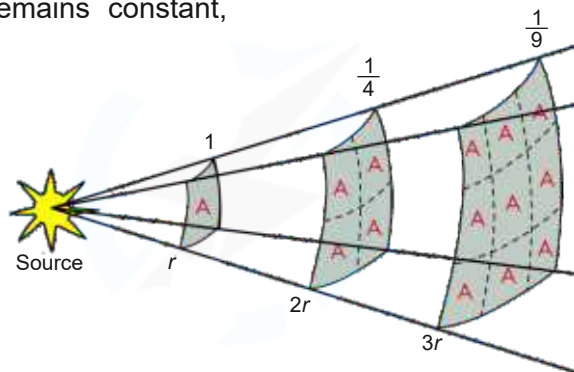


Fig. 21.1: Schematic illustration of the inverse-square law for radiant flux intensity

Example 21.2 The luminosity of a star is $4.50 \times 10^{26} \text{ W}$. It is located at a distance of $1.40 \times 10^{11} \text{ m}$ from a planet. Calculate the radiant flux intensity of the star at the surface of that planet.

Solution Luminosity of the star $L = 4.50 \times 10^{26} \text{ W}$

Distance between Earth and the star $r = 1.40 \times 10^{11} \text{ m}$

Radiant flux intensity $F = ?$

$$\begin{aligned} \text{Using the formula;} \quad F &= \frac{L}{4\pi r^2} \\ &= \frac{4.50 \times 10^{26} \text{ W}}{4 \times 3.14 (1.40 \times 10^{11} \text{ m})^2} \\ F &= 1.83 \times 10^3 \text{ W m}^{-2} \end{aligned}$$

Tidbit!

If Earth were twice as far from the Sun, the sunlight reaching it would be only one-fourth as intense. Such low energy might have made the evolution of life impossible.

Brain Teaser

Two stars appear equally bright in the night sky. However, one is 10 times farther away than the other. Which star is more luminous, and by how much?

Do you know?

The idea of standard candles dates back to the 1910s, when Henrietta Swan Leavitt discovered that Cepheid variable stars pulse in a predictable way. Her work laid the foundation for measuring cosmic distances and the scale of the universe. Without standard candles, we would not even know how big the universe is.

21.3 STANDARD CANDLES

In astronomy, measuring the distances to stars and galaxies is essential for understanding the structure and expansion of the universe.

The standard candles are astronomical objects with known intrinsic brightness that serve as vital tools for measuring cosmic distances.

By comparing how bright a standard candle appears from the Earth (apparent magnitude) to how bright it actually is, astronomers can calculate its distance using the inverse-square law of light. Examples of standard candles are Cepheid variable stars and Type Ia supernovae, both of which have predictable luminosities that make them reliable distance indicators.

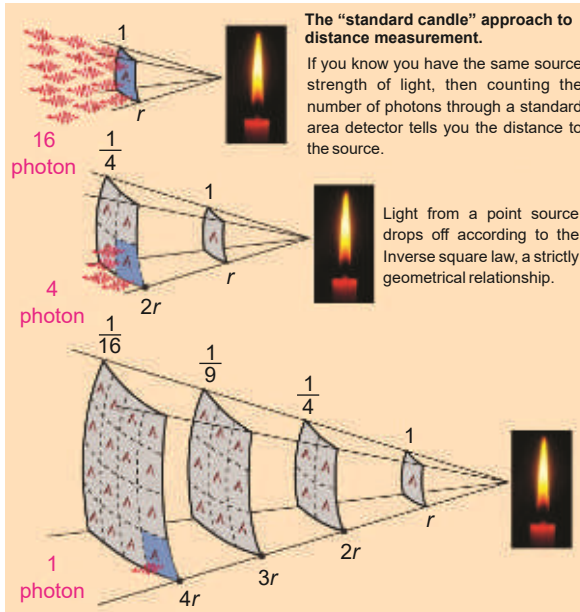


Fig. 21.2: Standard candles

21.4 BLACKBODY RADIATION

When an object is heated, it starts to glow and releases energy in the form of radiation. This radiation depends only on the temperature of the object, not on its material or shape. As the temperature increases, the colour of the object changes from dull red to bright yellow and then to white, indicating that the radiation is shifting from longer wavelengths (infrared) to shorter ones (ultraviolet). However, real objects do not absorb or emit all the radiations completely, which makes it difficult to study their behaviour accurately. To understand this better, scientists consider an ideal object, called a blackbody.

A blackbody is one that absorbs all the radiations falling on it and also emits radiations at all wavelengths perfectly.

In real life, a good example of a blackbody is a hollow cavity with a small hole and blackened inner walls as shown in Fig. 21.3. The inner walls of the cavity are coated with black carbon soot to help them absorb radiation effectively and reflect it internally.

Brain Teaser

Two stars appear equally bright in the sky. One is a standard candle whose true brightness is well known. The other is not. If the standard candle is farther away, what can you say about the true brightness of the second star?

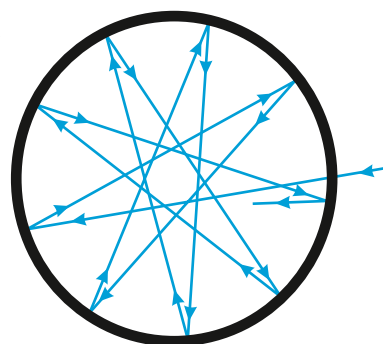
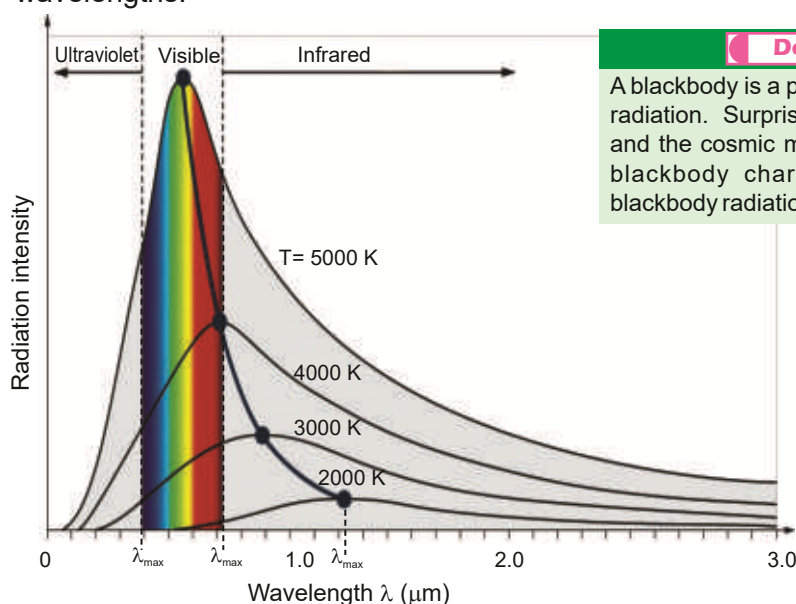


Fig. 21.3: Perfect or an ideal blackbody

When radiation enters through the small hole, it bounces around inside and gets trapped. If the cavity is heated to a high enough temperature, it starts to emit radiations of all wavelengths. This is known as blackbody radiation. The pattern of this radiation depends only on the temperature of the blackbody. As the temperature increases, the energy it gives off also increases. When we plot the intensity of this radiation against its wavelength (Fig. 21.4), we find that the peak of the curve shifts to shorter wavelengths with rising temperature. This means hotter blackbodies emit more energy at shorter wavelengths.



Do you know?

A blackbody is a perfect absorber and emitter of radiation. Surprisingly, both microwave oven and the cosmic microwave background exhibit blackbody characteristics. Even we emit blackbody radiation in the infrared range.

Brain Teaser

Two objects are heated to different temperatures. One glows red-hot, and the other glows white-hot. Which one is hotter, and why?

Fig. 21.4: Blackbody radiation spectra

Stars as Blackbody Radiators

Stars such as Sun behave like blackbody radiator. They generate and emit their own energy due to nuclear fusion and this energy spreads into space in the form of electromagnetic radiations. The radiation emitted by a star covers a wide range of wavelengths and forms a continuous spectrum without any gaps. This spectrum is similar to a blackbody spectrum. By studying the spectrum, scientists can estimate the surface temperature of the star accurately. As each star has its own spectrum, so its surface temperature is also different.

Wien's Displacement Law

Wien's displacement law describes the relationship between the surface temperature T of a blackbody and the wavelength λ_{max} at which it emits radiation most intensely. According to this law;

The wavelength corresponding to the radiation of maximum intensity emitted by a blackbody is inversely proportional to its surface temperature.

Tidbit!

The moment we switch on our mobile phone, our location can be tracked immediately by global positioning system.

i.e., $\lambda_{\max} \propto \frac{1}{T}$

or $\lambda_{\max} T = b \dots\dots\dots (21.3)$

where 'b' is a constant called Wien's constant whose value is 2.9×10^{-3} mK. This law demonstrates that, as the temperature of a blackbody increases, the maximum wavelength of its emitted radiation shifts towards shorter values. In simple words, hotter objects emit radiations with shorter wavelengths (appearing blue or violet), while cooler objects emit radiations with longer wavelengths (appearing red or infrared) as shown in Fig. 21.4. This shift explains the change in colour of hotter objects and helps to determine the temperature of stars and other radiating bodies.

Example 21.3 The surface temperature of a star is 6000 K, and its peak wavelength of emitted radiation is 480 nm. Another star has a peak wavelength of 320 nm. Calculate the surface temperature of this second star.

Solution Surface temperature a star-1 $T_1 = 6000$ K
 Wavelength at peak intensity of star-1 $\lambda_{\max 1} = 480$ nm
 Wavelength at peak intensity of star-2 $\lambda_{\max 2} = 320$ nm
 Surface temperature of star-2 $T_2 = ?$

Using the relationship for Wien's displacement law;

$$\lambda_{\max} T = b$$

i.e., $\lambda_{\max 1} T_1 = \lambda_{\max 2} T_2$ or $T_2 = \frac{\lambda_{\max 1} T_1}{\lambda_{\max 2}}$

$$T_2 = \frac{480 \times 10^{-9} \text{ m} \times 6000 \text{ K}}{320 \times 10^{-9} \text{ m}} = 9000 \text{ K}$$

So, the surface temperature of second star is 9000 K.

Brain Teaser

If a star's peak wavelength is shorter, what does that tell us about its temperature?

Tidbit!

Every star has its own light signature by studying the emission and absorption spectra, scientists can determine what elements a star is made of.

21.5 RADIUS OF A STAR

The radius of a star can be calculated using Wien's displacement law and the Stefan-Boltzmann law. For this purpose, the surface temperature T of a star is first determined using Wien's displacement law. The Luminosity of a star can be calculated using the radiant flux intensity and then Stefan-Boltzmann's law is used to determine the radius of star.

21.6 EMISSION AND ABSORPTION SPECTRA FROM DIFFERENT STARS

Stars emit a continuous spectrum of light because of their extremely hot, dense surfaces that behave like near-perfect blackbodies. As this light passes through the cooler outer atmosphere of a star, atoms in the gaseous layers absorb specific wavelengths corresponding to electronic transitions between energy levels. This creates dark lines in the continuous spectrum where specific colours of light have been absorbed. It is known

as the absorption spectrum. Although the absorbed energy is re-emitted at the same wavelengths, the emitted photons are scattered in various directions making them unlikely to reach observers on the Earth. Therefore, we detect a characteristic absorption spectrum with distinct missing wavelengths. In emission and absorption spectra from distant stars and galaxies, the lines in the spectra show an increase in wavelength from their known laboratory values (Fig. 21.5). This phenomenon is called redshift which occurs due to stretching of photons when they travel through the universe.

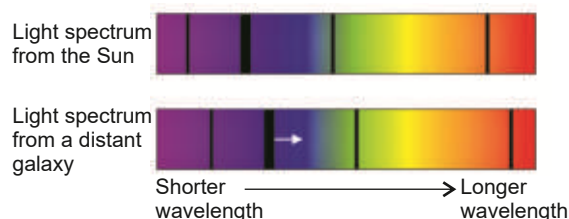


Fig. 21.5: Light spectra from Sun and distant galaxies, showing a redshift in wavelength

21.7 COSMIC REDSHIFT: EVIDENCE FOR AN EXPANDING UNIVERSE

Cosmic redshift refers to the effect in which light coming from distant galaxies across the universe is stretched toward the red end of the spectrum. This happens because the universe is expanding, causing galaxies to move away from us. As a result, the wavelength of light emitted by the galaxies increases. The greater the distance of a galaxy, the larger its redshift, i.e.; it is moving away faster. By studying cosmic redshift, scientists can measure how fast the universe is expanding.

The observed change in the wavelength of light is directly related to the relative velocity between the source and the observer.

$$\text{or} \quad Z = \frac{\Delta\lambda}{\lambda} \approx \frac{\Delta f}{f} \approx \frac{v}{c} \dots\dots\dots (21.4)$$

When the observed change in wavelength $\Delta\lambda$ is positive, it indicates a redshift Z . This illustrates that the celestial object is moving away from the Earth. On the other hand, a negative $\Delta\lambda$, represents a blue shift, showing that the object is approaching the Earth. Astronomers study the changes in the colour of light from stars and galaxies to understand how fast and in which direction they are moving. The light from far away galaxies appears more red than expected (redshift phenomenon). This happens because space itself is stretching as the universe expands, which also stretches the light waves and makes their wavelengths longer. The redshift pattern is observed in all directions which indicates the expansion of the universe. It is not just that galaxies are moving apart through space, rather the space between them is increasing. These findings strongly support the Big Bang Theory, which suggests that the universe originated from an extremely hot and dense state and is growing ever since. The relationship between redshift and distance enables scientists to determine the size of the universe and how it has changed over time.

21.8 HUBBLE'S LAW

Edwin Hubble established the foundation for modern cosmology by demonstrating that galaxies are moving away from the Earth, and their velocity increases with distance. In other words, the farther a galaxy is, the faster it appears to be moving away. This observation led to Hubble's law, which states that:

The recession velocity of a galaxy is directly proportional to its distance from the Earth.

Therefore $v \propto d$

or $v = H_0 d \dots\dots\dots (21.5)$

Here v is the recession velocity and d is the distance. H_0 is called Hubble's constant and its value is $2.3 \times 10^{-18} \text{ s}^{-1}$.

Figure 21.6 illustrates graphical representation of Hubble's law, revealing a direct and linear relationship between the redshift of galaxies and their distance from the Earth. As the distance increases, the recession velocity also increases, providing a strong evidence for an expanding universe. The slope of this linear trend corresponds to the Hubble constant H_0 , which quantifies the rate of cosmic expansion.

All the galaxies in the universe are moving away from each other. An observer in another galaxy will observe the same effect. This is happening because space itself is stretching. An inflated-balloon model is often considered by scientists to explain the Big Bang expansion, as shown in Fig. 21.7. The dots which are marked on the surface of a balloon represent the galaxies, and the skin of balloon represents four dimensional space-time. As the balloon is inflated, every dot "sees" the other dots moving away. This is similar to how galaxies move apart in space. This model helps us to understand two main ideas. First, no matter where we are in the universe, we would see galaxies moving away from us. This means the expansion is

Do you know?

Before Hubble's discovery, most scientists thought the universe was static. Hubble's law turned that idea on its head, suggesting the universe had a beginning, leading to the Big Bang Theory.

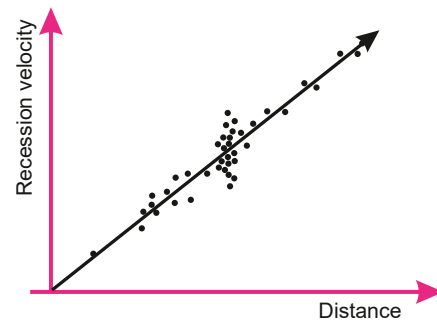


Fig. 21.6: Graphical representation of Hubble's law

Brain Teaser

A galaxy is moving away from us very fast. According to Hubble's law, what does that tell us about its distance?

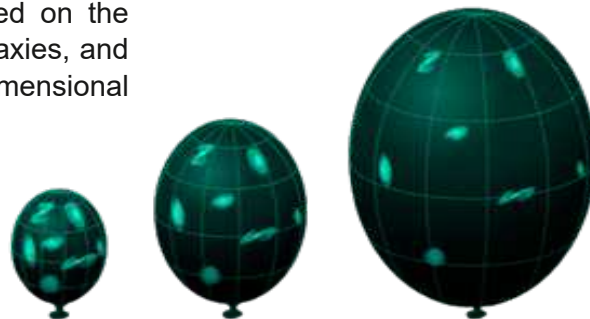


Fig. 21.7: Analogy of the expanding universe: dots move apart as the balloon inflates

happening everywhere, and there is no special centre. The Earth is not at the centre of the universe. The universe is expanding now, we can conceptually reverse time. Doing so implies:

- (i) in the past, galaxies were closer together.
- (ii) going further back, all matter was compressed into a very small, dense, and hot state.

This idea forms the basis of the Big Bang Theory, which proposes that the universe began from an initial singularity about 13.8 billion years ago and has been expanding ever since. Hubble's law provides strong observational evidence that the universe is not static but expanding, which is one of the key pillars supporting the Big Bang model.

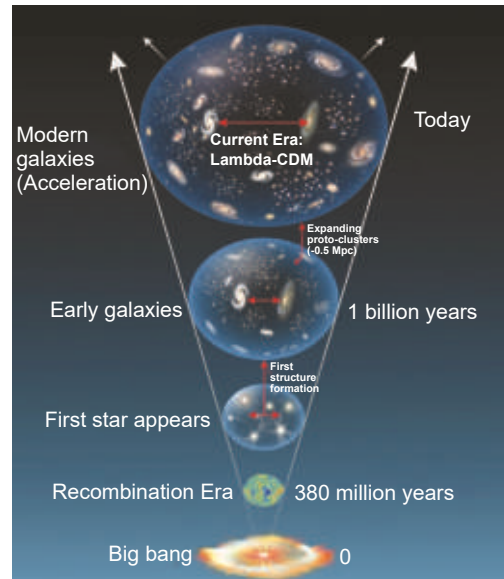


Fig. 21.8: Approximate timeline of the universe's evolution from the Big Bang to present

An approximate timeline for the evolution and expansion of the universe from the Big Bang till present is presented in Fig. 21.8.

21.9 EARTH'S CLIMATE SYSTEM

Earth's climate system is a long-term pattern of temperature, rainfall, wind, and other

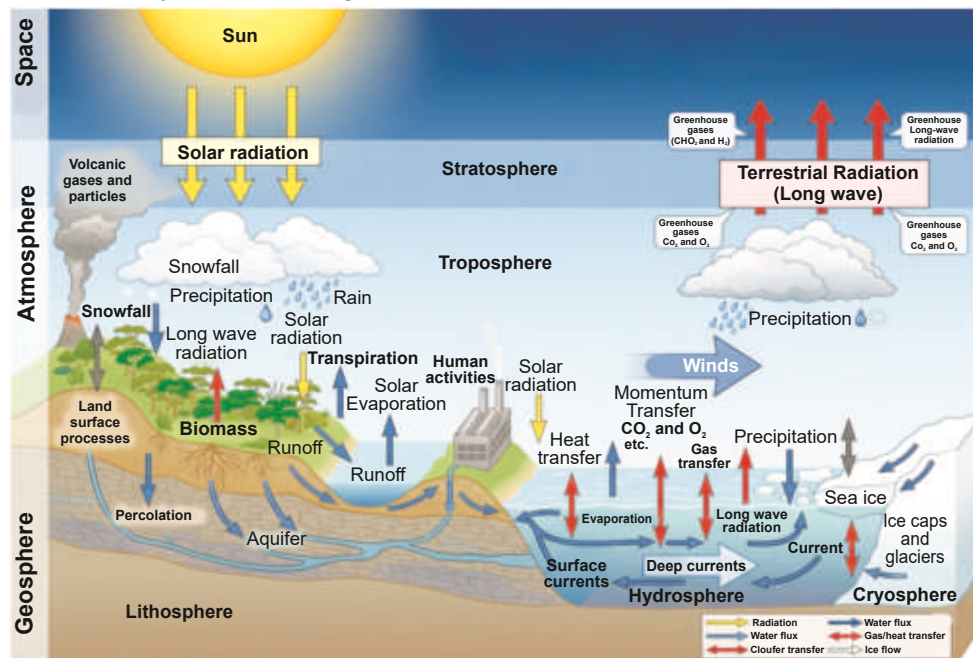


Fig. 21.9: Components of Earth's climate system and their interactions

weather conditions in a region. It affects all living things and natural systems on our planet. The climate is shaped by natural factors, like the Sun and ocean currents, and human activities, such as burning fossil fuels and cutting down forests. Understanding the climate of Earth helps us to prepare for and to reduce the impacts of climate change. The climate system of Earth is made up of five components which include atmosphere (air), the hydrosphere (water), the cryosphere (ice and permafrost), the lithosphere (Earth's upper rocky layer) and the biosphere (living things). These components work together and affect each other. Changes in one part can cause changes in the others.

21.10 OCEAN CURRENTS AND WIND PATTERNS

Ocean currents and wind patterns play an important role in regulating Earth's climate and weather systems. Ocean currents are generated due to large-scale movement of the seawater. The seawater circulates in two main ways: i.e., across the surface and between the deeper layers. Surface currents are mostly caused by wind, pushing water across the ocean's top layer. The vertical circulation of water is called thermohaline circulation. This is driven by differences in water temperature and salinity, which affect its density. The rotation of the Earth also shapes ocean movement, creating large circular current systems known as gyres. These spin clockwise in the northern hemisphere and anticlockwise in the southern hemisphere. The currents move warm water from the equator toward the poles and bring cold water back, helping in regulating the climate of the planet.

Gyres patterns on the surface of an ocean are shown in Fig. 21.10. Wind patterns are created due to an uneven heating of the Earth's surface. This causes warm air to rise and cooler air to move in, forming distinct wind belts like trade winds, westerlies, and polar easterlies. These winds not only move surface ocean water but also influence the direction and strength of ocean currents. Together, wind patterns and ocean currents help distribute heat and energy across the

Tidbit!

The trade winds that helped ancient sailors navigate are caused by Earth's rotation and solar heating. Ocean currents can change global weather patterns causing floods in one region and droughts in another.

Brain Teaser

A storm system is observed rotating clockwise in the southern hemisphere. If an identical system were observed in the northern hemisphere, which direction would it rotate, and why?

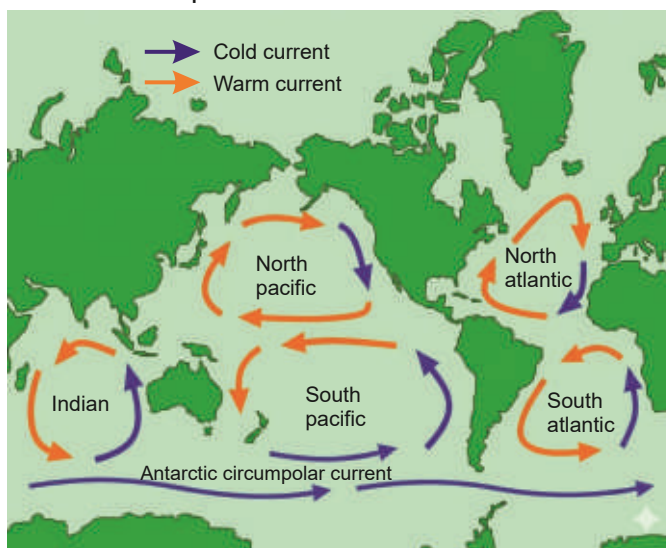


Fig. 21.10: Gyres patterns on the ocean surface

planet, affecting climate zones and weather patterns.

21.11 GLOBAL CLIMATE AND ENERGY TRANSFER FROM THE SUN

Energy received from the Sun plays an important role in producing changes in the global climate. Here, we will briefly describe how energy received from the Sun is distributed on Earth and how uneven heating of Earth between equator and poles can produce atmospheric changes.

1. Earth's Energy Budget

The Earth's energy budget refers to the balance between the solar energy absorbed by the Earth and the energy it radiates back into space.

It is a measure of how much solar energy is received at the Earth and how much of that energy is lost to space through reflection. When the Earth is colder, more ice and snow cover the surface. Ice and snow are bright and reflect a large amount of sunlight back into space, which makes the Earth even cooler. But when the Earth gets warmer, ice and snow start to melt. This exposes darker surfaces like land and oceans, which absorb more sunlight and reflect less. As a result, the Earth becomes warmer. This balance of incoming and outgoing energy is known as

Do you know?

The Earth only absorbs about 70% of the Sun's energy. The rest is reflected back into space by clouds, ice, and desert surfaces. A 1% imbalance in this budget can lead to major global temperature shifts.

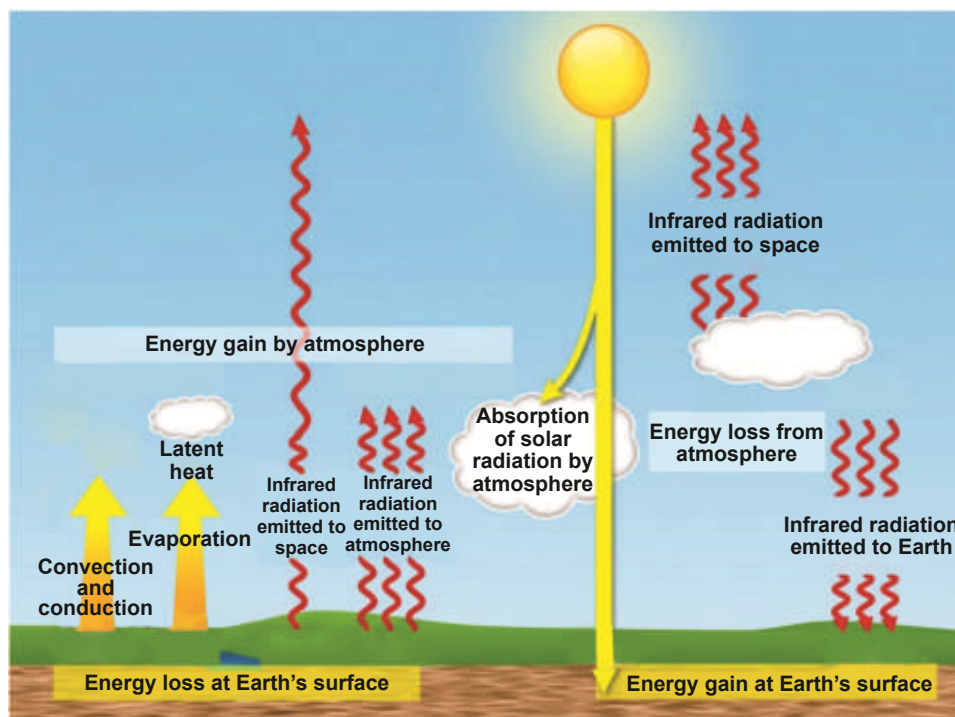


Fig. 21.11: Schematic diagram illustrating Earth's energy budget

the Earth's energy budget, which plays an important role in keeping Earth's temperature steady. Figure 21.11 depicts how sunlight reaches the Earth's surface and atmosphere. The energy coming in from the Sun is balanced by the energy going out from the Earth in the form of infrared (heat) radiations. When this balance is maintained, the temperature of Earth remains constant over time. However, greenhouse gases trap infrared radiations, preventing it from escaping into space. As a result, the Earth is warming up, a problem known as global warming.

2. Energy Imbalance between the Poles and Equator

The Earth's climate system is fundamentally shaped by the uneven distribution of solar energy across its surface. The equator receives direct, intense sunlight, causing warm air to rise and form low pressure zones, while the poles get slanted, weaker sunlight, leading to cold, sinking air and high pressure zones. This energy imbalance drives global atmospheric circulation, creating wind patterns like the trade winds near the equator, westerlies in mid-latitudes, and polar easterlies near the poles. These winds, along with ocean currents, work together to redistribute heat worldwide. However, human-induced greenhouse gas emissions are disrupting this natural balance, intensifying global warming and changing weather patterns. Figure 21.12 illustrates the Earth's energy imbalance, showing how solar radiation varies between the equator and poles. The equator receives direct sunlight, while the poles receive oblique sunlight at an angle, resulting in uneven heating.

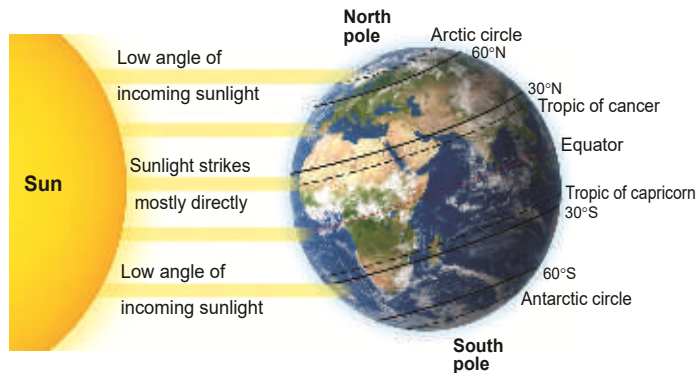


Fig. 21.12: Energy imbalance between poles and equator

21.12 ATMOSPHERIC CIRCULATION AND ATMOSPHERIC CELLS

The conservation of angular momentum of Earth plays an important role in shaping the moving air patterns. Atmospheric circulation is significantly affected by the conservation of angular momentum and Earth's rotation. The angular momentum of an air mass about Earth's rotation axis is given by $L = mvr$, where m is the mass of the air, r is the distance from the rotation axis, and v is the velocity. In atmospheric science, we usually consider specific angular momentum, which is the angular momentum per unit mass L' as follows:

$$L' = rv \dots\dots\dots (21.6)$$

As air moves from the equator toward the poles or vice versa, it retains its specific angular momentum

Tidbit!

The Dead Sea is so salty that people float effortlessly, the density is extremely high due to salt concentration.

at its original latitude. As the air mass moves, r decreases and so v increases to keep rv constant. As the Earth rotates faster at the equator than at higher latitudes, air moving toward the pole from the equator lags behind the rotation of the surface beneath it. This results in an apparent deflection of the moving air due to Earth's rotation and this type of effect is called Coriolis Effect. The deflection occurs toward right in the Coriolis Effect, northern hemisphere and to the left in the southern hemisphere. This deflection does not affect the speed of the air but changes its path, creating curved wind patterns. These deflected air flows lead to the formation of three distinct atmospheric circulation cells in each hemisphere: the Hadley cell, Ferrel cell, and Polar cell (Fig. 21.13).

Brain Teaser

In the polar regions, seawater becomes saltier as ice forms. How does this affect ocean circulation, and what process does it drive?

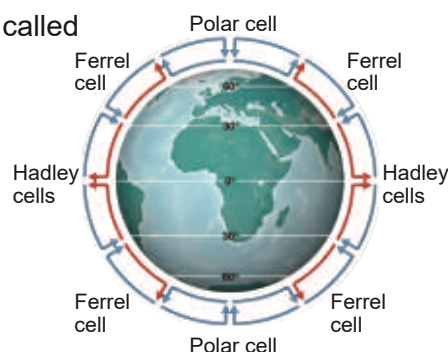


Fig. 21.13: Schematic representation of formation of atmospheric cells during atmospheric circulation of Earth.

21.13 ROLE OF SALT AND DENSITY IN OCEAN CIRCULATION

Ocean water contains a variety of dissolved substances, with salt being the most abundant. The amount of salt in ocean water is called salinity. When ocean water has more salt, it becomes heavier or denser. This happens because the dissolved salt increases the mass of the water without significantly increasing its volume. As a result, higher salinity leads to higher water density.

In the ocean, density differences between water masses are one of the key forces driving ocean circulation, especially in deep waters. These density differences arise mainly due to variations in temperature and salinity. Cold water is denser than warm water, and salty water is denser than less salty water. When surface waters become colder and saltier often due to evaporation or the formation of sea ice, they become denser and sink to deeper layers. As denser water sinks, it pushes other water out of the way, creating slow-moving but large-scale currents that circulate throughout the global oceans. This deep circulation helps transport heat, and gases (such as oxygen and carbon dioxide) around the globe. It connects surface and deep water systems and plays an important role in regulating the Earth's climate.

21.14 THERMOHALINE CIRCULATION

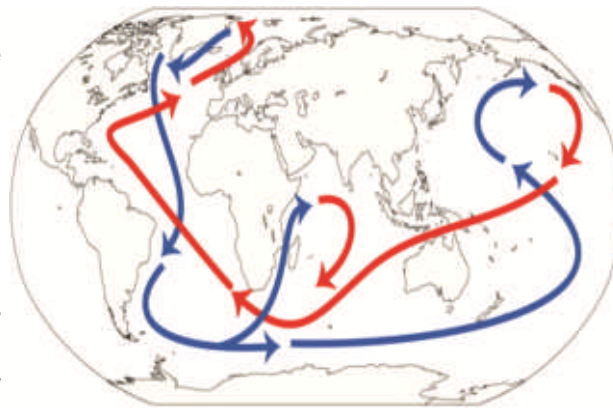
Thermohaline circulation, often referred to as the global conveyor belt, is a crucial process in the Earth's oceans that helps transport heat from the tropical regions to the poles. This circulation is driven by variations in water temperature and salinity, which together affect the

Tidbit!

Thermohaline circulation often called the "Global Conveyor Belt", it takes 1,000 years to make one full cycle around the globe. If global warming slows it down, Europe could face a mini-ice age, even while the rest of the world heats up.

density of the seawater. Warm, less dense water near the equator moves along the surface toward the polar regions, carrying heat energy. In polar areas, the water cools and becomes denser due to lower temperature and increased salinity from sea ice formation, causing it to sink and flow back toward the equator at deeper ocean levels. This continuous movement helps regulate the Earth's climate, influences atmospheric patterns, and supports marine ecosystems by redistributing heat and nutrients globally.

As shown in Fig. 21.14, warm surface currents (red ribbons) flow poleward, where cooling and ice formation increase their density, causing them to sink (blue ribbons). This cold, dense water spreads along the ocean floor before eventually rising again in other regions. The complete three-dimensional circulation of this global conveyor belt takes approximately 1,000 years to cycle water through all the world's ocean basins, playing a crucial role in regulating Earth's climate by redistributing heat and nutrients across the planet.



← Warm shallow current
← Cold deep current

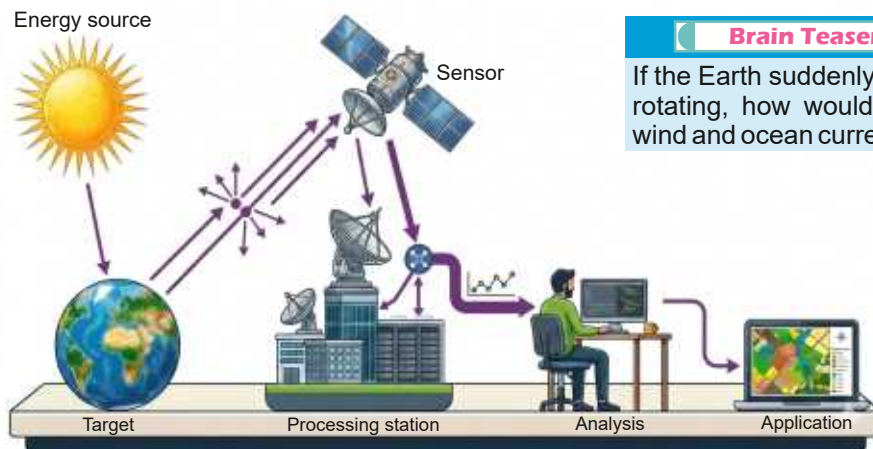
Fig. 21.14: Deep ocean circulation demonstrating thermohaline circulation

Brain Teaser

Why does the cold, salty water sink and create a deep current, and how does this relate to the real-world movement of ocean water in thermohaline circulation?

21.15 SATELLITE REMOTE SENSING

Satellite remote sensing is a powerful technology that uses sensors aboard orbiting spacecraft to monitor and analyze Earth's environment from space. It enables scientists



Brain Teaser

If the Earth suddenly stopped rotating, how would it affect wind and ocean currents?

Fig. 21.15: The remote sensing process

to track phenomena such as deforestation, climate change, ocean temperatures, and atmospheric composition. Unlike ground-based methods, it offers wide geographic coverage and repeated observations over time, making long-term environmental monitoring possible.

Following are the steps of satellite remote sensing:

1. **Source and Illumination:** The first and foremost element of SRS is source or illumination, which illuminates the target. It is in the form of electromagnetic radiation.
2. **Atmosphere:** After radiations are emitted from the source, they interact with atmosphere. This interaction may be in the form of scattering, reflection, etc.
3. **Interaction with Target:** After atmospheric interaction comes interaction with target of interest that this interaction depends on the properties of the target and the radiation.
4. **Recording with Sensor:** The emitted or reflected radiations from target are recorded by sensors which are remote (not in contact with target).
5. **Transmission and Reception:** The recorded radiations from sensor are then transmitted to the reception station to convert the data into understandable form of images, etc.
6. **Interpretation and Analysis:** The processed image is then interpreted (to get information) electronically or digitally.

QUESTIONS

Multiple Choice Questions

Choose the correct answer.

21.1 What does the term “luminosity” refer to in the context of stars?

- (a) The brightness of a star as seen from Earth
- (b) The total energy a star absorbs
- (c) The total power of radiation emitted by a star
- (d) The colour of a star's surface

21.2 Why are standard candles important in astronomy?

- (a) These reflect light from nearby stars
- (b) These are used to measure star temperatures
- (c) Their known luminosity helps to measure distances to galaxies
- (d) These are sources of blackbody radiation

21.3 According to Wien's displacement law, when the surface temperature of a star increases, its peak wavelength of emitted radiation:

- (a) increases
- (b) remains same
- (c) first increases then decreases
- (d) decreases

21.4 Which equation correctly represents the Stefan-Boltzmann law?

- (a) $L = 4\pi r^2 \sigma T^4$
- (b) $L = \sigma r T^2$
- (c) $L = \pi r^2 T$
- (d) $L = 4\pi r^2$

- 21.5 To estimate a star's radius, which law is most useful?
(a) Inverse-square law and Hubble's law (b) Stefan-Boltzmann's law
(c) Wien's law (d) Thermohaline circulation and Wien's law
- 21.6 What does a redshift in the spectral lines of a distant galaxy indicate?
(a) The galaxy is cooling down
(b) The galaxy is moving towards us
(c) The galaxy is rotating faster
(d) The wavelength of emitted light is increasing
- 21.7 Hubble's law states that the velocity of a galaxy is proportional to its:
(a) age (b) temperature (c) distance from the Earth (d) size
- 21.8 Which of the following is not one of the five interacting components of Earth's climate system?
(a) Atmosphere (b) Lithosphere (c) Magnetosphere (d) Cryosphere
- 21.9 Why does air deflect to the right in the northern hemisphere?
(a) Due to conservation of angular momentum (b) Due to Earth's axial tilt
(c) Due to magnetic field (d) Due to gravity
- 21.10 The thermohaline circulation is primarily driven by
(a) ocean tides and earthquakes
(b) differences in temperature and salinity of ocean water
(c) surface winds and volcanic activity
(d) Earth's gravitational pull

Short Answer Questions

- 21.1 Define luminosity in the context of a star.
- 21.2 State the inverse-square law for radiant flux intensity and explain its significance.
- 21.3 What is a standard candle, and how is it used to measure distances to galaxies?
- 21.4 State Stefan-Boltzmann's law and mention the physical quantities it relates.
- 21.5 How can you estimate the radius of a star using Wien's and Stefan-Boltzmann laws?
- 21.6 State Hubble's law, what is its connection to the Big Bang Theory?
- 21.7 List the five major components of Earth's climate system.

Constructed Response Questions

- 21.1 Why do astronomers use radiant flux intensity instead of total power to compare stars at different distances?
- 21.2 What is the role of standard candles in determining distances to far away galaxies? Explain why their known luminosity is critical for accurate measurements.
- 21.3 Why do the lines in the emission or absorption spectra of distant galaxies appear at longer wavelengths than expected, and what does this tell us about their motion?
- 21.4 How would the spectrum of light from galaxies look different if the universe were not expanding?
- 21.5 Describe how energy transfer from the Sun influences global climate patterns.

How does the imbalance of solar energy between the equator and the poles affect atmospheric circulation?

- 21.6 Explain how human activities, such as increasing greenhouse gases, can create an imbalance in Earth's energy budget.

Comprehensive Questions

- 21.1 How does the concept of luminosity help astronomers understand the nature of stars, and how is the inverse-square law used to calculate how bright a star appears from Earth? Support your answer with an example.
- 21.2 Explain how astronomers use standard candles to measure distances to galaxies. What makes a star or object suitable to be used as a standard candle? Provide a relevant example.
- 21.3 How does Wien's displacement law help determine the temperature of a star from its emitted light? Describe the relationship between the peak wavelength and temperature with suitable examples or calculations.
- 21.4 State the Stefan-Boltzmann law. How can this law be applied to calculate a star's luminosity and radius if its temperature and energy output are known? Provide an example.
- 21.5 What is redshift, and how does the shifting of spectral lines support the idea that the universe is expanding? Explain using the redshift formula and describe how this leads to Hubble's law and the Big Bang Theory.
- 21.6 Describe how Earth's climate system functions as a complex system involving the atmosphere, oceans, ice, land, and living things. How do ocean salinity and thermohaline circulation help transport heat from the equator toward the poles?

Numerical Problems

- 21.1 A star radiates with a total power of 3.6×10^{26} W. What is the radiant flux received by Earth located 1.5×10^{11} m away from the star? (Ans: 1.3×10^3 W m⁻²)
- 21.2 A supernova, acting as a standard candle, has a luminosity of 1.0×10^{36} W. If the observed flux at Earth is 2.0×10^{-12} W m⁻², calculate the distance between the Earth and the supernova. (Ans: 2×10^{23} m)
- 21.3 The peak wavelength of radiation from a star is observed to be 500 nm. Estimate the surface temperature of the star. (Ans: 5800 K)
- 21.4 Calculate the luminosity of a star having a radius of 7.0×10^8 m and a surface temperature of 6000 K. (Ans: 4.5×10^{26} W)
- 21.5 A star emits a total power of 3.9×10^{26} W and has a surface temperature of 5800 K. Estimate the radius of the star. (Ans: 7.0×10^8 m)
- 21.6 A spectral line that normally appears at 656 nm is observed from a distant galaxy at 700 nm. Calculate the redshift of the galaxy. (Ans: 0.067)
- 21.7 Light from a galaxy is redshifted, showing a shift value of 0.1. Determine the speed at which the galaxy is receding from the Earth. (Ans: 30,000 km s⁻¹)
- 21.8 A galaxy is receding from the Earth at a speed of 2.1×10^6 m s⁻¹. Using a Hubble constant of 70 km s⁻¹/Mpc, find the distance to the galaxy. (Ans: 9.26×10^{23} m)