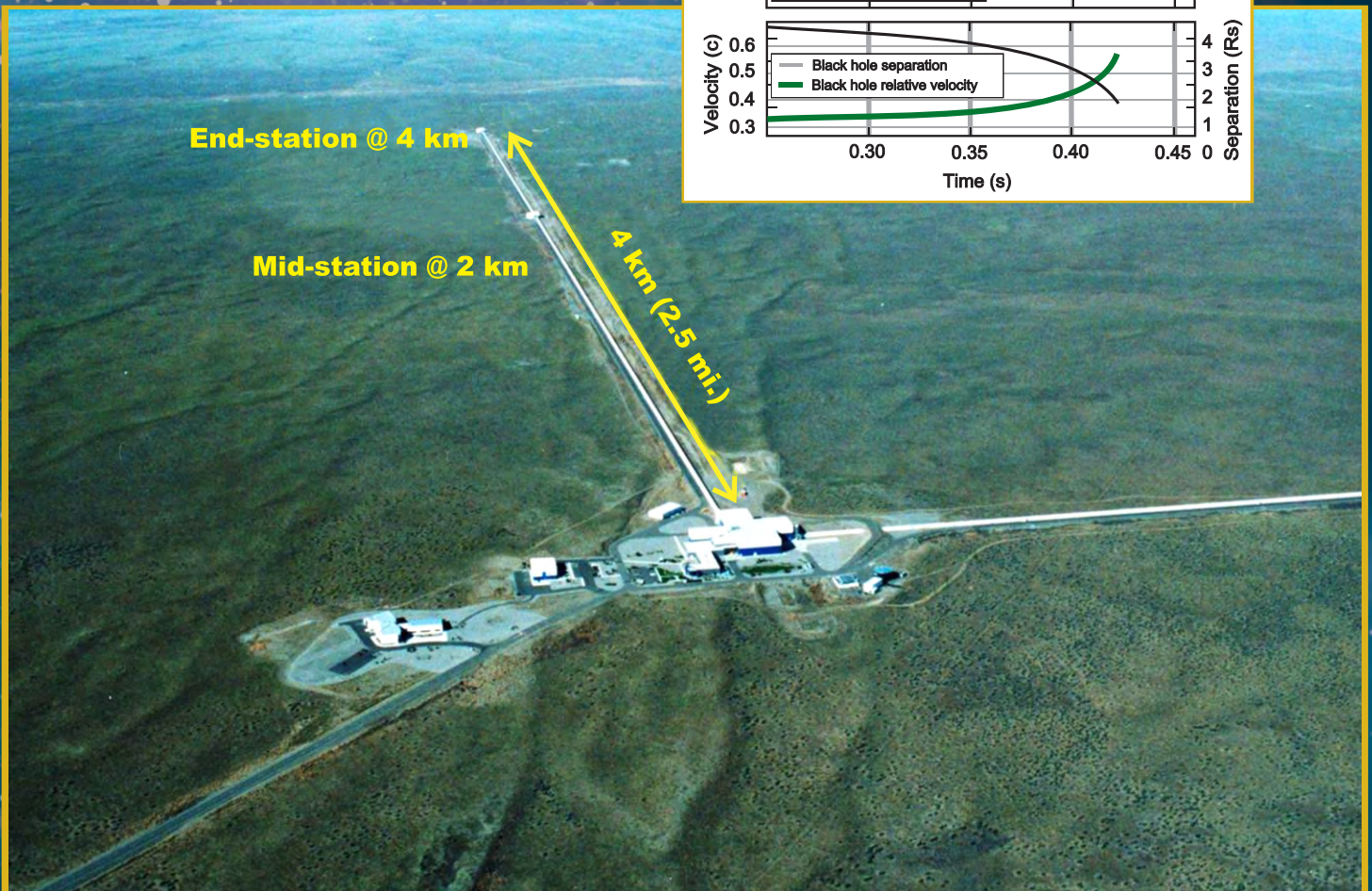
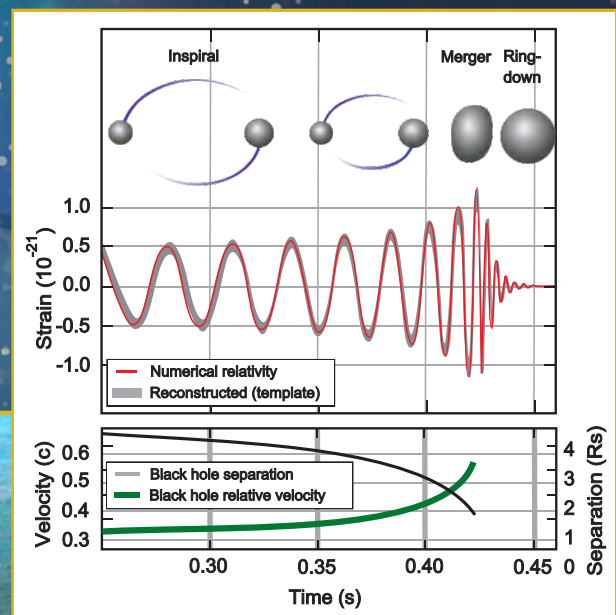




PHYSICS

Standard XI



The Constitution of India

Chapter IV A

Fundamental Duties

ARTICLE 51A

Fundamental Duties- It shall be the duty of every citizen of India—

- (a) to abide by the Constitution and respect its ideals and institutions, the National Flag and the National Anthem;
- (b) to cherish and follow the noble ideals which inspired our national struggle for freedom;
- (c) to uphold and protect the sovereignty, unity and integrity of India;
- (d) to defend the country and render national service when called upon to do so;
- (e) to promote harmony and the spirit of common brotherhood amongst all the people of India transcending religious, linguistic and regional or sectional diversities, to renounce practices derogatory to the dignity of women;
- (f) to value and preserve the rich heritage of our composite culture;
- (g) to protect and improve the natural environment including forests, lakes, rivers and wild life and to have compassion for living creatures;
- (h) to develop the scientific temper, humanism and the spirit of inquiry and reform;
- (i) to safeguard public property and to abjure violence;
- (j) to strive towards excellence in all spheres of individual and collective activity so that the nation constantly rises to higher levels of endeavour and achievement;
- (k) who is a parent or guardian to provide opportunities for education to his child or, as the case may be, ward between the age of six and fourteen years.

The Coordination Committee formed by GR No. Abhyas - 2116/(Pra.Kra.43/16) SD - 4 Dated 25.4.2016 has given approval to prescribe this textbook in its meeting held on 20.06.2019 and it has been decided to implement it from academic year 2019-20.

PHYSICS

Standard XI



T8P5A8

Download DIKSHA App on your smartphone. If you scan the Q.R.Code on this page of your textbook, you will be able to access full text. If you scan the Q.R.Code provided, you will be able to access audio-visual study material relevant to each lesson, provided as teaching and learning aids.



2019

**Maharashtra State Bureau of Textbook Production and
Curriculum Research, Pune.**

First Edition : © Maharashtra State Bureau of Textbook Production and Curriculum Research, Pune - 411 004.
2019

Reprint: 2021

The Maharashtra State Bureau of Textbook Production and Curriculum Research reserves all rights relating to the book. No part of this book should be reproduced without the written permission of the Director, Maharashtra State Bureau of Textbook Production and Curriculum Research, 'Balbharati', Senapati Bapat Marg, Pune 411004.

Subject Committee:

Dr. Chandrashekhar V. Murumkar, Chairman
Dr. Dilip Sadashiv Joag, Member
Dr. Pushpa Avinash Khare, Member
Dr. Rajendra Shankar Mahamuni, Member
Dr. Anjali Lalit Kshirsagar, Member
Dr. Rishi Baboo Sharma, Member
Shri. Rajiv Arun Patole, Member Secretary

Study group:

Dr. Umesh Anant Palnitkar
Dr. Vandana Laxmanrao Jadhav Patil
Dr. Neelam Sunil Shinde
Dr. Radhika Gautamkumar Deshmukh
Dr. Prabhakar Nagnath Kshirsagar
Dr. Bari Anil Ramdas
Dr. Sutar Milind Madhusudan
Dr. Bodade Archana Balasaheb
Dr. Chavan Jayashri Kalyanrao
Smt. Chokshi Falguni Manish
Shri. Ramesh Devidas Deshpande
Shri. Vinayak Shripad Katdare
Smt. Pratibha Pradeep Pandit
Shri. Dinesh Madhusudan Joshi
Shri. Kolase Prashant Panditrao
Shri. Brijesh Pandey
Shri. Ramchandra Sambhaji Shinde
Smt. Taksale Mugdha Milind

Illustration

Shri. Pradeep Ghodke
Shri. Shubham Chavan

Cover

Shri. Vivekanand S. Patil

Typesetting

DTP Section, Textbook Bureau,
Pune

Co-ordination :

Shri. Rajiv Arun Patole
Special Officer for Physics

Paper :

70 GSM Creamwove

Print Order :

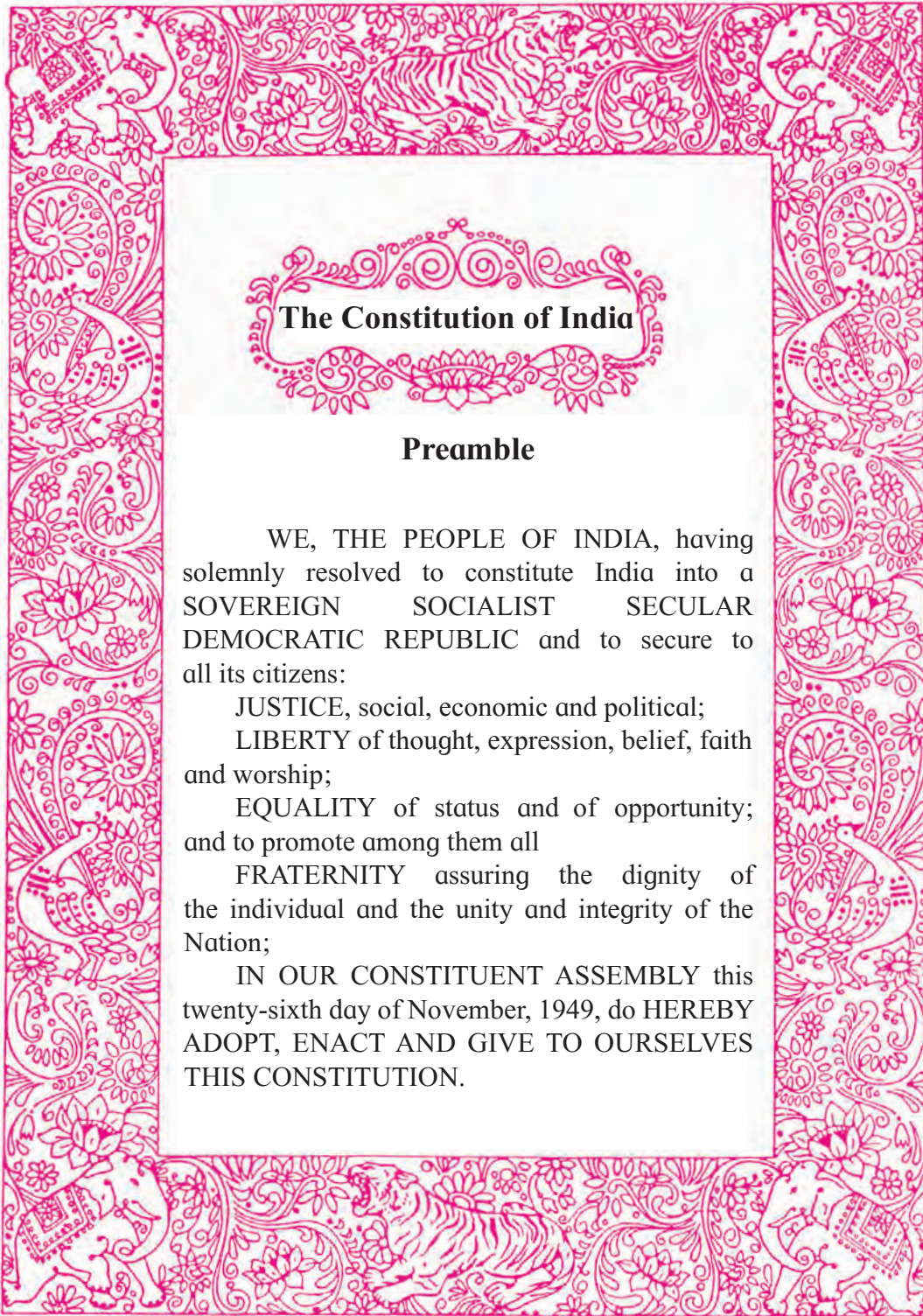
Printer

Production :

Shri Sachchitanand Aphale
Chief Production Officer
Shri Liladhar Atram
Production Officer

Publisher :

Shri Vivek Uttam Gosavi
Controller
Maharashtra State Textbook
Bureau, Prabhadevi,
Mumbai - 400 025



The Constitution of India

Preamble

WE, THE PEOPLE OF INDIA, having solemnly resolved to constitute India into a SOVEREIGN SOCIALIST SECULAR DEMOCRATIC REPUBLIC and to secure to all its citizens:

JUSTICE, social, economic and political;

LIBERTY of thought, expression, belief, faith and worship;

EQUALITY of status and of opportunity; and to promote among them all

FRATERNITY assuring the dignity of the individual and the unity and integrity of the Nation;

IN OUR CONSTITUENT ASSEMBLY this twenty-sixth day of November, 1949, do HEREBY ADOPT, ENACT AND GIVE TO OURSELVES THIS CONSTITUTION.

NATIONAL ANTHEM

Jana-gana-mana-adhināyaka jaya hē
Bhārata-bhāgya-vidhātā,

Panjāba-Sindhu-Gujarāta-Marāthā
Drāvida-Utkala-Banga

Vindhya-Himāchala-Yamunā-Gangā
uchchala-jaladhi-taranga

Tava subha nāmē jāgē, tava subha āsisa māgē,
gāhē tava jaya-gāthā,

Jana-gana-mangala-dāyaka jaya hē
Bhārata-bhāgya-vidhātā,

Jaya hē, Jaya hē, Jaya hē,
Jaya jaya jaya, jaya hē.

PLEDGE

India is my country. All Indians
are my brothers and sisters.

I love my country, and I am proud
of its rich and varied heritage. I shall
always strive to be worthy of it.

I shall give my parents, teachers
and all elders respect, and treat
everyone with courtesy.

To my country and my people,
I pledge my devotion. In their
well-being and prosperity alone lies
my happiness.

Preface

Dear Students,

It is a matter of pleasure and pride to place this exposition on basic physics in the hands of the young generation. This is not only textbook of physics for standard XI class, but embodies material which will be useful for self-study.

This textbook aims to create awareness about Physics. The National Curriculum Framework (NCF) was formulated in the year 2005, followed by the State Curriculum Framework (SCF) in 2010. Based on the given two frameworks, reconstruction of the curriculum and preparation of a revised syllabus has been undertaken which will be introduced from the academic year 2019-20. The textbook incorporating the revised syllabus has been prepared and designed by the Maharashtra State Bureau of Textbook Production and Curriculum Research, (Balharati), Pune.

The purpose of the book is to prepare a solid foundation for further studies in physics at the standard XII class. Proficiency in science in general and physics in particular is a basic requirement for the professional courses such as engineering and medicine etc., apart from the graduation courses in science itself. With this point of view, each chapter is prepared with elementary level and encompassing the secondary school level physics to the higher secondary level. Most of the topics are explained lucidly and in sufficient details, so that the students understand them well. A number of illustrative examples and figures are included to enlighten the student proficiency. With this background, the student is expected to solve the exercises given at the end of the chapters. For students who want more, Internet sites for many topics have been provided. They can enjoy further reading.

After all, physics is a conceptual subject. Knowledge about physical phenomena is gained as a natural consequence of observation, experience and revelation upon problem solving.

The book is written with this mind-set. The curriculum and syllabus conforms to the maxims of teaching such as moving from concrete to abstract, known to unknown and from part to the whole. For the first time, in this textbook of Physics, various activities have been introduced. These activities will not only help to develop understanding the content but also provide scope of the for gaining relevant and additional knowledge on your own efforts. A detailed information of all concepts is also given for a better understanding of the subject. QR Codes have been introduced for gaining additional information, abstracts of chapters and practice questions/ activities.

The efforts taken to prepare the textbook will not only enrich the learning experiences of the students, but also benefit other stakeholders such as teachers, parents as well as candidates aspiring for the competitive examinations.

We look forward to a positive response from the teachers and students.

Our best wishes to all!



(Dr. Sunil Magar)

Director

Maharashtra State Bureau of
Textbook Production and
Curriculum Research, Pune 4

Pune

Date : 20 June 2019

Bhartiya Saur : 30 Jyeshtha 1941

- For Teachers -

Dear Teachers,

We are happy to introduce the revised textbook of Physics for Std XI. This book is a sincere attempt to follow the maxims of teaching as well as develop a 'constructivist' approach to enhance the quality of learning. The demand for more activity based, experiential and innovative learning opportunities is the need of the hour. The present curriculum has been restructured so as to bridge the credibility gap that exists between what is taught and what students learn from direct experience in the outside world. Guidelines provided below will help to enrich the teaching-learning process and achieve the desired learning outcomes.

- ✓ To begin with, get familiar with the textbook yourself, and encourage the students to read each chapter carefully.
- ✓ The present book has been prepared for constructivist and activity-based teaching, including problem solving exercises.
- ✓ Use teaching aids as required for proper understanding of the subject.
- ✓ Do not finish the chapter in short.
- ✓ Follow the order of the chapters strictly as listed in the contents because the units are introduced in a graded manner to facilitate knowledge building.

- ✓ 'Error in measurements' is an important topic in physics. Please ask the students to use this in estimating errors in their measurements. This must become an integral part of laboratory practices.
 - ✓ Major concepts of physics have a scientific base. Encourage group work, learning through each other's help etc. Facilitate peer learning as much as possible by reorganizing the class structure frequently.
 - ✓ Do not use the boxes titled 'Do you know?' for evaluation. However, teachers must ensure that students read this extra information.
 - ✓ For evaluation, equal weightage should be assigned to all the topics. Use different combinations of questions. Stereotype questions should be avoided.
 - ✓ Use QR Code given in the textbook. Keep checking the QR Code for updated information. Certain important links, websites have been given for references. Also a list of reference books is given. Teachers as well as the students can use these references for extra reading and in-depth understanding of the subject.
- Best wishes for a wonderful teaching experience!

References:

1. Fundamentals of Physics - Halliday, Resnick, Walker; John Wiley (sixth ed.).
2. Sears and Zeemansky's University Physics - Young and Freedman, Pearson Education (12th ed.)
3. Physics for Scientists and Engineers - Lawrence S. Lerner; Jones and Bartlett Publishers, UK.

Front Page : Figure shows the LIGO laboratory in the United States of America and the inset shows the trace of gravity waves detected upon the merger of two black holes. In the background is the artist's impression of planets and galaxies.

Since ages, mankind is awed by the sheer scale of the universe and is trying to understand the laws governing the same. Today we observe the events in the universe with highly sophisticated instruments and laboratories such as the LIGO project seen on the cover. Picture Credit: Caltech/ MIT/ LIGO laboratory.

Figure Credit : B. P. Abott et al. Physical Review Lett 116, 061102, 2016

DISCLAIMER Note : All attempts have been made to contact copy right/s (©) but we have not heard from them. We will be pleased to acknowledge the copy right holder (s) in our next edition if we learn from them.

Competency Statements Standard XI

Area/ Unit/ Lesson	Competency Statements After studying the content in Textbook student ...
Units and Mathematical Tools	- Distinguish between fundamental and derived quantities. - Distinguish between different system of units and their use. - Identify methods to be used for measuring lengths and distances of varying magnitudes. - Check correctness of physical equations using dimensional analysis. - Establish the relation between related physical quantities using dimensional analysis. - Find conversion factors between the units of the same physical quantity in two different sets of units. - Identify different types of errors in measurement of physical quantities and estimate them. - Identify the order of magnitude of a given quantity and the significant figures in them. - Distinguish between scalar and vector quantities. - Perform addition, subtraction and multiplication (scalar and vector product) of vectors. - Determine the relative velocity between two objects. - Obtain derivatives and integrals of simple functions. - Obtain components of vectors. - Apply mathematical tools to analyze physics problems.
Motion and Gravitation	- Visualize motions in daily life in one, two and three dimensions. - Explain the necessity of Newton's first law of motion. - Categorize various forces of nature into four fundamental forces. - State various conservation principles and use these in daily life situations. - Derive expressions and evaluate work done by a constant force and variable force. - Organize/categorize the common principles between collisions and explosions. - Explain the necessity of defining impulse and apply it to collisions, etc. - Elaborate the limitations of Newton's laws of motion. - Elaborate different types of mechanical equilibria with suitable examples. - Apply the Kepler's laws of planetary motion to solar system. - Elaborate Newton's law of gravitation. - Calculate the values of acceleration due to gravity at any height above and depth below the earth's surface. - Distinguish between different orbits of earth's satellite. - Explain how escape velocity varies from planet. - Explain weightlessness in a satellite.
Properties of Matter	- Explain the difference between elasticity and plasticity - Identify elastic limit for a given material. - Differentiate between different types of elasticity modules. - Judge the suitability of materials for specific applications in daily life appliances. - Identify the role of force of friction in daily life. - Differentiate between good and bad conductors of heat. - Relate underlying physics for use of specific materials for use in thermometers for specific applications.
Sound and Optics	- Apply and relate various parameters related to wave motion. - Compare various types of waves with common features and distinguishing features. - Analytically relate the factors on which the speed of sound and speed of light depends. - Explain the essential factor to describe wave propagation and relate it with phase angle. - Apply the laws of reflection to light. - Mathematically describe the Doppler effect for sound waves. - Apply the laws of refraction to common phenomena in daily life like, a mirage or a rainbow. - Identify the defects in images obtained by mirrors and lenses, with their cause and ways of reducing or eliminating them. - Explain the construction and use of various optical instruments such as a microscope, a telescope, etc. - Relate dispersion of light with colour and apply it analytically with the help of prisms.

	- Describe dispersive power as a basic property of transparent materials and relate it with their refractive indices. - Analyze the time taken to receive an echo and calculate distance to the reflecting object. - Explain reverberation and acoustics.
Electricity and Magnetism	- Distinguish between conductors and insulators. - Apply coulomb's law and obtain the electric field due to a certain distribution of charges. - Define dipole, obtain the dipolar field. - Relate the drift of electrons in a conductor to resistivity - Calculate resistivity at various temperature. - Connect resistors in series and parallel combination. - Compare electric and magnetic fields. - Draw electric and magnetic lines of force. - Obtain magnetic parameters of the Earth. - Solve numerical and analytical problems.
Communication and Semiconductors	- Explain the properties of an electromagnetic wave. - Distinguish between mechanical waves and electromagnetic waves. - Identify different types of electromagnetic radiations from γ- rays to radio waves. - Distinguish between different modes of propagation of EM waves through earth's atmosphere. - Identify different elements of a communication system. - Explain different types of modulation and identify the types of modulation needed in given situation. - Distinguish between conductors, insulators and semiconductors based on band structure. - Differentiate between p type and n type semiconductors and their uses. - Explain working of forward and reverse biased junction. - Explain the working of semiconductor diode.

CONTENTS

Sr. No	Title	Page No
1	Units and Measurements	1-15
2	Mathematical Methods	16-29
3	Motion in a Plane	30-46
4	Laws of Motion	47-77
5	Gravitation	78-99
6	Mechanical Properties of Solids	100-113
7	Thermal Properties of Matter	114-141
8	Sound	142-158
9	Optics	159-187
10	Electrostatics	188-206
11	Electric Current Through Conductors	207-220
12	Magnetism	221-228
13	Electromagnetic Waves and Communication System	229-241
14	Semiconductors	242-256



Can you recall?

1. What is a unit?
2. Which units have you used in the laboratory for measuring
(i) length (ii) mass (iii) time (iv) temperature?
3. Which system of units have you used?

1.1 Introduction:

Physics is a quantitative science, where we measure various physical quantities during experiments. In our day to day life, we need to measure a number of quantities, e.g., size of objects, volume of liquids, amount of matter, weight of vegetables or fruits, body temperature, length of cloth, etc. A measurement always involves a comparison with a standard measuring unit which is internationally accepted. For example, for measuring the mass of a given fruit we need standard mass units of 1 kg, 500 g, etc. These standards are called units. The measured quantity is expressed in terms of a number followed by a corresponding unit, e.g., the length of a wire is written as 5 m where m (metre) is the unit and 5 is the value of the length in that unit. Different quantities are measured in different units, e.g. length in metre (m), time in seconds (s), mass in kilogram (kg), etc. The standard measure of any quantity is called the unit of that quantity.

1.2 System of Units:

In our earlier standards we have come across various systems of units namely

- (i) CGS: Centimetre Gram Second system
- (ii) MKS: Metre Kilogram Second system
- (iii) FPS: Foot Pound Second system.
- (iv) SI: System International

The first three systems namely CGS, MKS and FPS were used extensively till recently. In 1971, the 14th International general conference on weights and measures recommended the use of 'International system' of units. This international system of units is called the SI units. As the SI units use decimal system, conversion within the system is very simple and convenient.

1.2.1 Fundamental Quantities and Units:

The physical quantities which do not depend on any other physical quantities for their measurements are known as fundamental quantities. There are seven fundamental quantities: length, mass, time, temperature, electric current, luminous intensity and amount of substance.

Fundamental units: The units used to measure fundamental quantities are called fundamental units. The fundamental quantities, their units and symbols are shown in the Table 1.1.

Table 1.1: Fundamental Quantities with their SI Units and Symbols

Fundamental quantity	SI units	Symbol
1) Length	metre	m
2) Mass	kilogram	kg
3) Time	second	s
4) Temperature	kelvin	K
5) Electric current	ampere	A
6) Luminous Intensity	candela	cd
7) Amount of substance	mole	mol

1.2.2 Derived Quantities and Units:

In physics, we come across a large number of quantities like speed, momentum, resistance, conductivity, etc. which depend on some or all of the seven fundamental quantities and can be expressed in terms of these quantities. These are called derived quantities and their units, which can be expressed in terms of the fundamental units, are called derived units.

For example,

SI unit of velocity

$$= \frac{\text{Unit of displacement}}{\text{Unit of time}} = \frac{\text{m}}{\text{s}} = \text{m s}^{-1}$$

Unit of momentum = (Unit of mass) × (Unit of velocity)

$$= \text{kg m/s} = \text{kg m s}^{-1}$$

The above two units are derived units.

Supplementary units : Besides, the seven fundamental or basic units, there are two more units called supplementary units: (i) Plane angle $d\theta$ and (ii) Solid angle $d\Omega$

(i) **Plane angle ($d\theta$) :** This is the ratio of the length of an arc of a circle to the radius of the circle as shown in Fig. 1.1 (a). Thus $d\theta = ds/r$ is the angle subtended by the arc at the centre of the circle. It is measured in radian (rad). An angle θ in radian is denoted as θ^c .

(ii) **Solid angle ($d\Omega$) :** This is the 3-dimensional analogue of $d\theta$ and is defined as the area of a portion of surface of a sphere to the square of radius of the sphere. Thus $d\Omega = dA/r^2$ is the solid angle subtended by area ds at O as shown in Fig. 1.1 (b). It is measured in steradians (sr). A sphere of radius r has surface area $4\pi r^2$. Thus, the solid angle subtended by the entire sphere at its centre is $\Omega = 4\pi r^2/r^2 = 4\pi \text{ sr}$.

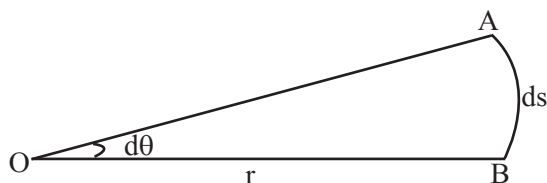


Fig 1.1 (a): Plane angle $d\theta$.

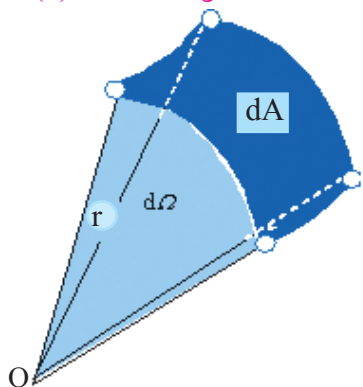


Fig 1.1 (b): Solid angle $d\Omega$.

Example 1.1: What is the solid angle subtended by the moon at any point of the Earth, given the diameter of the moon is 3474 km and its distance from the Earth $3.84 \times 10^8 \text{ m}$.

Solution: Solid angle subtended by the moon at the Earth

$$\begin{aligned} &= \frac{\text{Area of the disc of the moon}}{(\text{moon - earth distance})^2} \\ &= \frac{\pi \times (1.737 \times 10^3)^2}{(3.84 \times 10^5)^2} \\ &= 6.425 \times 10^{-5} \text{ sr} \end{aligned}$$



Do you know ?

The relation between radian and degree is

$$\pi \text{ radians} = \pi^c = 180^\circ$$

$$\therefore 1 \text{ radian} = \frac{180}{\pi} = \frac{180}{3.1415} = 57.297^\circ$$

$$\text{Similarly } 1^\circ = \frac{\pi}{180} = \frac{3.1415}{180} = 1.745 \times 10^{-2} \text{ rad}$$

$$1^\circ = 60', 1' = 2.91 \times 10^{-4} \text{ rad}$$

$$\text{and } 1' = 60'', 1'' = 4.847 \times 10^{-6} \text{ rad}$$

1.2.3 Conventions for the use of SI Units:

- (1) Unit of every physical quantity should be represented by its symbol.
- (2) Full name of a unit always starts with smaller letter even if the name is after a person, e.g., 1 newton, 1 joule, etc. But symbol for unit named after a person should be in capital letter, e.g., N after scientist Newton, J after scientist Joule, etc.
- (3) Symbols for units do not take plural form for example, force of 20 N and not 20 newtons or not 20 Ns.
- (4) Symbols for units do not contain any full stops at the end of recommended letter, e.g., 25 kg and not 25 kg..
- (5) The units of physical quantities in numerator and denominator should be written as one ratio for example the SI unit of acceleration is m/s^2 or m s^{-2} but not m/s/s .
- (6) Use of combination of units and symbols for units is avoided when physical quantity is expressed by combination of two. e.g., The unit J/kg K is correct while joule/kg K is not correct.
- (7) A prefix symbol is used before the symbol of the unit.

Thus prefix symbol and units symbol constitute a new symbol for the unit which can be raised to a positive or negative power of 10.

1ms = 1 millisecond = 10^{-3} s

1 μ s = 1 microsecond = 10^{-6} s

1ns = 1 nanosecond = 10^{-9} s

Use of double prefixes is avoided when single prefix is available

10^{-6} s = 1 μ s and not 1mms.

10^{-9} s = 1ns and not 1m μ s

- (8) Space or hyphen must be introduced while indicating multiplication of two units e.g., m/s should be written as m s⁻¹ or m·s⁻¹ and Not as ms⁻¹ (because ms will be read as millisecond).

1.3 Measurement of Length:

One fundamental quantity which we have

discussed earlier is length. To measure the length or distance the SI unit used is metre (m). In 1960, a standard for the metre based on the wavelength of orange-red light emitted by atoms of krypton was adopted. By 1983 a more precise measurement was developed. It says that a metre is the length of the path travelled by light in vacuum during a time interval of 1/299792458 second. This was possible as by that time the speed of light in vacuum could be measured precisely as $c = 299792458$ m/s

Some typical distances/lengths are given in Table 1.2.

Table 1.2: Some Useful Distances

Measurement	Length in metre
Distance to Andromeda Galaxy (from Earth)	2×10^{22} m
Distance to nearest star (after Sun) Proxima Centuari (from Earth)	4×10^{16} m
Distance to Pluto (from Earth)	6×10^{12} m
Average Radius of Earth	6×10^6 m
Height of Mount Everest	9×10^3 m
Thickness of this paper	1×10^{-4} m
Length of a typical virus	1×10^{-8} m
Radius of hydrogen atom	5×10^{-11} m
Radius of proton	1×10^{-15} m

1.3.1 Measurements of Large Distance:

Parallax method

Large distance, such as the distance of a planet or a star from the Earth, cannot be measured directly with a metre scale, so a parallax method is used for it.

Let us do a simple experiment to understand what is parallax.

Hold your hand in front of you and look at it with your left eye closed and then with your right eye closed. You will find that your hand appears to move against the background. This effect is called parallax. Parallax is defined as the apparent change in position of an object due to a change in the position of the observer. By measuring the parallax angle (θ) and knowing the distance between the eyes E_1E_2 as shown in Fig. 1.2, we can determine the distance of the object from us, i.e., OP as E_1E_2/θ .

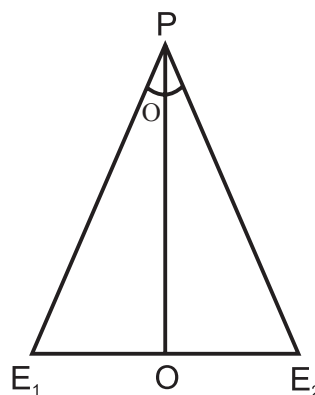


Fig.1.2: Parallax method for determining distance.

As the distances of planets from the Earth are very large, we can not use two eyes method as discussed here. In order to make simultaneous observations of an astronomical object, we select two distant points on the Earth.

Consider two positions A and B on the surface of Earth, separated by a straight line at

distance b as shown in Fig. 1.3. Two observers at these two points observe a distant planet S simultaneously. We measure the angle $\angle ASB$ between the two directions along which the planet is viewed at these two points. This angle, represented by symbol θ , is the parallax angle.

As the planet is far away, i.e., the distance of the planet from the Earth is very large in comparison to b , $b/D \ll 1$ and, therefore, θ is very small.

We can thus consider AB as the arc of length b of the circle and D its radius.

$AB = b$ and $AS = BS = D$ and $\theta \cong AB/D$, where θ is in radian

$$D = b / \theta$$

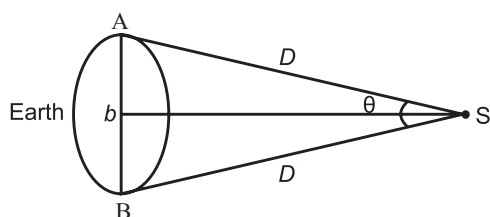


Fig.1.3: Measurement of distances of planets

1.3.2 Measurement of Distance to Stars:

Sun is the star which is closest to the Earth. The next closest star is at a distance of 4.29 light years. The parallax measured from two most distance points on the Earth for stars will be too small to be measured and for this purpose we measure the parallax between two farthest points (i.e. 2 AU apart, see box below) along the orbit of the Earth around the Sun (see figure in example 1.2 below).

1.3.3 Measurement of the Size of a Planet or a Star:

If d is the diameter of a planet, the angle subtended by it at any single point on the Earth is called angular diameter of the planet. Let α be the angle between the two directions when two diametrically opposite points of the planet are viewed through a telescope as shown in Fig. 1.4. As the distance D of the planet is large (assuming it has been already measured), we can calculate the diameter of the planet as

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

$$\therefore d = \alpha D \quad \text{--- (1.2)}$$

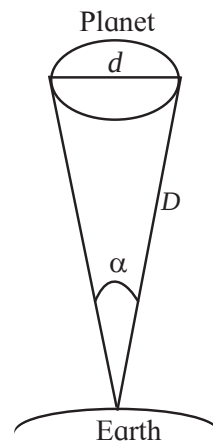


Fig. 1.4: Measurement of size of a planet

1.3.4 Measurement of Very Small Distances:

When we intend to measure the size of the atoms and molecules, the conventional apparatus like Vernier calliper or screw gauge will not be useful. Therefore, we use electron microscope or tunnelling electron microscope to measure the size of atoms.



Do you know ?

For measuring large distances, astronomers use the following units.

$$1 \text{ astronomical unit, (AU)} = 1.496 \times 10^{11} \text{ m}$$

$$1 \text{ light year} = 9.46 \times 10^{15} \text{ m}$$

$$1 \text{ parsec (pc)} = 3.08 \times 10^{16} \text{ m} \cong 3.26 \text{ light years}$$

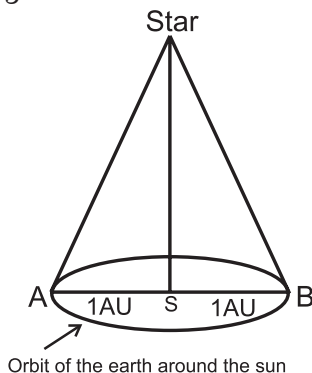
A light year is the distance travelled by light in one year. The astronomical unit (AU) is the mean distance between the centre of the Earth and the centre of the Sun.

A parsec (pc) is the distance from where 1AU subtends an angle of 1 second of arc.

$$r = \frac{1 \text{ AU}}{(1'')^c} = \frac{1.496 \times 10^{11}}{4.847 \times 10^{-6}} = 3.08 \times 10^{16} \text{ m}$$

Example 1.2: A star is 5.5 light years away from the Earth. How much parallax in arcsec will it subtend when viewed from two opposite

points along the orbit of the Earth?



Solution: Two opposite points A and B along the orbit of the Earth are 2 AU apart. The angle subtended by AB at the position of the star is = AB/distance of the star from the Earth

$$= \frac{2\text{AU}}{5.5 \text{ ly}} = \frac{2 \times 1.496 \times 10^{11} \text{ m}}{5.5 \times 9.46 \times 10^{15} \text{ m}} = 5.75 \times 10^{-6} \text{ rad}$$

$$= 5.75 \times 10^{-6} \times 57.297 \times 60 \times 60 \text{ arcsec}$$

$$= 1.186 \text{ arcsec}$$



Do you know ?

Small distances are measured in units of (i) fermi = $1\text{F} = 10^{-15} \text{ m}$ in SI system. Thus, 1F is one femtometre (fm) (ii) Angstrom = $1 \text{ \AA} = 10^{-10} \text{ m}$

For measuring sizes using a microscope we need to select the wavelength of light to be used in the microscope such that it is smaller than the size of the object to be measured. Thus visible light (wavelength from $4000 \text{ \AA}$ to $7000 \text{ \AA}$) can measure sizes upto about $4000 \text{ \AA}$. If we want to measure even smaller sizes we need to use even smaller wavelength and so the use of electron microscope is necessary. As you will study in the XIIth standard, each material particle corresponds to a wave. The approximate wavelength of the electrons in an electron microscope is about $0.6 \text{ \AA}$ so that one can measure atomic sizes $\approx 1 \text{ \AA}$ using this microscope.

Example 1.3: The moon is at a distance of $3.84 \times 10^8 \text{ m}$ from the Earth. If viewed from two diametrically opposite points on the Earth, the angle subtended at the moon is $1^\circ 54'$. What is the diameter of the Earth?

Solution: Angle subtended

$$\theta = 1^\circ 54' = 114' = 114 \times 2.91 \times 10^{-4} \text{ rad}$$

$$= 3.317 \times 10^{-2} \text{ rad}$$

Diameter of the Earth = $\theta \times$ distance to the moon from the Earth

$$= 3.317 \times 10^{-2} \times 3.84 \times 10^8 \text{ m}$$

$$= 1.274 \times 10^7 \text{ m}$$

1.4 Measurement of Mass:

Since 1889, a kilogram was the mass of a shiny piece of platinum-iridium alloy kept in a special glass case at the International Bureau of weights and measures. This definition of mass has been modified on 20th May 2019, the reason being that the carefully kept platinum-iridium piece is seen to pick up micro particles of dirt and is also affected by the atmosphere causing its mass to change. The new measure of kilogram is defined in terms of magnitude of electric current. We know that electric current can be used to make an electromagnet. An electromagnet attracts magnetic materials and is thus used in research and in industrial applications such as cranes to lift heavy pieces of iron/steel. Thus the kilogram mass can be described in terms of the amount of current which has to be passed through an electromagnet so that it can pull down one side of an extremely sensitive balance to balance the other side which holds one standard kg mass.

While dealing with mass of atoms and molecules, kg is an inconvenient unit. Therefore, their mass is measured in atomic mass unit. It will be easy to compare mass of any atom in terms of mass of some standard atom which has been decided internationally to be C^{12} atom. The $(1/12)^{\text{th}}$ mass of an unexcited atom of C^{12} is called atomic mass unit (amu).

$1 \text{ amu} = 1.6605402 \times 10^{-27} \text{ kg}$ with an uncertainty of 10 in the last two decimal places.

1.5 Measurement of Time:

The SI unit of time is the second (s). For many years, duration of one mean Solar day was considered as reference. A mean Solar day is the average time interval from one noon to the next noon. Average duration of a day is taken as 24 hours. One hour is of 60 minutes

and each minute is of 60 seconds. Thus a mean Solar day = 24 hours = $24 \times 60 \times 60 = 86400$ s. Accordingly a second was defined as $1/86400$ of a mean Solar day.

It was later observed that the length of a Solar day varies gradually due to the gradual slowing down of the Earth's rotation. Hence, to get more standard and nonvarying (constant) unit for measurement of time, a cesium atomic clock is used. It is based on periodic vibrations produced in cesium atom. In cesium atomic clock, a second is taken as the time needed for 9,192,631,770 vibrations of the radiation (wave) emitted during a transition between two hyperfine states of Cs^{133} atom.



Do you know ?

Why is only carbon used and not any other element for defining atomic mass unit? Carbon 12 (C^{12}) is the most abundant isotope of carbon and the most stable one. Around 98% of the available carbon is C^{12} isotope.

Earlier, oxygen and hydrogen were used as the standard atoms. But various isotopes of oxygen and hydrogen are present in higher proportion and therefore it is more accurate to use C^{12} .

1.6 Dimensions and Dimensional Analysis:

As mentioned earlier, a derived physical quantity can be expressed in terms of some combination of seven basic or fundamental

quantities. For convenience, the basic quantities are represented by symbols as 'L' for length, 'M' for mass, 'T' for time, 'K' For temperature, 'I' for current, 'C' for luminous intensity and 'mol' for amount of mass.

The dimensions of a physical quantity are the powers to which the concerned fundamental units must be raised in order to obtain the unit of the given physical quantity.

When we represent any derived quantity with appropriate powers of symbols of the fundamental quantities, then such an expression is called dimensional formula. This dimensional formula is expressed by square bracket and no comma is written in between any of the symbols.

Illustration:

(i) Dimensional formula of velocity

$$\text{Velocity} = \frac{\text{displacement}}{\text{time}}$$

$$\text{Dimensions of velocity} = \frac{[L]}{[T]} = [L^1 M^0 T^{-1}]$$

ii) Dimensional formula of velocity gradient

$$\text{velocity gradient} = \frac{\text{velocity}}{\text{distance}}$$

Dimensions of velocity gradient

$$= \frac{[LT^{-1}]}{[L]} = [L^0 M^0 T^{-1}]$$

iii) Dimensional formula for charge.

$$\text{charge} = \text{current} \times \text{time}$$

$$\text{Dimensions of charge} = [I] [T] = [L^0 M^0 T^1 I^1]$$

Table 1.3: Some Common Physical Quantities their, SI Units and Dimensions

S. No	Physical quantity	Formula	SI unit	Dimensional formula
1	Density	$\rho = M/V$	kilogram per cubic metre (kg/m^3)	$[L^{-3} M^1 T^0]$
2	Acceleration	$a = v/t$	metre per second square (m/s^2)	$[L^1 M^0 T^{-2}]$
3	Momentum	$P = mv$	kilogram metre per second (kg m/s)	$[L^1 M^1 T^{-1}]$
4	Force	$F = ma$	kilogram metre per second square (kg m/s^2) or newton (N)	$[L^1 M^1 T^{-2}]$
5	Impulse	$J = F \cdot t$	newton second (Ns)	$[L^1 M^1 T^{-1}]$
6	Work	$W = F \cdot s$	joule (J)	$[L^2 M^1 T^{-2}]$
7	Kinetic Energy	$KE = 1/2 mv^2$	joule (J)	$[L^2 M^1 T^{-2}]$
8	Pressure	$P = F/A$	kilogram per metre second square (kg/ms^2)	$[L^{-1} M^1 T^{-2}]$

Table 1.3 gives the dimensions of various physical quantities commonly used in mechanics.

1.6.1 Uses of Dimensional Analysis:

- (i) **To check the correctness of physical equations:** In any equation relating different physical quantities, if the dimensions of all the terms on both the sides are the same then that equation is said to be dimensionally correct. This is called the principle of homogeneity of dimensions. Consider the first equation of motion.

$$v = u + at$$

$$\text{Dimension of L.H.S} = [v] = [LT^{-1}]$$

$$[u] = [LT^{-1}]$$

$$[at] = [LT^{-2}] \quad [T] = [LT^{-1}]$$

$$\text{Dimension of R.H.S} = [LT^{-1}] + [LT^{-1}]$$

$$[L.H.S] = [R.H.S]$$

As the dimensions of L.H.S and R.H.S are the same, the given equation is dimensionally correct.

- (ii) **To establish the relationship between related physical quantities:** The period T of oscillation of a simple pendulum depends on length l and acceleration due to gravity g . Let us derive the relation between T, l, g :

$$\text{Suppose } T \propto l^a$$

$$\text{and } T \propto g^b$$

$$\therefore T \propto l^a g^b$$

$$T = k l^a g^b,$$

where k is constant of proportionality and it is a dimensionless quantity and a and b are rational numbers. Equating dimensions on both sides,

$$[M^0 L^0 T^1] = k [L^1]^a [LT^{-2}]^b \\ = k [L^{a+b} T^{-2b}]$$

$$[L^0 T^1] = k [L^{a+b} T^{-2b}]$$

Comparing the dimensions of the corresponding quantities on both the sides we get

$$a + b = 0$$

$$\therefore a = -b$$

and

$$-2b = 1$$

$$\therefore b = -1/2$$

$$\therefore a = -b = -(-1/2)$$

$$\therefore a = 1/2$$

$$\therefore T = k l^{1/2} g^{-1/2}$$

$$\therefore T = k \sqrt{l/g}$$

The value of k is determined experimentally and is found to be 2π

$$\therefore T = 2\pi \sqrt{l/g}$$

- (iii) **To find the conversion factor between the units of the same physical quantity in two different systems of units:** Let us use dimensional analysis to determine the conversion factor between joule (SI unit of work) and erg (CGS unit of work).

$$\text{Let } 1 \text{ J} = x \text{ erg}$$

Dimensional formula for work is $[M^1 L^2 T^{-2}]$
Substituting in the above equation, we can write

$$[M_1^1 L_1^2 T_1^{-2}] = x [M_2^1 L_2^2 T_2^{-2}]$$

$$x = \frac{[M_1^1 L_1^2 T_1^{-2}]}{[M_2^1 L_2^2 T_2^{-2}]}$$

$$\text{or, } x = \left(\frac{M_1}{M_2} \right)^1 \left(\frac{L_1}{L_2} \right)^2 \left(\frac{T_1}{T_2} \right)^{-2}$$

where suffix 1 indicates SI units and 2 indicates CGS units.

In SI units, L, M, T are expressed in m, kg and s and in CGS system L, M, T are represented in cm, g and s respectively.

$$\therefore x = \left(\frac{\text{kg}}{\text{g}} \right)^1 \left(\frac{\text{m}}{\text{cm}} \right)^2 \left(\frac{\text{s}}{\text{s}} \right)^{-2}$$

$$\text{or } x = \left(10^3 \frac{\text{g}}{\text{g}} \right)^1 \left(100 \frac{\text{cm}}{\text{cm}} \right)^2 (1)^{-2}$$

$$\therefore x = (10^3)(10^4) = 10^7$$

$$\therefore 1 \text{ joule} = 10^7 \text{ erg}$$

Example 1.4: A calorie is a unit of heat and it equals 4.2 J, where $1 \text{ J} = \text{kg m}^2 \text{ s}^{-2}$. A distant civilisation employs a system of units in which the units of mass, length and time are $\alpha \text{ kg}, \beta \text{ m}$

and γ s. Also J' is their unit of energy. What will be the magnitude of calorie in their units?

Solution: Let us write the new units as A, B and C for mass, length and time respectively.

We are given

$$1 \text{ A} = \alpha \text{ kg}$$

$$1 \text{ B} = \beta \text{ m}$$

$$1 \text{ C} = \gamma \text{ s}$$

$$1 \text{ cal} = 4.2 \text{ J} = 4.2 \text{ kg m}^2 \text{ s}^{-2}$$

$$= 4.2 \left(\frac{\text{A}}{\alpha} \right) \left(\frac{\text{B}}{\beta} \right)^2 \left(\frac{\text{C}}{\gamma} \right)^{-2}$$

$$= \frac{4.2 \gamma^2}{\alpha \beta^2} \text{ AB}^2 \text{ C}^{-2}$$

$$= \frac{4.2 \gamma^2}{\alpha \beta^2} \text{ J'}$$

$$\text{Thus in the new units, 1 calorie is} = \frac{4.2 \gamma^2}{\alpha \beta^2} \text{ J'}$$

1.6.2 Limitations of Dimensional Analysis:

- 1) The value of dimensionless constant can be obtained with the help of experiments only.
- 2) Dimensional analysis can not be used to derive relations involving trigonometric, exponential, and logarithmic functions as these quantities are dimensionless.
- 3) This method is not useful if constant of proportionality is not a dimensionless quantity.

Illustration : Gravitational force between two point masses is directly proportional to product of the two masses and inversely proportional to square of the distance between the two

$$\therefore F \propto \frac{m_1 m_2}{r^2}$$

$$\text{Let } F = G \frac{m_1 m_2}{r^2}$$

The constant of proportionality 'G' is NOT dimensionless. Thus earlier method will not work.

- 4) If the correct equation contains some more terms of the same dimension, it is not possible to know about their presence using dimensional equation. For example, with

standard symbols, the equation $S = \frac{1}{2} at^2$ is dimensionally correct. However, the complete equation is $S = ut + \frac{1}{2} at^2$

1.7 Accuracy, Precision and Uncertainty in Measurement:

Physics is a science based on observations and experiments. Observations of various physical quantities are made during an experiment. For example, during the atmospheric study we measure atmospheric pressure, wind velocity, humidity, etc. All the measurements may be accurate, meaning that the measured values are the same as the true values. Accuracy is how close a measurement is to the actual value of that quantity. These measurements may be precise, meaning that multiple measurements give nearly identical values (i.e., reproducible results). In actual measurements, an observation may be both accurate and precise or neither accurate nor precise. The goal of the observer should be to get accurate as well as precise measurements.

Possible uncertainties in an observation may arise due to following reasons:

- 1) Quality of instrument used.
- 2) Skill of the person doing the experiment.
- 3) The method used for measurement.
- 4) External or internal factors affecting the result of the experiment.



Can you tell?

If ten students are asked to measure the length of a piece of cloth up to a mm, using a metre scale, do you think their answers will be identical? Give reasons.

1.8 Errors in Measurements:

Faulty measurements of physical quantity can lead to errors. The errors are broadly divided into the following two categories :

a) Systematic errors : Systematic errors are errors that are not determined by chance but are introduced by an inaccuracy (involving

either the observation or measurement process) inherent to the system. Sources of systematic error may be due to imperfect calibration of the instrument, and sometimes imperfect method of observation.

Each of these errors tends to be in one direction, either positive or negative. The sources of systematic errors are as follows:

(i) **Instrumental error:** This type of error arises due to defective calibration of an instrument, for example an incorrect zeroing of an instrument will lead to such kind of error ('zero' of a thermometer not graduated at proper place, the pointer of weighing balance in the laboratory already indicating some value instead of showing zero when no load is kept on it, an ammeter showing a current of 0.5 amp even when not connected in circuit, etc).

(ii) **Error due to imperfection in experimental technique:** This is an error due to defective setting of an instrument. For example the measured volume of a liquid in a graduated tube will be inaccurate if the tube is not held vertical.

(iii) **Personal error:** Such errors are introduced due to fault of the observer. Bias of the observer, carelessness in taking observations etc. could result in such errors. For example, while measuring the length of an object with a ruler, it is necessary to look at the ruler from directly above. If the observer looks at it from an angle, the measured length will be wrong due to parallax.

Systematic errors can be minimized by using correct instrument, following proper experimental procedure and removing personal error.

b) Random errors: These are the errors which are introduced even after following all the procedures to minimize systematic errors. These type of errors may be positive or negative. These errors can not be eliminated completely but we can minimize them by repeated observations and then taking their mean (average). Random errors occur due to variation in conditions in

which experiment is performed. For example, the temperature may change during the course of an experiment, pressure of any gas used in the experiment may change, or the voltage of the power supply may change randomly, etc.

1.8.1 Estimation of error:

Suppose the readings recorded repeatedly for a physical quantity during a measurement are

$$a_1, a_2, a_3, \dots, a_n.$$

Arithmetic mean a_{mean} is given by

$$a_{\text{mean}} = \frac{a_1 + a_2 + a_3 + \dots + a_n}{n}$$

$$a_{\text{mean}} = \frac{1}{n} \sum_{i=1}^n a_i \quad \text{--- (1.3)}$$

This is the most probable value of the quantity. The magnitude of the difference between mean value and each individual value is called **absolute error** in the observations.

Thus for ' a_1 ', the absolute error Δa_1 is given by

$$\Delta a_1 = |a_{\text{mean}} - a_1|,$$

for a_2 ,

$$\Delta a_2 = |a_{\text{mean}} - a_2|$$

and so for a_n it will be

$$\Delta a_n = |a_{\text{mean}} - a_n|$$

The arithmetic mean of all the absolute errors is called **mean absolute error** in the measurement of the physical quantity.

$$\Delta a_{\text{mean}} = \frac{\Delta a_1 + \Delta a_2 + \dots + \Delta a_n}{n}$$

$$= \frac{1}{n} \sum_{i=1}^n \Delta a_i \quad \text{--- (1.4)}$$

The measured value of the physical quantity a can then be represented by

$a = a_{\text{mean}} \pm \Delta a_{\text{mean}}$ which tells us that the actual value of ' a ' could be between $a_{\text{mean}} - \Delta a_{\text{mean}}$ and $a_{\text{mean}} + \Delta a_{\text{mean}}$. The ratio of mean absolute error to its arithmetic mean value is called **relative error**.

$$\text{Relative error} = \frac{\Delta a_{\text{mean}}}{a_{\text{mean}}} \quad \text{--- (1.5)}$$

When relative error is represented as percentage it is called **percentage error**.

$$\text{Percentage error} = \frac{\Delta a_{\text{mean}}}{a_{\text{mean}}} \times 100 \quad \text{--- (1.6)}$$



Activity :

Perform an experiment using a Vernier callipers of least count 0.01cm to measure the external diameter of a hollow cylinder. Take 3 readings at different position on the cylinder and find (i) the mean diameter (ii) the absolute mean error and (iii) the percentage error in the measurement of diameter.

Example 1.5: The radius of a sphere measured repeatedly yields values 5.63 m, 5.54 m, 5.44 m, 5.40 m and 5.35 m. Determine the most probable value of radius and the mean absolute, relative and percentage errors.

Solution: Most probable value of radius is its arithmetic mean

$$= \frac{5.63 + 5.54 + 5.44 + 5.40 + 5.35}{5} \text{ m}$$

$$= 5.472 \text{ m}$$

Mean absolute error

$$= \frac{1}{5} \left\{ \begin{array}{l} |5.63 - 5.472| + |5.54 - 5.472| \\ + |5.44 - 5.472| + |5.40 - 5.472| \\ + |5.35 - 5.472| \end{array} \right\} \text{ m}$$

$$= \frac{0.452}{5} = 0.0904 \text{ m}$$

$$\text{Relative error} = \frac{0.0904}{5.472} = 0.017$$

$$\% \text{ error} = 1.7\%$$

1.8.2 Combination of errors:

When we do an experiment and measure various physical quantities associated with the experiment, we must know how the errors from individual measurement combine to give errors in the final result. For example, in the measurement of the resistance of a conductor using Ohms law, there will be an error in the measurement of potential difference and that of current. It is important to study how these errors combine to give the error in the measurement of

resistance.

a) Errors in sum and in difference:

Suppose two physical quantities A and B have measured values $A \pm \Delta A$ and $B \pm \Delta B$, respectively, where ΔA and ΔB are their mean absolute errors. We wish to find the absolute error ΔZ in their sum.

$$Z = A + B$$

$$Z \pm \Delta Z = (A \pm \Delta A) + (B \pm \Delta B)$$

$$= (A + B) \pm \Delta A \pm \Delta B$$

$$\pm \Delta Z = \pm \Delta A \pm \Delta B,$$

For difference, i.e., if $Z = A - B$,

$$Z \pm \Delta Z = (A \pm \Delta A) - (B \pm \Delta B)$$

$$= (A - B) \pm \Delta A \pm \Delta B$$

$$\pm \Delta Z = \pm \Delta A \pm \Delta B,$$

There are four possible values for ΔZ , namely $(+ \Delta A - \Delta B)$, $(+\Delta A + \Delta B)$, $(-\Delta A - \Delta B)$, $(-\Delta A + \Delta B)$. Hence maximum value of absolute error is $\Delta Z = \Delta A + \Delta B$ in both the cases.

When two quantities are added or subtracted, the absolute error in the final result is the *sum* of the absolute errors in the individual quantities.

b) Errors in product and in division:

Suppose $Z = AB$ and measured values of A and B are $(A \pm \Delta A)$ and $(B \pm \Delta B)$ Then

$$Z \pm \Delta Z = (A \pm \Delta A) (B \pm \Delta B)$$

$$= AB \pm A\Delta B \pm B\Delta A \pm \Delta A\Delta B$$

Dividing L.H.S by Z and R.H.S. by AB we get

$$\left(1 \pm \frac{\Delta Z}{Z} \right) = \left(1 \pm \frac{\Delta B}{B} \pm \frac{\Delta A}{A} \pm \left(\frac{\Delta A}{A} \right) \left(\frac{\Delta B}{B} \right) \right)$$

Since $\Delta A/A$ and $\Delta B/B$ are very small we shall neglect their product. Hence maximum *relative* error in Z is

$$\frac{\Delta Z}{Z} = \frac{\Delta A}{A} + \frac{\Delta B}{B} \quad \text{--- (1.7)}$$

This formula also applies to the division of two quantities.

Thus, when two quantities are multiplied or divided, the maximum relative error in the result is the sum of *relative* errors in each quantity.

c) Errors due to the power (index) of measured quantity:

Suppose

$$Z = A^3 = A.A.A$$

$$\frac{\Delta Z}{Z} = \frac{\Delta A}{A} + \frac{\Delta A}{A} + \frac{\Delta A}{A}$$

from the multiplication rule above.

Hence the relative error in $Z = A^3$ is three times the relative error in A . So if $Z = A^n$

$$\frac{\Delta Z}{Z} = n \frac{\Delta A}{A} \quad \text{--- (1.8)}$$

$$\text{In general, if } Z = \frac{A^p B^q}{C^r}$$

$$\frac{\Delta Z}{Z} = p \frac{\Delta A}{A} + q \frac{\Delta B}{B} + r \frac{\Delta C}{C} \quad \text{--- (1.9)}$$

The quantity in the formula which has large power is responsible for maximum error.

Example 1.6: In an experiment to determine the volume of an object, mass and density are recorded as $m = (5 \pm 0.15) \text{ kg}$ and $\rho = (5 \pm 0.2) \text{ kg m}^{-3}$ respectively. Calculate percentage error in the measurement of volume.

Solution : We know,

$$\text{Density} = \frac{\text{Mass}}{\text{Volume}}$$

$$\therefore \text{Volume} = \frac{\text{Mass}}{\text{Density}} = \frac{M}{\rho}$$

$$\begin{aligned} \text{Percentage error in volume} &= \left(\frac{\Delta m}{m} + \frac{\Delta \rho}{\rho} \right) \times 100 \\ &= \left(\frac{0.15}{5} + \frac{0.2}{5} \right) \times 100 \\ &= (0.03 + 0.04) \times 100 \\ &= (0.07) \times 100 = 7\% \end{aligned}$$

Example 1.7: The acceleration due to gravity is determined by using a simple pendulum of length $l = (100 \pm 0.1) \text{ cm}$. If its time period is $T = (2 \pm 0.01) \text{ s}$, find the maximum percentage error in the measurement of g .

Solution: The time period of oscillation of a simple pendulum is given by

$$T = 2\pi \sqrt{\frac{l}{g}}$$

Squaring both sides

$$T^2 = 4\pi^2 l / g$$

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

$$\text{Now } \Delta l = 0.1, l = 100 \text{ cm}, \Delta T = 0.01 \text{ s}, T = 2 \text{ s}$$

$$\text{Percentage error} = \frac{\Delta g \times 100}{g}$$

$$= \left(\frac{\Delta l}{l} + \frac{2\Delta T}{T} \right) \times 100$$

$$= \left(\frac{0.1}{100} + \frac{2 \times 0.01}{2} \right) \times 100$$

$$= (0.001 + 0.01) \times 100 = 1.1$$

Percentage error in measurement of g is 1.1%

1.9 Significant Figures:

In the previous sections, we have studied various types of errors, their origins and the ways to minimize them. Our accuracy is limited to the least count of the instrument used during the measurement. Least count is the smallest measurement that can be made using the given instrument. For example with the usual metre scale, one can measure 0.1 cm as the least value. Hence its least count is 0.1cm.

Suppose we measure the length of a metal rod using a metre scale of least count 0.1cm. The measurement is done three times and the readings are 15.4, 15.4, and 15.5 cm. The most probable length which is the arithmetic mean as per our earlier discussion is 15.43. Out of this we are certain about the digits 1 and 5 but are not certain about the last 2 digits because of the least count limitation.

The number of digits in a measurement about which we are certain, plus one additional digit, the first one about which we are not certain is known as significant figures or significant digits.

Thus in above example, we have 3 significant digits 1, 5 and 4.

The larger the number of significant figures obtained in a measurement, the greater is the accuracy of the measurement. If one uses the instrument of smaller least count, the number of significant digits increases.

Rules for determining significant figures

- 1) All the nonzero digits are significant, for example if the volume of an object is 178.43 cm^3 , there are five significant digits which are 1, 7, 8, 4 and 3.
- 2) All the zeros between two nonzero digits are significant, eg., $m = 165.02 \text{ g}$ has 5 significant digits.
- 3) If the number is less than 1, the zero/zeros on the right of the decimal point and to the left of the first nonzero digit are not significant e.g. in 0.001405, the underlined zeros are not significant. Thus the above number has four significant digits.
- 4) The zeros on the right hand side of the last nonzero number are significant (but for this, the number must be written with a decimal point), e.g. 1.500 or 0.01500 have both 4 significant figures each.

On the contrary, if a measurement yields length L given as

$L = 125 \text{ m} = 12500 \text{ cm} = 125000 \text{ mm}$, it has only three significant digits.

To avoid the ambiguities in determining the number of significant figures, it is necessary to report every measurement in scientific notation (i.e., in powers of 10) i.e., by using the concept of order of magnitude.

The magnitude of any physical quantity can be expressed as $A \times 10^n$ where 'A' is a number such that $0.5 \leq A < 5$ and 'n' is an integer called the **order of magnitude**.

(i) radius of Earth = 6400 km
 $= 0.64 \times 10^7 \text{ m}$

The order of magnitude is 7 and the number of significant figures are 2.

(ii) Magnitude of the charge on electron e
 $= 1.6 \times 10^{-19} \text{ C}$

Here the order of magnitude is -19 and the number of significant digits are 2.



Internet my friend

1. videlectures.net/mit801f99_lewin_lec01/
2. hyperphysics.phy-astr.gsu.edu/hbase/hframe.html

Definitions of SI Units

Till May 20, 2019 the *kilogram* did not have a definition; it was mass of the prototype cylinder kept under controlled conditions of temperature and pressure at the SI museum at Paris. A rigorous and meticulous experimentation has shown that the mass of the *standard* prototype for the *kilogram* has changed in the course of time. This shows the acute necessity for standardisation of units. The new definitions aim to improve the SI without changing the size of any units, thus ensuring continuity with existing measurements. In November 2018, the 26th General Conference on Weights and Measures (CGPM) unanimously approved these changes, which the International Committee for Weights and Measures (CIPM) had proposed earlier that year. These definitions came in force from 20 May 2019.

- (i) As per new SI units, each of the seven fundamental units (metre, kilogram, etc.) uses **one** of the following 7 constants which are proposed to be having **exact values** as given below:

The Planck constant,
 $h = 6.62607015 \times 10^{-34} \text{ joule-second}$
(J s **or** $\text{kg m}^2 \text{s}^{-1}$).

The elementary charge,
 $e = 1.602176634 \times 10^{-19} \text{ coulomb}$ (**C or A s**).

The Boltzmann constant,
 $k = 1.380649 \times 10^{-23} \text{ joule per kelvin}$
(J K^{-1} **or** $\text{kg m}^2 \text{s}^{-2} \text{K}^{-1}$).

The Avogadro constant (number),
 $N_A = 6.02214076 \times 10^{23} \text{ reciprocal mole}$
(mol^{-1}).

The speed of light in vacuum,
 $c = 299792458 \text{ metre per second}$ (m s^{-1}).

The ground state hyperfine structure transition frequency of Caesium-133 atom,

$\Delta\nu_{\text{Cs}} = 9192631770 \text{ hertz}$ (Hz **or** s^{-1}).

The luminous efficacy of monochromatic radiation of frequency $540 \times 10^{12} \text{ Hz}$, $K_{\text{cd}} = 683 \text{ lumen per watt}$ ($\text{lm} \cdot \text{W}^{-1}$) = $683 \text{ cd sr s}^3 \text{ kg}^{-1} \text{ m}^{-2}$, where sr is steradian; the SI unit of solid angle.

- (ii) Definitions of the units *second* and *mole* are based only upon their respective constants, for example (a) the *second* uses ground state hyperfine structure transition frequency of Caesium-133 atom to be exactly 9192631770 hertz. Thus, the *second* is defined as 9192631770 periods of that transition, (b) the *mole* uses Avogadro's number to be $N_A = 6.02214076 \times 10^{23}$. Thus, one *mole* is that amount of substance which contains exactly $6.02214076 \times 10^{23}$ molecules.
- (iii) Definitions of all the other fundamental units use one constant each and at least one other fundamental unit, for example, the *metre* makes use of speed of light in vacuum as a constant and *second* as fundamental unit. Thus, *metre* is defined as the distance traveled by the light in vacuum in exactly $1/299792458$ *second*.
- (iv) The figures show the dependency of various units upon their respective constants and other units (wherever

used). The arrows arriving at that unit refer to the constant and the fundamental unit (or units, wherever used) for defining that unit. The arrows going away from a unit indicate other units which use this unit for their definition.

For example, as described above, in fig (a) i) the arrows directed to *metre* are from *second* and *c*. The arrows going away from the *metre* indicate that the *metre* is used in defining the *kilogram* the *candela* and the *kelvin*, (ii) the newly defined unit *kilogram* uses Planck constant, the *metre* and the *second*, while the *kilogram* itself is used in defining the *kelvin* and the *candela*. This definition relates the *kilogram* to the equivalent mass of the energy of a photon given its frequency, via the Planck constant.

Figure (a) refers to new definitions while the figure (b) refers to the corresponding definitions before 20 May 2019. Interested students may compare them to know which definitions are modified and how.

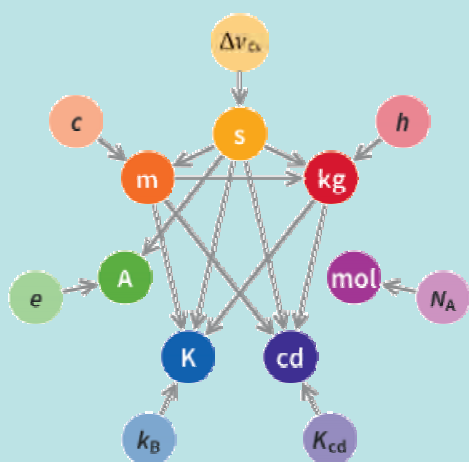


Fig (a) New SI

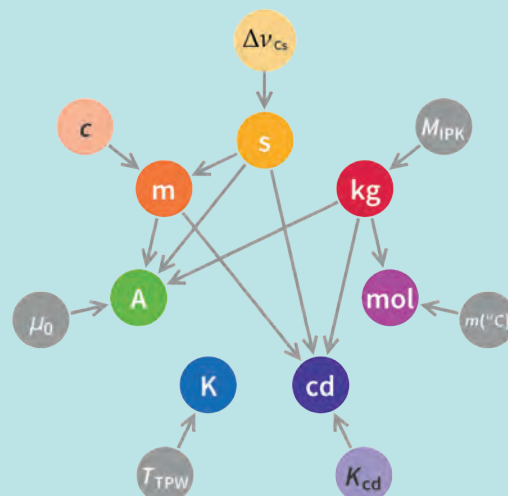


Fig (b) Old SI



Exercises

1. Choose the correct option.

- $[L^1M^1T^{-2}]$ is the dimensional formula for
(A) Velocity (B) Acceleration
(C) Force (D) Work
- The error in the measurement of the sides of a rectangle is 1%. The error in the measurement of its area is
(A) 1% (B) $1/2\%$
(C) 2% (D) None of the above.
- Light year is a unit of
(A) Time (B) Mass
(C) Distance (D) Luminosity
- Dimensions of kinetic energy are the same as that of
(A) Force (B) Acceleration
(C) Work (D) Pressure
- Which of the following is not a fundamental unit?
(A) cm (B) kg
(C) centigrade (D) volt

2. Answer the following questions.

- Star A is farther than star B. Which star will have a large parallax angle?
- What are the dimensions of the quantity $l\sqrt{l/g}$, l being the length and g the acceleration due to gravity?
- Define absolute error, mean absolute error, relative error and percentage error.
- Describe what is meant by significant figures and order of magnitude.
- If the measured values of two quantities are $A \pm \Delta A$ and $B \pm \Delta B$, ΔA and ΔB being the mean absolute errors. What is the maximum possible error in $A \pm B$? Show that if $Z = \frac{A}{B}$
$$\frac{\Delta Z}{Z} = \frac{\Delta A}{A} + \frac{\Delta B}{B}$$
- Derive the formula for kinetic energy of a particle having mass m and velocity v using dimensional analysis

3. Solve numerical examples.

- The masses of two bodies are measured to be 15.7 ± 0.2 kg and 27.3 ± 0.3 kg. What is the total mass of the two and the error in it?
[Ans : 43 kg, ± 0.5 kg]
- The distance travelled by an object in time (100 ± 1) s is (5.2 ± 0.1) m. What is the speed and its relative error?
[Ans : 0.052 ms^{-1} , $\pm 0.0292 \text{ ms}^{-1}$]
- An electron with charge e enters a uniform magnetic field $\vec{B}$ with a velocity $\vec{v}$. The velocity is perpendicular to the magnetic field. The force on the charge e is given by $|\vec{F}| = Bev$. Obtain the dimensions of $\vec{B}$.
[Ans: $[L^0M^1T^{-2}I^{-1}]$]
- A large ball 2 m in radius is made up of a rope of square cross section with edge length 4 mm. Neglecting the air gaps in the ball, what is the total length of the rope to the nearest order of magnitude?
[Ans : $\approx 10^6 \text{ m} = 10^3 \text{ km}$]
- Nuclear radius R has a dependence on the mass number (A) as $R = 1.3 \times 10^{-16} A^{1/3}$ m. For a nucleus of mass number $A=125$, obtain the order of magnitude of R expressed in metre.
[Ans : -15]
- In a workshop a worker measures the length of a steel plate with a Vernier callipers having a least count 0.01 cm. Four such measurements of the length yielded the following values: 3.11 cm, 3.13 cm, 3.14 cm, 3.14 cm. Find the mean length, the mean absolute error and the percentage error in the measured value of the length.
[Ans: 3.13 cm, 0.01 cm, 0.32%]

- vii) Find the percentage error in kinetic energy of a body having mass 60.0 ± 0.3 g moving with a velocity 25.0 ± 0.1 cm/s.
[Ans: 1.3%]
- viii) In Ohm's experiments, the values of the unknown resistances were found to be 6.12Ω , 6.09Ω , 6.22Ω , 6.15Ω . Calculate the mean absolute error, relative error and percentage error in these measurements.
[Ans: 0.04Ω , 0.0065Ω , 0.65%]
- ix) An object is falling freely under the gravitational force. Its velocity after travelling a distance h is v . If v depends upon gravitational acceleration g and distance, prove with dimensional analysis that $v = k\sqrt{gh}$ where k is a constant.
- x) $v = at + \frac{b}{t+c} + v_0$ is a dimensionally valid equation. Obtain the dimensional formula for a , b and c where v is velocity, t is time and v_0 is initial velocity.
[Ans: a - $[L^1 M^0 T^{-2}]$, b - $[L^1 M^0 T^0]$, c - $[L^0 M^0 T^1]$]
- xi) The length, breadth and thickness of a rectangular sheet of metal are 4.234 m, 1.005 m, and 2.01 cm respectively. Give the area and volume of the sheet to correct significant figures.
[Ans: 4.255 m^2 , 8.552 m^3]
- xii) If the length of a cylinder is $l = (4.00 \pm 0.001)$ cm, radius $r = (0.0250 \pm 0.001)$ cm and mass $m = (6.25 \pm 0.01)$ gm. Calculate the percentage error in the determination of density.
[Ans: 8.185%]
- xiii) When the planet Jupiter is at a distance of 824.7 million kilometers from the Earth, its angular diameter is measured to be $35.72''$ of arc. Calculate the diameter of the Jupiter.
[Ans: 1.428×10^5 km]
- xiv) If the formula for a physical quantity is $X = \frac{a^4 b^3}{c^{1/3} d^{1/2}}$ and if the percentage error in the measurements of a , b , c and d are 2% , 3% , 3% and 4% respectively. Calculate percentage error in X .
[Ans: 20%]
- xv) Write down the number of significant figures in the following: 0.003 m^2 , $0.1250 \text{ gm cm}^{-2}$, $6.4 \times 10^6 \text{ m}$, $1.6 \times 10^{-19} \text{ C}$, $9.1 \times 10^{-31} \text{ kg}$.
[Ans: 1, 4, 2, 2, 2]
- xvi) The diameter of a sphere is 2.14 cm. Calculate the volume of the sphere to the correct number of significant figures.
[Ans: 5.13 cm^3]



Can you recall?

1. What is the difference between a scalar and a vector?
2. Which of the following are scalars or vectors?
 - (i) displacements (ii) distance travelled (iii) velocity
 - (iv) speed (v) force (vi) work done (vii) energy

2.1 Introduction:

You will need certain mathematical tools to understand the topics covered in this book. Vector analysis and elementary calculus are two among these. You will learn calculus in details, in mathematics, in the XIIth standard. In this Chapter, you are going to learn about vector analysis and will have a preliminary introduction to calculus which should be sufficient for you to understand the physics that you will learn in this book.

2.2 Vector Analysis:

In the previous Chapter, you have studied different aspects of physical quantities like their division into fundamental and derived quantities and their units and dimensions. You also need to understand that all physical quantities may not be fully described by their magnitudes and units alone. For example if you are given the time for which a man has walked with a certain speed, you can find the distance travelled by the man, but you cannot find out where exactly the man has reached unless you know the direction in which the man has walked.

Therefore, you can say that some physical quantities, which are called scalars, can be described with magnitude alone, whereas some other physical quantities, which are called vectors, need to be described with magnitude as well as direction. In the above example the distance travelled by the man is a scalar quantity while the final position of the man relative to his initial position, i.e., his displacement can be described by magnitude and direction and is a vector quantity. In this Chapter you will study different aspects of scalar and vector quantities.

2.2.1 Scalars:

Physical quantities which can be completely

described by their magnitude are called scalars, i.e. they are specified by a number and a unit. For example when we say that a given metal rod has a length 2 m, it indicates that the rod is two times longer than a certain standard unit *metre*. Thus the number 2 is the magnitude and metre is the unit; together they give us a complete idea about the length of the rod. Thus length is a scalar quantity. Similarly mass, time, temperature, density, etc., are examples of scalars. Scalars can be added or subtracted by rules of simple algebra.

2.2.2 Vectors:

Physical quantities which need magnitude as well as direction for their complete description are called vectors. Examples of vectors are displacement, velocity, force etc.

A vector can be represented by a directed line segment or by an arrow. The length of the line segment drawn to scale gives the magnitude of the vector, e.g., displacement of a body from P to Q can be represented as $P \longrightarrow Q$, where the starting point P is called the tail and the end point Q (arrow head) is called the head of the vector. Symbolically we write it as $\overrightarrow{PQ}$. Symbolically vectors are also represented by a single capital letter with an arrow above it, e.g., $\vec{X}$, $\vec{A}$, etc. Magnitude of a vector $\vec{X}$ is written as $|\vec{X}|$.

Let us see a few examples of different types of vectors.

(a) Zero vector (Null vector): A vector having zero magnitude with a particular direction (arbitrary) is called zero vector. Symbolically it is represented as $\vec{0}$.

(1) Velocity vector of a stationary particle is a zero vector.

(2) The acceleration vector of an object

moving with uniform velocity is a zero vector.

- (b) **Resultant vector:** The resultant of two or more vectors is that single vector, which produces the same effect, as produced by all the vectors together.
- (c) **Negative vector (opposite vector):** A negative vector of a given vector is a vector of the same magnitude but opposite in direction to that of the given vector.

In Fig. 2.1, $\vec{B}$ is a negative vector to $\vec{A}$.

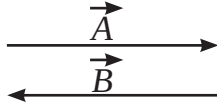


Fig. 2.1: Negative vector.

- (d) **Equal vector:** Two vectors A and B representing same physical quantity are said to be equal if and only if they have the same magnitude and direction. Two equal vectors are shown in Fig. 2.2.

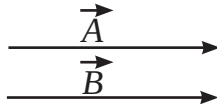


Fig. 2.2: Equal vectors.

- (e) **Position vector:** A vector which gives the position of a particle at a point with respect to the origin of a chosen co-ordinate system is called the position vector of the particle.

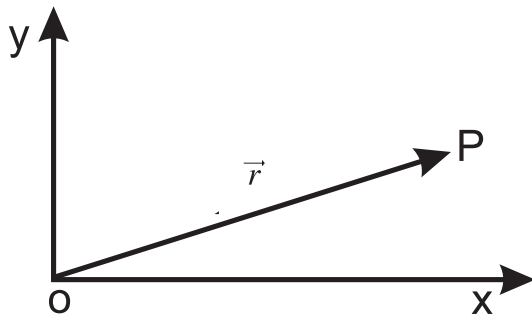


Fig 2.3: Position vector.

In Fig 2.3, $\vec{r} = \overrightarrow{OP}$ is the position vector of the particle present at P.

- (f) **Unit vector:** A vector having unit magnitude in a given direction is called a unit vector in that direction. If $\vec{M}$ is a non-zero vector i.e. its magnitude $M = |\vec{M}|$ is not zero, the unit vector along

$\vec{M}$ is written as $\hat{u}_M$ and is given by

$$\vec{M} = \hat{u}_M M \quad \text{--- (2.1)}$$

$$\text{or, } \hat{u}_M = \frac{\vec{M}}{M} \quad \text{--- (2.2)}$$

Hence $\hat{u}_M$ has magnitude unity and has the same direction as that of $\vec{M}$. We use $\hat{i}$, $\hat{j}$, and $\hat{k}$, respectively, as unit vectors along the x, y and z directions of a rectangular (three dimensional) coordinate system.

$$\begin{aligned} \hat{u}_x = \hat{i}, \hat{u}_y = \hat{j} \text{ and } \hat{u}_z = \hat{k} \\ \therefore \hat{i} = \frac{\vec{x}}{x}, \hat{j} = \frac{\vec{y}}{y} \text{ and } \hat{k} = \frac{\vec{z}}{z} \end{aligned} \quad \text{--- (2.3)}$$

Here $\vec{x}$, $\vec{y}$ and $\vec{z}$ are vectors along x, y and z axes, respectively.

2.3 Vector Operations:

2.3.1 Multiplication of a Vector by a Scalar:

Multiplying a vector $\vec{P}$ by a scalar quantity, say s , yields another vector. Let us write

$$\vec{Q} = s\vec{P} \quad \text{--- (2.4)}$$

$\vec{Q}$ will be a vector whose direction is the same as that of $\vec{P}$ and magnitude is s times the magnitude of $\vec{P}$.

2.3.2 Addition and Subtraction of Vectors:

The addition or subtraction of two or more vectors of the same type, i.e., describing the same physical quantity, gives rise to a single vector, such that the effect of this single vector is the same as the net effect of the vectors which have been added or subtracted.

It is important to understand that only vectors of the same type (describing same physical quantity) can be added or subtracted e.g. force $\vec{F}_1$ and force $\vec{F}_2$ can be added to give the resultant force $\vec{F} = \vec{F}_1 + \vec{F}_2$. But a force vector can not be added to a velocity vector.

It is easy to find addition of vectors $\vec{AB}$ and $\vec{BC}$ having the same or opposite direction but different magnitudes. If individual vectors are parallel (i.e., in the same direction), the magnitude of their resultant is the addition of individual magnitudes, i.e., $|\vec{AC}| = |\vec{AB}| + |\vec{BC}|$

and direction of the resultant is the same as that of the individual vectors as shown in Fig 2.4 (a). However, if the individual vectors are anti-parallel (i.e., in the opposite direction), the magnitude of their resultant is the difference of the individual magnitudes, and the direction is that of the larger vector i.e., $|\vec{AC}| = ||\vec{AB}| - |\vec{BC}||$ as shown in Fig. 2.4 (b).

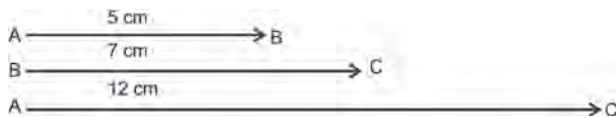


Fig. 2.4 (a): Resultant of parallel displacements.

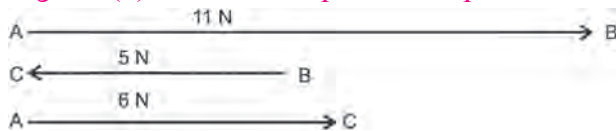


Fig 2.4 (b): Resultant of anti-parallel forces.

2.3.3 Triangle Law for Vector Addition:

When vectors of a given type do not act in the same or opposite directions, the resultant can be determined by using the triangle law of vector addition which is stated as follows:

If two vectors describing the same physical quantity are represented in magnitude and direction by the two sides of a triangle taken in order, then their resultant is represented in magnitude and direction by the third side of the triangle drawn in the opposite sense (from the starting point of first vector to the end point of the second vector).

Let $\vec{A}$ and $\vec{B}$ be two vectors in the plane of paper as shown in Fig. 2.5 (a). The sum of these two vectors can be obtained by using the triangle law described above as shown in Fig. 2.5 (b). The resultant vector is indicated by $\vec{C}$.

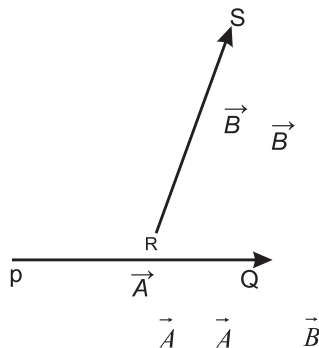


Fig. 2.5 (a): Two vectors $\vec{A}$ and $\vec{B}$ in a plane,

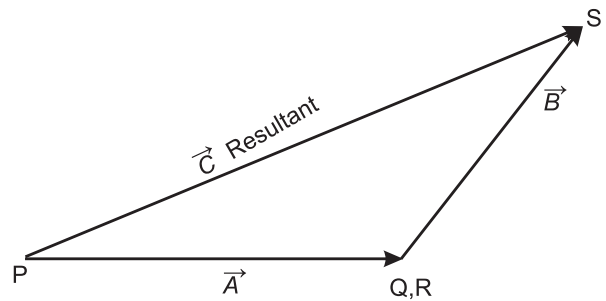


Fig. 2.5 (b): Resultant vector $\vec{C} = \vec{A} + \vec{B}$.

We can use the triangle law for showing that

(a) Vector addition is commutative.

For any two vectors $\vec{P}$ and $\vec{Q}$,

$$\vec{P} + \vec{Q} = \vec{Q} + \vec{P} \quad \text{--- (2.5)}$$

Figure 2.6 (a) shows addition of the two vector $\vec{P}$ and $\vec{Q}$ in two different ways. Triangle OAB shows $\vec{P} + \vec{Q} = \vec{R} = \vec{OB}$, while triangle OCB shows $\vec{Q} + \vec{P} = \vec{R} = \vec{OB}$.

$$\therefore \vec{P} + \vec{Q} = \vec{Q} + \vec{P}$$

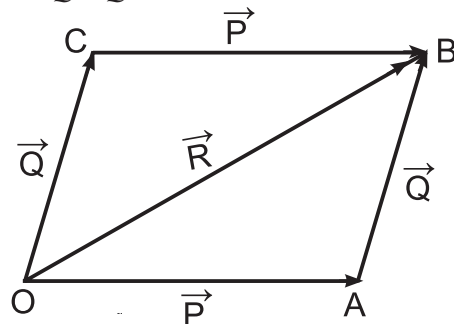


Fig. 2.6 (a): Commutative law.

(b) Vector addition is associative

If $\vec{A}$, $\vec{B}$ and $\vec{C}$ are three vectors then

$$(\vec{A} + \vec{B}) + \vec{C} = \vec{A} + (\vec{B} + \vec{C})$$

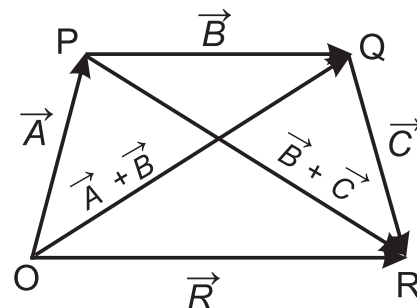


Fig. 2.6 (b): Associative law.

Figure 2.6 (b) shows addition of 3 vectors

$\vec{A}$, $\vec{B}$ and $\vec{C}$ in two different ways to give resultant $\vec{R}$.

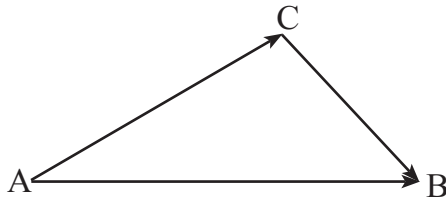
$$\vec{R} = (\vec{A} + \vec{B}) + \vec{C} \quad \text{--- from triangle OQR}$$

$$\vec{R} = \vec{A} + (\vec{B} + \vec{C}) \quad \text{--- from triangle OPR}$$

$$\text{i.e., } (\vec{A} + \vec{B}) + \vec{C} = \vec{A} + (\vec{B} + \vec{C}) \quad \text{--- (2.6)}$$

Thus the Associative law is proved.

Example 2.1: Express vector $\vec{AC}$ in terms of vectors $\vec{AB}$ and $\vec{CB}$ shown in the following figure.



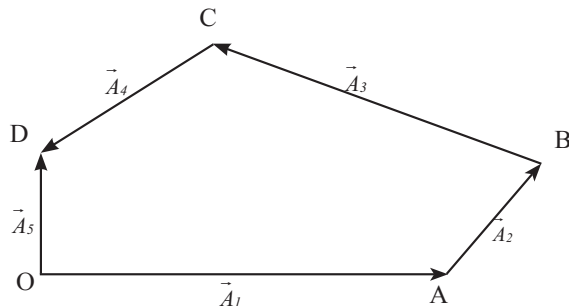
Solution: Using the triangle law of addition of vectors we can write

$$\vec{AC} + \vec{CB} = \vec{AB}$$

$$\therefore \vec{AC} = \vec{AB} - \vec{CB}$$

Example 2.2: From the following figure, determine the resultant of four forces

$$\vec{A}_1, \vec{A}_2, \vec{A}_3 \text{ and } \vec{A}_4$$



Solution: Join $\vec{OB}$ to complete ΔOAB as shown in (a)

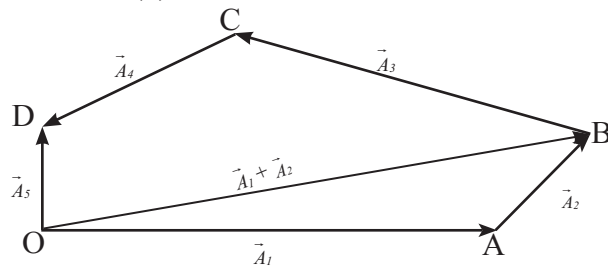


Fig. (a)

$$\text{Now, } \vec{OB} = \vec{OA} + \vec{AB} = \vec{A}_1 + \vec{A}_2$$

Join $\vec{OC}$ to complete triangle OBC as shown in (b).

$$\text{Now, } \vec{OC} = \vec{OB} + \vec{BC} = \vec{A}_1 + \vec{A}_2 + \vec{A}_3$$

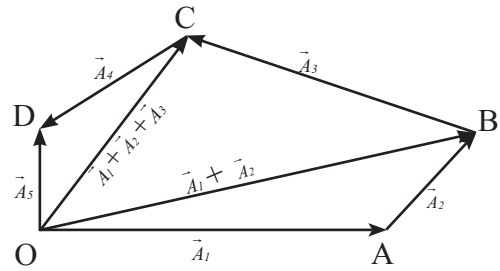


Fig. (b)

From triangle OCD,

$$\vec{OD} = \vec{A}_4 = \vec{OC} + \vec{CD} = \vec{A}_1 + \vec{A}_2 + \vec{A}_3 + \vec{A}_4$$

Thus $\vec{OD}$ is the resultant of the four vectors, $\vec{A}_1, \vec{A}_2, \vec{A}_3$ and $\vec{A}_4$, represented by

$$\vec{OA}, \vec{AB}, \vec{BC} \text{ and } \vec{CD}, \text{ respectively.}$$

2.3.4 Law of parallelogram of vectors:

Another geometrical method of adding two vectors is called parallelogram law of vector addition which is stated as follows:

If two vectors of the same type, originating from the same point (tails at the same point) are represented in magnitude and direction by two adjacent sides of a parallelogram, their resultant vector is given in magnitude and direction by the diagonal of the parallelogram starting from the same point as shown in Fig. 2.7.

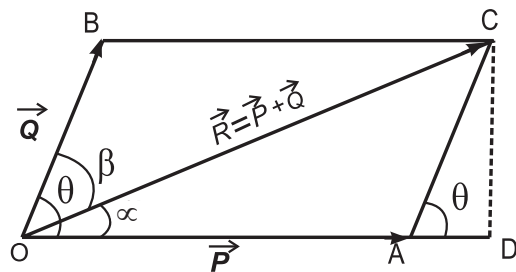


Fig 2.7: Parallelogram law of vector addition.

In Fig. 2.7, vector $\vec{OA} = \vec{P}$ and vector $\vec{OB} = \vec{Q}$, represent two vectors originating from point O, inclined to each other at an angle θ . If we complete the parallelogram, then according to this law, the diagonal $\vec{OC} = \vec{R}$ represents the resultant vector.

To find the magnitude of $\vec{R}$, drop a

perpendicular from C to reach OA (extended) at D. In right angled triangle ODC, by application by Pythagoras theorem,

$$\begin{aligned} OC^2 &= OD^2 + DC^2 \\ &= (OA + AD)^2 + DC^2 \\ OC^2 &= OA^2 + 2OA \cdot AD + AD^2 + DC^2 \end{aligned}$$

In the right angled triangle ADC, by application of Pythagoras theorem

$$\begin{aligned} AD^2 + DC^2 &= AC^2 \\ \therefore OC^2 &= OA^2 + 2OA \cdot AD + AC^2 \quad \text{--- (2.7)} \end{aligned}$$

Also,

$$\overrightarrow{OA} = \vec{P}, \overrightarrow{AC} = \overrightarrow{OB} = \vec{Q} \text{ and } \overrightarrow{OC} = \vec{R}$$

In ΔADC , $\cos \theta = AD/AC$

$$\therefore AD = AC \cos \theta = Q \cos \theta$$

Substituting in Eq. (2.7)

$$R^2 = P^2 + Q^2 + 2PQ \cos \theta.$$

$$\therefore R = \sqrt{P^2 + Q^2 + 2PQ \cos \theta} \quad \text{--- (2.8)}$$

Equation (2.8) gives us the magnitude of resultant vector $\vec{R}$.

To find the direction of the resultant vector $\vec{R}$, we will have to find the angle (α) made by $\vec{R}$ with $\vec{P}$.

$$\begin{aligned} \text{In } \Delta ODC, \tan \alpha &= \frac{DC}{OD} \\ &= \frac{DC}{OA + AD} \quad \text{--- (2.9)} \end{aligned}$$

$$\text{From the figure, } \sin \theta = \frac{DC}{AC}$$

$$\therefore DC = AC \sin \theta = Q \sin \theta$$

Also,

$$AD = AC \cos \theta = Q \cos \theta$$

and $OA = \vec{P}$,

Substituting in Eq. (2.9), we get

$$\begin{aligned} \tan \alpha &= \frac{Q \sin \theta}{P + Q \cos \theta} \\ \therefore \alpha &= \tan^{-1} \left(\frac{Q \sin \theta}{P + Q \cos \theta} \right) \quad \text{--- (2.10)} \end{aligned}$$

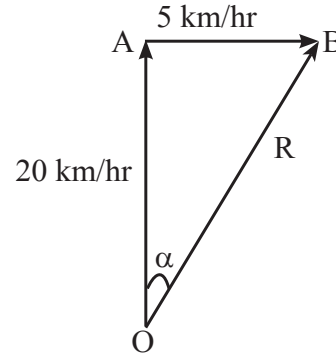
Equation (2.10) gives us the direction of resultant vector $\vec{R}$.

If β is the angle between $\vec{R}$ and $\vec{Q}$, it can be

$$\text{similarly derived that } \beta = \tan^{-1} \left(\frac{P \sin \theta}{Q + P \cos \theta} \right)$$

Example 2.3: Water is flowing in a stream with velocity 5 km/hr in an easterly direction relative to the shore. Speed of a boat is relative to still water is 20 km/hr. If the boat enters the stream heading North, with what velocity will the boat actually travel?

Solution: The resultant velocity $\vec{R}$ of the boat can be obtained by adding the two velocities using ΔOAB shown in the figure. Magnitude of the resultant velocity is calculated as follows:



$$\begin{aligned} R &= \sqrt{20^2 + 5^2} \\ &= \sqrt{425} = 20.61 \text{ km/hr} \end{aligned}$$

The direction of the resultant velocity is

$$\begin{aligned} &= \tan^{-1} \left(\frac{5}{20} \right) = \tan^{-1}(0.25) \\ &= 14^\circ 04' \end{aligned}$$

$\therefore$ The velocity of the boat is 20.61 km/hr in a direction $14^\circ 04'$ east of north.

2.4 Resolution of vectors:

A vector can be written as a sum of two or more vectors along certain fixed directions.

Thus a vector $\vec{V}$ can be written as

$$\vec{V} = V_1 \hat{\alpha} + V_2 \hat{\beta} + V_3 \hat{\gamma} \quad \text{--- (2.11)}$$

where $\hat{\alpha}, \hat{\beta}, \hat{\gamma}$ are unit vectors along chosen directions. V_1, V_2 and V_3 are known as components of $\vec{V}$ along the three directions $\hat{\alpha}, \hat{\beta}$ and $\hat{\gamma}$.

The process of splitting a given vector into its components is called resolution of the vector. The components can be found along

directions at any required angles, but if these components are found along the directions which are mutually perpendicular, they are called rectangular components.

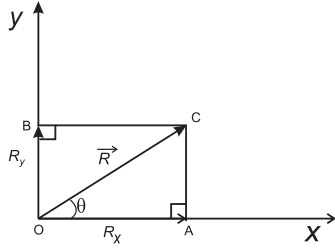


Fig. 2.8 : Resolution of a vector.

Let us see how to find rectangular components in two dimensions.

Consider a vector $\vec{R} = \vec{OC}$, originating from the origin of a rectangular co-ordinate system as shown in Fig. 2.8.

Drop perpendiculars from C that meet the x-axis at A and y-axis of at B.

$\vec{OA} = \vec{R}_x$ and $\vec{OB} = \vec{R}_y$; $\vec{R}_x$ and $\vec{R}_y$ being the components of $\vec{OC}$ along the x and y axes, respectively.

Then by the law of parallelogram of vectors,

$$\vec{R} = \vec{R}_x + \vec{R}_y \quad \text{--- (2.12)}$$

$$\vec{R} = R_x \hat{i} + R_y \hat{j} \quad \text{--- (2.13)}$$

where $\hat{i}$ and $\hat{j}$ are unit vectors along the x and y axes respectively, and R_x and R_y are the magnitudes of the two components of $\vec{R}$.

Let θ be the angle made by $\vec{R}$ with the x-axis, then

$$\begin{aligned} \cos \theta &= \frac{R_x}{R} \\ \therefore R_x &= R \cos \theta \end{aligned} \quad \text{--- (2.14)}$$

$$\begin{aligned} \sin \theta &= \frac{R_y}{R} \\ \therefore R_y &= R \sin \theta \end{aligned} \quad \text{--- (2.15)}$$

Squaring and adding Eqs. (2.14) and (2.15), we get

$$\begin{aligned} R^2 \cos^2 \theta + R^2 \sin^2 \theta &= R_x^2 + R_y^2 \\ \therefore R^2 &= R_x^2 + R_y^2 \\ \text{or, } R &= \sqrt{R_x^2 + R_y^2} \end{aligned} \quad \text{--- (2.16)}$$

Equation (2.16) gives the magnitude of $\vec{R}$. To find the direction of $\vec{R}$, from Fig. 2.8,

$$\begin{aligned} \tan \theta &= \frac{R_y}{R_x} \\ \therefore \theta &= \tan^{-1} \left(\frac{R_y}{R_x} \right) \end{aligned} \quad \text{--- (2.17)}$$

Similarly, if $\vec{R}_x$, $\vec{R}_y$ and $\vec{R}_z$ are the rectangular components of $\vec{R}$ along the x, y and z axes of the rectangular Cartesian co-ordinate system in three dimensions, then

$$\begin{aligned} \vec{R} &= \vec{R}_x + \vec{R}_y + \vec{R}_z = R_x \hat{i} + R_y \hat{j} + R_z \hat{k} \\ \text{or, } |\vec{R}| &= \sqrt{R_x^2 + R_y^2 + R_z^2} \end{aligned} \quad \text{---- (2.18)}$$

If two vectors are equal, it means that their corresponding components are also equal and vice versa.

$$\text{If } \vec{A} = \vec{B}$$

i.e., if $A_x \hat{i} + A_y \hat{j} + A_z \hat{k} = B_x \hat{i} + B_y \hat{j} + B_z \hat{k}$, then $A_x = B_x$, $A_y = B_y$ and $A_z = B_z$

Example 2.4: Find a unit vector in the direction of the vector $3\hat{i} + 4\hat{j}$

Solution:

$$\text{Let } \vec{V} = 3\hat{i} + 4\hat{j}$$

$$\text{Magnitude of } \vec{V} = |\vec{V}| = \sqrt{3^2 + 4^2} = \sqrt{25} = 5$$

$\vec{V} = \hat{\alpha} |\vec{V}|$, where $\hat{\alpha}$ is a unit vector along $\vec{V}$.

$$\hat{\alpha} = \frac{\vec{V}}{|\vec{V}|} = \frac{3}{5}\hat{i} + \frac{4}{5}\hat{j}$$

Example 2.5: Given $\vec{a} = \hat{i} + 2\hat{j}$ and $\vec{b} = 2\hat{i} + \hat{j}$, what are the magnitudes of the two vectors? Are these two vectors equal?

Solution:

$$|\vec{a}| = \sqrt{1^2 + 2^2} = \sqrt{5}$$

$$|\vec{b}| = \sqrt{2^2 + 1^2} = \sqrt{5}$$

The magnitudes of $\vec{a}$ and $\vec{b}$ are equal. However, their corresponding components are not equal i.e., $a_x \neq b_x$ and $a_y \neq b_y$. Hence, the two vectors are not equal.

2.5 Multiplication of Vectors:

We saw that we can add or subtract vectors of the same type to get resultant vectors of the same type. However, when we multiply vectors of the same or different types, we get a new physical quantity which may either be a scalar (scalar product) or a vector (vector product). Also note that the multiplication of a scalar with a scalar is always a scalar and the multiplication of scalar with a vector is always a vector. Let us now study the characteristics of a scalar product and vector product of two vectors.

2.5.1 Scalar Product (Dot Product):

The scalar product or dot product of two nonzero vectors $\vec{P}$ and $\vec{Q}$ is defined as the product of magnitudes of the two vectors and the cosine of the angle θ between the two vectors. The scalar product of $\vec{P}$ and $\vec{Q}$ is written as,

$$\vec{P} \cdot \vec{Q} = PQ \cos \theta, \quad \text{--- (2.19)}$$

where θ is the angle between $\vec{P}$ and $\vec{Q}$.

Characteristics of scalar product

(1) The scalar product of two vectors is equivalent to the product of magnitude of one vector with the magnitude of the component of the other vector in the direction of the first.

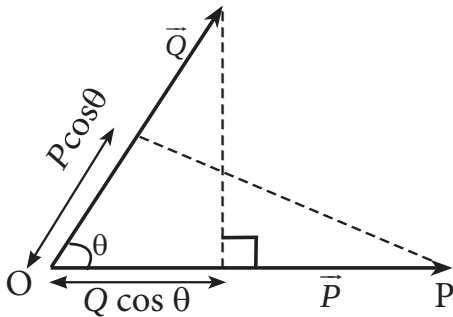


Fig. 2.9: Projection of vectors.

From Fig. 2.9,

$$\begin{aligned} \vec{P} \cdot \vec{Q} &= PQ \cos \theta \\ &= P (Q \cos \theta) \\ &= P (\text{component of } \vec{Q} \text{ in the direction of } \vec{P}) \end{aligned}$$

$$\begin{aligned} \text{Similarly } \vec{P} \cdot \vec{Q} &= Q (P \cos \theta) \\ &= Q (\text{component of } \vec{P} \text{ in the direction of } \vec{Q}) \end{aligned}$$

(2) Scalar product obeys the commutative law of vector multiplication.

$$\vec{P} \cdot \vec{Q} = P Q \cos \theta = Q P \cos \theta = \vec{Q} \cdot \vec{P}$$

(3) Scalar product obeys the distributive law of multiplication

$$\vec{P} \cdot (\vec{Q} + \vec{R}) = \vec{P} \cdot \vec{Q} + \vec{P} \cdot \vec{R}$$

(4) Special cases of scalar product $\vec{P} \cdot \vec{Q} = P Q \cos \theta$

(i) If $\theta = 0$, i.e., the two vectors $\vec{P}$ and $\vec{Q}$ are parallel to each other, then

$$\vec{P} \cdot \vec{Q} = P Q \cos \theta = P Q$$

$$\text{Thus, } \hat{i} \cdot \hat{i} = \hat{j} \cdot \hat{j} = \hat{k} \cdot \hat{k} = 1$$



Do you know ?

Scalar and vector products are very useful in physics. They make mathematical formulae and their derivation very elegant.

Figure below shows a toy car pulled through a displacement $\vec{S}$. The force $\vec{F}$ responsible for this is not in the direction of $\vec{S}$ but is at an angle θ to it. Component of displacement along the direction of force $\vec{F}$ is $S \cos \theta$. According to the definition, the work done by a force is the product of the force and the displacement in the direction of force. $\therefore W = FS \cos \theta$. According to the definition of scalar product,

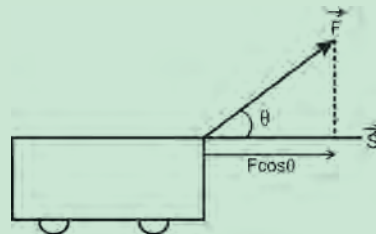
$$\vec{F} \cdot \vec{S} = F S \cos \theta$$

$$\therefore W = \vec{F} \cdot \vec{S}$$

$$\text{Also } W = F (S \cos \theta) = (F \cos \theta) S$$

Hence dot or scalar product is the product of magnitude of one of the vectors and component of the other vector in the direction of the first.

Power is the rate of doing work on a body by an external force $\vec{F}$ assumed to be constant in time. If $\vec{v}$ is the velocity of the body under the action of the force then power P is given by the scalar product of $\vec{F}$ and $\vec{v}$ i.e., $P = \vec{F} \cdot \vec{v}$.



(ii) If $\theta = 180^\circ$, i.e., the two vectors $\vec{P}$ and $\vec{Q}$ are anti-parallel, then

$$\vec{P} \cdot \vec{Q} = P Q \cos 180^\circ = -P Q$$

(iii) If $\theta = 90^\circ$, i.e., the two vectors are perpendicular to each other, then

$$\vec{P} \cdot \vec{Q} = P Q \cos 90^\circ = 0$$

$$\text{Thus, } \hat{i} \cdot \hat{j} = \hat{j} \cdot \hat{k} = \hat{k} \cdot \hat{i} = 0$$

(5) If $\vec{P} = \vec{Q}$ then $\vec{P} \cdot \vec{Q} = P^2 = Q^2$

(6) Scalar product of vectors expressed in terms of rectangular components :

$$\text{Let } \vec{P} = P_x \hat{i} + P_y \hat{j} + P_z \hat{k}$$

$$\text{and } \vec{Q} = Q_x \hat{i} + Q_y \hat{j} + Q_z \hat{k}$$

$$\text{Then } \vec{P} \cdot \vec{Q} = P_x Q_x + P_y Q_y + P_z Q_z$$

Proof :

$$\vec{P} \cdot \vec{Q} = (P_x \hat{i} + P_y \hat{j} + P_z \hat{k}) \cdot (Q_x \hat{i} + Q_y \hat{j} + Q_z \hat{k})$$

$$= P_x \hat{i} \cdot (Q_x \hat{i} + Q_y \hat{j} + Q_z \hat{k})$$

$$+ P_y \hat{j} \cdot (Q_x \hat{i} + Q_y \hat{j} + Q_z \hat{k})$$

$$+ P_z \hat{k} \cdot (Q_x \hat{i} + Q_y \hat{j} + Q_z \hat{k})$$

$$= (\hat{i} \cdot \hat{i}) P_x Q_x + (\hat{i} \cdot \hat{j}) P_x Q_y + (\hat{i} \cdot \hat{k}) P_x Q_z$$

$$+ (\hat{j} \cdot \hat{i}) P_y Q_x + (\hat{j} \cdot \hat{j}) P_y Q_y + (\hat{j} \cdot \hat{k}) P_y Q_z$$

$$+ (\hat{k} \cdot \hat{i}) P_z Q_x + (\hat{k} \cdot \hat{j}) P_z Q_y + (\hat{k} \cdot \hat{k}) P_z Q_z$$

$$\text{Since, } \hat{i} \cdot \hat{i} = \hat{j} \cdot \hat{j} = \hat{k} \cdot \hat{k} = 1$$

$$\text{and } \hat{i} \cdot \hat{j} = \hat{j} \cdot \hat{k} = \hat{k} \cdot \hat{i} = \hat{i} \cdot \hat{k} = \hat{j} \cdot \hat{i} = \hat{k} \cdot \hat{j} = 0$$

$$\therefore \vec{P} \cdot \vec{Q} = P_x Q_x + 0 + 0$$

$$+ 0 + P_y Q_y + 0$$

$$+ 0 + 0 + P_z Q_z$$

$$\therefore \vec{P} \cdot \vec{Q} = P_x Q_x + P_y Q_y + P_z Q_z$$

(7) If $\vec{a} \cdot \vec{b} = \vec{a} \cdot \vec{c}$, where $\vec{a} \neq 0$, it is not necessary that $\vec{b} = \vec{c}$. Using the distributive law, we can write $\vec{a} \cdot (\vec{b} - \vec{c}) = 0$. It implies that either $\vec{b} - \vec{c} = 0$ or $\vec{a}$ is perpendicular to $\vec{b} - \vec{c}$. It does not necessarily imply that $\vec{b} - \vec{c} = 0$

Example 2.6: Find the scalar product of the two vectors

$$\vec{v}_1 = \hat{i} + 2\hat{j} + 3\hat{k} \text{ and } \vec{v}_2 = 3\hat{i} + 4\hat{j} - 5\hat{k}$$

Solution:

$$\begin{aligned} \vec{v}_1 \cdot \vec{v}_2 &= (\hat{i} + 2\hat{j} + 3\hat{k}) \cdot (3\hat{i} + 4\hat{j} - 5\hat{k}) \\ &= 1 \times 3 + 2 \times 4 + 3 \times (-5) \\ &= -4 \end{aligned}$$

$$\text{as } \hat{i} \cdot \hat{i} = \hat{j} \cdot \hat{j} = \hat{k} \cdot \hat{k} = 1,$$

$$\text{and } \hat{i} \cdot \hat{j} = \hat{i} \cdot \hat{k} = \hat{j} \cdot \hat{k} = \hat{j} \cdot \hat{i} = \hat{k} \cdot \hat{i} = \hat{k} \cdot \hat{j} = 0$$

2.5.2 Vector Product (cross product):

The vector product or cross product of two vectors ($\vec{P}$ and $\vec{Q}$) is a vector whose magnitude is equal to the product of magnitudes of the two vectors and sine of the smaller angle (θ) between the two vectors. The direction of the product vector is given by $\hat{u}_r$ which is a unit vector perpendicular to the plane containing the two vectors and is given by the right hand screw rule. This is shown in Fig. 2.10 (a) and (b)

$$\text{a) } \vec{R} = \vec{P} \times \vec{Q} = PQ \sin \theta \hat{u}_r \quad \text{--- (2.18)}$$

$$\text{b) } \vec{S} = \vec{Q} \times \vec{P} = PQ \sin \theta \hat{u}_s \quad \text{--- (2.19)}$$

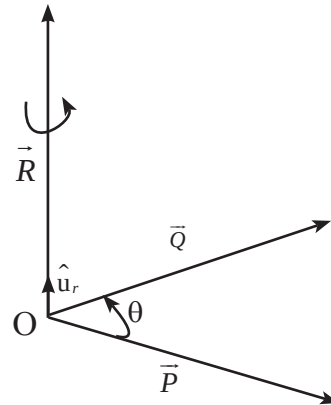


Fig. 2.10 (a): Vector product $\vec{R} = \vec{P} \times \vec{Q}$.

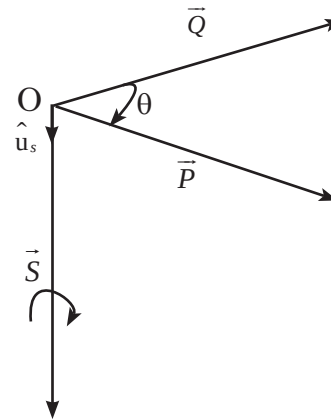


Fig. 2.10 (b): Vector product $\vec{S} = \vec{Q} \times \vec{P}$.

According to the right hand screw rule, if the screw is rotated in a direction from $\vec{P}$ to $\vec{Q}$ through the smaller angle, then the direction in which the tip of the screw advances is the direction of $\vec{R}$, perpendicular to the plane containing $\vec{P}$ and $\vec{Q}$. One example of vector or cross product is the force $\vec{F}$ experienced by a charge q moving with velocity $\vec{v}$ through a uniform magnetic field of magnetic induction $\vec{B}$. It is an empirical law (experimentally determined) given by $\vec{F} = q\vec{v} \times \vec{B}$.



Do you know ?

1. As linear displacement $\vec{x}$ is the distance travelled by a body along the line of travel, angular displacement $\vec{\theta}$ is the angle swept by a body about a given axis. The rate of change of angular displacement is the angular velocity denoted by $\vec{\omega}$. If a body is rotating about an axis, it possesses an angular velocity $\vec{\omega}$. If at a point at a distance $\vec{r}$ from the axis of rotation the body has linear velocity $\vec{v}$, then $\vec{v} = \vec{\omega} \times \vec{r}$.
2. An external force is needed to move a body from one point to other. Similarly to rotate a body about an axis passing through it, torque is required. Torque is a vector with its direction along the axis of rotation and magnitude describing the turning effect of force $\vec{F}$ acting on the body to rotate it about the given axis. Torque $\vec{\tau}$ is given as $\vec{\tau} = \vec{r} \times \vec{F}$, $\vec{r}$ being the perpendicular distance of a point on the body where the force is applied from the axis of rotation.

Characteristics of Vector Product:

- (1) Vector product does not obey commutative law of multiplication.

$$\vec{P} \times \vec{Q} \neq \vec{Q} \times \vec{P} \quad \text{--- (2.20)}$$

However, $|\vec{P} \times \vec{Q}| = |\vec{Q} \times \vec{P}|$ i.e., the magnitudes are the same but the directions are opposite to each other.

- (2) The vector product obeys the distributive law of multiplication.

$$\vec{A} \times (\vec{B} + \vec{C}) = \vec{A} \times \vec{B} + \vec{A} \times \vec{C} \quad \text{--- (2.21)}$$

(3) Special cases of cross product

$$|\vec{P} \times \vec{Q}| = PQ \sin \theta \quad \text{--- (2.22)}$$

- (i) If $\theta = 0^\circ$ i.e., if the two nonzero vectors are parallel to each other, their vector product is a zero vector $|\vec{P} \times \vec{Q}| = PQ \cdot 0 = 0$

- (ii) If $\theta = 180^\circ$ i.e., if the two nonzero vectors are anti-parallel, their vector product is a zero vector $|\vec{P} \times \vec{Q}| = PQ \sin 180^\circ = PQ \sin \pi = 0$

- (iii) If $\theta = 90^\circ$ i.e., if the two nonzero vectors are perpendicular to each other, the magnitude of their vector product is equal to the product of magnitudes of the two vectors.

$$|\vec{P} \times \vec{Q}| = PQ \sin 90^\circ = PQ$$

$$\text{Thus } \hat{i} \times \hat{j} = \hat{k}, \hat{j} \times \hat{k} = \hat{i} \text{ and } \hat{k} \times \hat{i} = \hat{j}$$

$$(4) \text{ If } \vec{P} = \vec{Q} \text{ then } |\vec{P} \times \vec{Q}| = |\vec{P} \times \vec{P}| = |\vec{Q} \times \vec{Q}| = 0$$

$$\text{Thus } \hat{i} \times \hat{i} = \hat{j} \times \hat{j} = \hat{k} \times \hat{k} = 0$$

$$(5) \text{ Let } \vec{P} = P_x \hat{i} + P_y \hat{j} + P_z \hat{k}$$

$$\text{and } \vec{Q} = Q_x \hat{i} + Q_y \hat{j} + Q_z \hat{k}$$

$$\begin{aligned} \vec{P} \times \vec{Q} &= (P_x \hat{i} + P_y \hat{j} + P_z \hat{k}) \times (Q_x \hat{i} + Q_y \hat{j} + Q_z \hat{k}) \\ &= P_x Q_x (\hat{i} \times \hat{i}) + P_x Q_y (\hat{i} \times \hat{j}) + P_x Q_z (\hat{i} \times \hat{k}) \\ &\quad + P_y Q_x (\hat{j} \times \hat{i}) + P_y Q_y (\hat{j} \times \hat{j}) + P_y Q_z (\hat{j} \times \hat{k}) \\ &\quad + P_z Q_x (\hat{k} \times \hat{i}) + P_z Q_y (\hat{k} \times \hat{j}) + P_z Q_z (\hat{k} \times \hat{k}) \end{aligned}$$

$$\text{Now } \hat{i} \times \hat{i} = \hat{j} \times \hat{j} = \hat{k} \times \hat{k} = 0, \text{ and}$$

$$\hat{i} \times \hat{k} = -\hat{j}, \hat{j} \times \hat{i} = -\hat{k}, \hat{k} \times \hat{j} = -\hat{i}$$

$$\hat{i} \times \hat{j} = \hat{k}, \hat{j} \times \hat{k} = \hat{i}, \hat{k} \times \hat{i} = \hat{j}$$

$$\therefore \vec{P} \times \vec{Q} = 0 + P_x Q_y \hat{k} - P_x Q_z \hat{j}$$

$$- P_y Q_x \hat{k} + 0 + P_y Q_z \hat{i}$$

$$+ P_z Q_x \hat{j} - P_z Q_y \hat{i} + 0$$

$$= (P_y Q_z - P_z Q_y) \hat{i}$$

$$+ (P_z Q_x - P_x Q_z) \hat{j}$$

$$+ (P_x Q_y - P_y Q_x) \hat{k}$$

This can be written in a determinant form as

$$\vec{P} \times \vec{Q} = \begin{vmatrix} \hat{i} & \hat{j} & \hat{k} \\ P_x & P_y & P_z \\ Q_x & Q_y & Q_z \end{vmatrix} \quad \text{--- (2.23)}$$

(6) The magnitude of cross product of two vectors is numerically equal to the area of a parallelogram whose adjacent sides represent the two vectors.

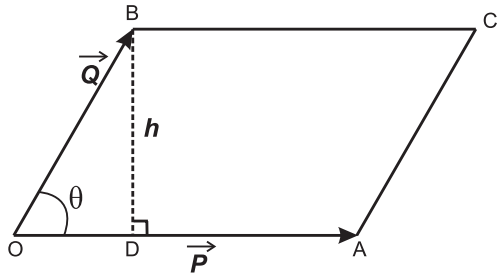


Fig 2.11: Area of parallelogram and vector product.

As shown in fig. 2.11,

$\vec{P} = \vec{OA}$, $\vec{Q} = \vec{OB}$, $\vec{P}$ and $\vec{Q}$ are inclined at an angle θ .

Perpendicular BD, of length h drawn on OA, gives the height of the parallelogram with OA as base.

Area of parallelogram

= base $\times$ height

$$= OA \times BD, \text{ as } \sin \theta = \frac{BD}{OB}$$

$$= PQ \sin \theta$$

$$= |\vec{P} \times \vec{Q}|$$

= magnitude of the vector product --- (2.24)

Example 2.7: The angular momentum $\vec{L} = \vec{r} \times \vec{p}$, where $\vec{r}$ is a position vector and $\vec{p}$ is linear momentum of a body.

If $\vec{r} = 4\hat{i} + 6\hat{j} - 3\hat{k}$ and $\vec{p} = 2\hat{i} + 4\hat{j} - 5\hat{k}$, find $\vec{L}$

Solution:

$$\vec{L} = \vec{r} \times \vec{p} = \begin{vmatrix} \hat{i} & \hat{j} & \hat{k} \\ 4 & 6 & -3 \\ 2 & 4 & -5 \end{vmatrix}$$

$$\therefore \vec{L} = (-30 + 12)\hat{i} + (-6 + 20)\hat{j} + (16 - 12)\hat{k}$$

$$= -18\hat{i} + 14\hat{j} + 4\hat{k}.$$

Example 2.8: If $\vec{A} = 5\hat{i} + 6\hat{j} + 4\hat{k}$ and

$\vec{B} = 2\hat{i} - 2\hat{j} + 3\hat{k}$, determine the angle between $\vec{A}$ and $\vec{B}$.

Solution: $\vec{A} \cdot \vec{B} = AB \cos \theta = A_x B_x + A_y B_y + A_z B_z$

$$\cos \theta = \frac{A_x B_x + A_y B_y + A_z B_z}{AB}$$

$$\cos \theta = \frac{A_x B_x + A_y B_y + A_z B_z}{\sqrt{A_x^2 + A_y^2 + A_z^2} \sqrt{B_x^2 + B_y^2 + B_z^2}}$$

$$\begin{aligned} \cos \theta &= \frac{(5)(2) + (6)(-2) + (4)(3)}{\sqrt{25 + 36 + 16} \sqrt{4 + 4 + 9}} \\ &= \frac{10}{\sqrt{77} \sqrt{17}} = 0.2764 \end{aligned}$$

$$\theta = \cos^{-1} 0.2765 = 73^\circ 58'$$

Example 2.9: Given $\vec{P} = 4\hat{i} - \hat{j} + 8\hat{k}$ and

$\vec{Q} = 2\hat{i} - m\hat{j} + 4\hat{k}$, find m if $\vec{P}$ and $\vec{Q}$ have the same direction.

Solution: Since $\vec{P}$ and $\vec{Q}$ have the same direction, their corresponding components must be in the same proportion, i.e.,

$$\frac{P_x}{Q_x} = \frac{P_y}{Q_y} = \frac{P_z}{Q_z}$$

$$\frac{4}{2} = \frac{-1}{-m} = \frac{8}{4}$$

$$\therefore m = \frac{1}{2}$$

2.6 Introduction to Calculus:

Calculus is the study of continuous (not discrete) changes in mathematical quantities. This branch of mathematics was first developed by G.W Leibnitz and Sir Issac Newton in the 17th century and is extensively used in several branches of science. You will study calculus in mathematics in XIIth standard. Here we will learn the basics of the two branches of calculus namely differential and integral calculus. These are necessary to understand the topics covered in this book.

2.6.1 Differential Calculus:

Let us consider a function $y = f(x)$. Here x is called an independent variable and $f(x)$ gives the value of y for different values of x and is the

dependent variable. For example x could be the position of a particle moving along x -axis and $y = f(x)$ could be its velocity at that position x . We can thus draw a graph of y against x as shown in Fig. 2.12 (a). Let A and B be two points on the curve giving values of y at $x = x_0$ and $x = x_0 + \Delta x$, where Δx is a small increment in x . The slope of the straight line joining A and B is given by $\tan \theta = \frac{\Delta y}{\Delta x}$.

If we make Δx smaller, the point B will come closer to A and if we keep making Δx smaller and smaller, we will ultimately reach a stage when B will coincide with A. This process is called taking the limit Δx going to zero and is written as $\lim_{\Delta x \rightarrow 0}$. In this limit the line AB extended on both sides to P and Q will become the tangent to the curve at A, i.e., at

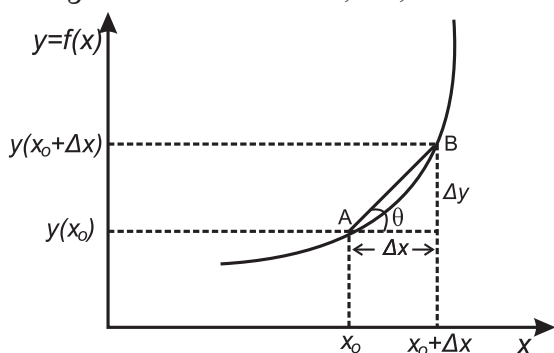


Fig. 2.12 (a): Average rate of change of y with respect to x .

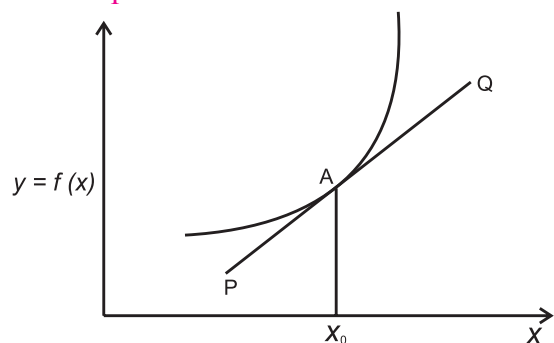


Fig. 2.12 (b): Rate of change of y with respect to x at x_0

$x = x_0$. In this limit both Δx and Δy will go to zero. However, when two quantities tend to zero, their ratio need not go to zero. In fact $\lim_{\Delta x \rightarrow 0} \left(\frac{\Delta y}{\Delta x} \right)$ becomes the slope of the tangent shown by PQ in Fig. 2.12 (b). This is written as dy/dx at $x = x_0$.

Thus,

$$\left. \frac{dy}{dx} \right|_{x_0} = \lim_{\Delta x \rightarrow 0} \frac{(y + \Delta y) - y}{\Delta x}$$

$$\left. \frac{df(x)}{dx} \right|_{x_0} = \lim_{\Delta x \rightarrow 0} \frac{f(x_0 + \Delta x) - f(x_0)}{\Delta x}$$

We can drop the subscript zero and write a general formula which will be valid for all values of x as

$$\frac{dy}{dx} = \lim_{\Delta x \rightarrow 0} \frac{f(x + \Delta x) - f(x)}{\Delta x} = \frac{df(x)}{dx} \quad \text{--- (2.25)}$$

In XIIth standard you will learn about the properties of derivatives and how to find derivatives of different functions. Here we will just list the properties as we will need them in later Chapter s. dy/dx is called the derivative of y with respect to x (which is the rate of change of y with respect to change in x) and the process of finding the derivative is called differentiation. Let $f_1(x)$ and $f_2(x)$ be two different functions of x and let s be a constant. Some of the properties of differentiation are

$$1. \frac{d(sf(x))}{dx} = s \frac{df(x)}{dx} \quad \text{--- (2.26)}$$

$$2. \frac{d}{dx} (f_1(x) + f_2(x)) = \frac{df_1(x)}{dx} + \frac{df_2(x)}{dx} \quad \text{--- (2.27)}$$

$$3. \frac{d}{dx} (f_1(x) \times f_2(x)) = f_1(x) \frac{df_2(x)}{dx} + f_2(x) \frac{df_1(x)}{dx} \quad \text{--- (2.28)}$$

$$4. \frac{d}{dx} \left(\frac{f_1(x)}{f_2(x)} \right) = \frac{1}{f_2(x)} \frac{df_1(x)}{dx} - \frac{f_1(x)}{f_2^2(x)} \frac{df_2(x)}{dx} \quad \text{--- (2.29)}$$

5. If x depends on time another variable t then,

$$\frac{df(x)}{dt} = \frac{df(x)}{dx} \frac{dx}{dt} \quad \text{--- (2.30)}$$

6.

$$\frac{d}{dx} f(g(x)) = f'(g(x)) \times g'(x)$$

$$\text{where } f'(g(x)) = \frac{df}{dg}$$

$$\text{or } \frac{dy}{dx} = \frac{dy}{dv} \frac{dv}{dx}$$

The derivatives of some simple functions of x are given below.

$$1. \frac{d}{dx}(x^n) = n x^{n-1} \quad \text{--- (2.31)}$$

$$2. \frac{d(e^x)}{dx} = e^x \text{ and } \frac{d(e^{ax})}{dx} = ae^{ax} \quad \text{--- (2.32)}$$

$$3. \frac{d}{dx}(\ln x) = \frac{1}{x} \quad \text{--- (2.33)}$$

$$4. \frac{d}{dx}(\sin x) = \cos x \quad \text{--- (2.34)}$$

$$5. \frac{d}{dx}(\cos x) = -\sin x \quad \text{--- (2.35)}$$

$$6. \frac{d}{dx}(\tan x) = \sec^2 x \quad \text{--- (2.36)}$$

$$7. \frac{d}{dx}(\cot x) = -\operatorname{cosec}^2 x \quad \text{--- (2.37)}$$

$$8. \frac{d}{dx}(\sec x) = \tan x \sec x \quad \text{--- (2.38)}$$

$$9. \frac{d}{dx}(\operatorname{cosec} x) = -\operatorname{cosec} x \cot x \quad \text{--- (2.39)}$$

Example 2.10: Find the derivatives of the functions.

$$(a) f(x) = x^8 \quad (b) f(x) = x^3 + \sin x$$

$$(c) f(x) = x^3 \sin x$$

Solution :

$$(a) \text{ Using } \frac{dx^n}{dx} = nx^{n-1},$$

$$\frac{d(x^8)}{dx} = 8x^7$$

(b) Using

$$\frac{d}{dx}(f_1(x) + f_2(x)) = \frac{df_1(x)}{dx} + \frac{df_2(x)}{dx} \text{ and}$$

$$\frac{d(\sin x)}{dx} = \cos x$$

$$\begin{aligned} \frac{d}{dx}(x^3 + \sin x) &= \frac{d(x^3)}{dx} + \frac{d(\sin x)}{dx} \\ &= 3x^2 + \cos x \end{aligned}$$

(c) Using

$$\frac{d}{dx}(f_1(x) f_2(x)) = f_1(x) \frac{df_2(x)}{dx} + \frac{df_1(x)}{dx} f_2(x)$$

$$\text{and } \frac{d(\sin x)}{dx} = \cos x$$

$$\begin{aligned} \frac{d}{dx}(x^3 \sin x) &= x^3 \frac{d(\sin x)}{dx} + \frac{d(x^3)}{dx} \sin x \\ &= x^3 \cos x + 3x^2 \sin x \end{aligned}$$

2.6.2 Integral calculus

Integral calculus is the branch of mathematics dealing with properties of integrals and their applications. Physical interpretation of integral of a function $f(x)$, i.e., $\int f(x)dx$ is the area under the curve $f(x)$ versus x . It is the reverse process of differentiation as we will see below.

We know how to find the area of a rectangle, triangle etc. In Fig. 2.13(a) we have shown y which is a function of x , A and B being two points on it.

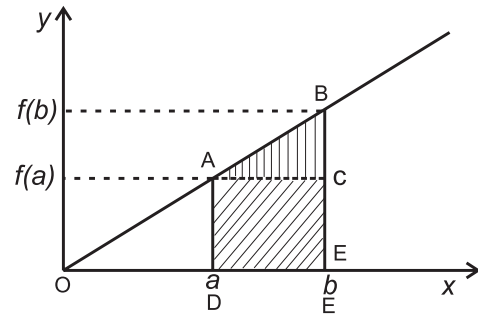


Fig. 2.13 (a): Area under a straight line.

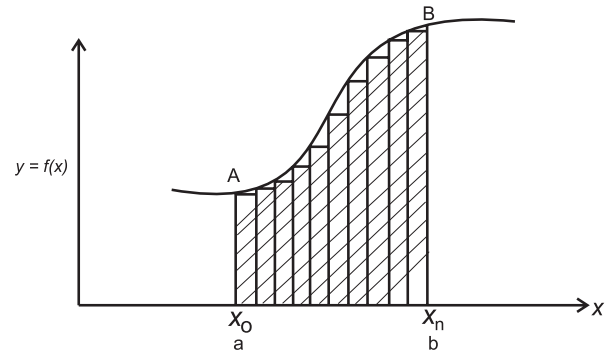


Fig. 2.13 (b): Area under a curve.

The area under the curve (straight line) from $x = a$ to $x = b$ is shown by shaded area. This can be obtained as sum of the area of the rectangle $ADEC = f(a)(b-a)$ and the area of the triangle $ABC = \frac{1}{2}(b-a)(f(b)-f(a))$

Figure 2.13(b) shows another function of x . We do not have a simple formula to calculate the area under this curve. For this calculation, we use a simple trick. We divide the area into a large number of vertical strips as shown in the figure. We assume thickness (width) of each strip to be so small that it can be assumed to be a rectangle as shown in the figure and add the areas of these rectangles. Thus the area under the curve is given by

Area under the curve

$$= \sum_{i=1}^n \Delta A_i = \sum_{i=1}^n (x_i - x_{i-1}) f(x_i)$$

where n is the number of strips and ΔA_i is the area of the i^{th} strip.

As the strips are not really rectangles, the area calculated above is not exactly equal to the area under the curve. However as we increase n , the sum of areas of rectangles gets closer to the actual area under the curve and becomes equal to it in the limit $n \rightarrow \infty$. Thus we can write,

Area under the curve

$$= \lim_{n \rightarrow \infty} \sum_{i=1}^n (x_i - x_{i-1}) f(x_i) \quad \text{--- (2.40)}$$

Integration helps us in getting exact area if the change is really continuous, i.e., n is really

infinite. It is represented as $\int_{x=a}^{x=b} f(x) dx$ and is

called the definite integral of $f(x)$ from $x = a$ to $x = b$.

$$\text{Thus, } \int_{x=a}^{x=b} f(x) dx = \lim_{n \rightarrow \infty} \sum_{i=1}^n (x_i - x_{i-1}) f(x_i) \quad \text{--- (2.41)}$$

The process of obtaining the integral is called integration. We can also write

$$F(x) = \int f(x) dx \quad \text{--- (2.42)}$$

$F(x)$ is called the indefinite (without any limits on x) integral of $f(x)$. Differentiation is the reverse process to that of integration. Therefore,

$$f(x) = \frac{d}{dx} (F(x)) \quad \text{--- (2.43)}$$

$$\therefore F(x) \Big|_a^b = F(b) - F(a) = \int_a^b f(x) dx \quad \text{--- (2.44)}$$

Properties of integration

$$1. \int (f_1(x) + f_2(x)) dx = \int f_1(x) dx + \int f_2(x) dx \quad \text{--- (2.45)}$$

$$2. \int K f(x) dx = K \int f(x) dx \text{ for } K = \text{constant} \quad \text{--- (2.46)}$$

Indefinite integrals of some basic functions are given below. Their definite integrals can be obtained by using the Eq. (2.44)

$$1. \int x^n dx = \frac{x^{n+1}}{n+1} \quad \text{--- (2.47)}$$

$$2. \int \frac{1}{x} dx = \ln x \quad \text{--- (2.48)}$$

$$3. \int \sin x dx = -\cos x \quad \text{--- (2.49)}$$

$$4. \int \cos x dx = \sin x \quad \text{--- (2.50)}$$

$$5. \int e^x dx = e^x \quad \text{--- (2.51)}$$

Example 2.11: Evaluate the following integrals:

$$(a) \int x^8 dx$$

$$(b) \int_2^5 x^2 dx$$

$$(c) \int (x + \sin x) dx$$

Solution: (a) Using formula

$$\int x^n dx = \frac{x^{n+1}}{n+1}, \quad \int x^8 dx = \frac{x^9}{9}$$

(b) Using Eq. (2.44),

$$\int_2^5 x^2 dx = \frac{x^3}{3} \Big|_2^5 = \frac{5^3}{3} - \frac{2^3}{3} = \frac{125 - 8}{3} = \frac{117}{3}$$

(c) Using Eq. (2.45),

$$\int (f_1(x) + f_2(x)) dx = \int f_1(x) dx + \int f_2(x) dx$$

and $\int \sin x dx = -\cos x$, we get $\int (x + \sin x) dx$

$$\int x dx + \int \sin x dx = \frac{x^2}{2} - \cos x$$



Internet my friend

1. hyperphysics.phy-astr.gsu.edu/hbase/vect.html#veccon
2. hyperphysics.phy-astr.gsu.edu/hbase/hframe.html



Exercises

1. Choose the correct option.

- The resultant of two forces 10 N and 15 N acting along +x and -x-axes respectively, is
(A) 25 N along +x-axis
(B) 25 N along -x-axis
(C) 5 N along +x-axis
(D) 5 N along -x-axis
- For two vectors to be equal, they should have the
(A) same magnitude
(B) same direction
(C) same magnitude and direction
(D) same magnitude but opposite direction
- The magnitude of scalar product of two unit vectors perpendicular to each other is
(A) zero (B) 1
(C) -1 (D) 2
- The magnitude of vector product of two unit vectors making an angle of 60° with each other is
(A) 1 (B) 2
(C) $3/2$ (D) $\sqrt{3}/2$
- If $\vec{A}, \vec{B}$ and $\vec{C}$ are three vectors, then which of the following is not correct?
(A) $\vec{A} \cdot (\vec{B} + \vec{C}) = \vec{A} \cdot \vec{B} + \vec{A} \cdot \vec{C}$
(B) $\vec{A} \cdot \vec{B} = \vec{B} \cdot \vec{A}$
(C) $\vec{A} \times \vec{B} = \vec{B} \times \vec{A}$
(D) $\vec{A} \times (\vec{B} + \vec{C}) = \vec{A} \times \vec{B} + \vec{B} \times \vec{C}$

2. Answer the following questions.

- Show that $\vec{a} = \frac{\hat{i} - \hat{j}}{\sqrt{2}}$ is a unit vector.
- If $\vec{v}_1 = 3\hat{i} + 4\hat{j} + \hat{k}$ and $\vec{v}_2 = \hat{i} - \hat{j} - \hat{k}$, determine the magnitude of $\vec{v}_1 + \vec{v}_2$.
[Ans: 5]
- For $\vec{v}_1 = 2\hat{i} - 3\hat{j}$ and $\vec{v}_2 = -6\hat{i} + 5\hat{j}$, determine the magnitude and direction of $\vec{v}_1 + \vec{v}_2$.
[Ans: $2\sqrt{5}$, $\theta = \tan^{-1}\left(-\frac{1}{2}\right)$ with x-axis]

- Find a vector which is parallel to $\vec{v} = \hat{i} - 2\hat{j}$ and has a magnitude 10.
[Ans: $\frac{10}{\sqrt{5}}\hat{i} - \frac{20}{\sqrt{5}}\hat{j}$]

- Show that vectors $\vec{a} = 2\hat{i} + 5\hat{j} - 6\hat{k}$ and $\vec{b} = \hat{i} + \frac{5}{2}\hat{j} - 3\hat{k}$ are parallel.

3. Solve the following problems.

- Determine $\vec{a} \times \vec{b}$, given $\vec{a} = 2\hat{i} + 3\hat{j}$ and $\vec{b} = 3\hat{i} + 5\hat{j}$.
[Ans: $\hat{k}$]
- Show that vectors $\vec{a} = 2\hat{i} + 3\hat{j} + 6\hat{k}$, $\vec{b} = 3\hat{i} - 6\hat{j} + 2\hat{k}$ and $\vec{c} = 6\hat{i} + 2\hat{j} - 3\hat{k}$ are mutually perpendicular.
- Determine the vector product of $\vec{v}_1 = 2\hat{i} + 3\hat{j} - \hat{k}$ and $\vec{v}_2 = \hat{i} + 2\hat{j} - 3\hat{k}$,
[Ans: $-7\hat{i} + 5\hat{j} + \hat{k}$]
- Given $\vec{v}_1 = 5\hat{i} + 2\hat{j}$ and $\vec{v}_2 = a\hat{i} - 6\hat{j}$ are perpendicular to each other, determine the value of a.
[Ans: $\frac{12}{5}$]
- Obtain derivatives of the following functions:
(i) $x \sin x$ (ii) $x^4 + \cos x$
(iii) $x/\sin x$

- Using the rule for differentiation for quotient of two functions, prove that
 $\frac{d}{dx} \left(\frac{\sin x}{\cos x} \right) = \sec^2 x$
- Evaluate the following integral:
(i) $\int_0^{\pi/2} \sin x \, dx$ (ii) $\int_1^5 x \, dx$
[Ans: (i) 1, (ii) 12]



Can you recall?

1. What is meant by motion?
2. What is rectilinear motion?
3. What is the difference between displacement and distance travelled?
4. What is the difference between uniform and nonuniform motion?

3.1 Introduction:

We see objects moving all around us. Motion is a change in the position of an object with time. We have come across the motion of a toy car when pushed along some particular direction, the motion of a cricket ball hit by a batsman for a sixer and the motion of an aeroplane from one place to another. The motion of objects can be divided in three categories: (1) motion along a straight line, i.e., rectilinear motion, (2) motion in two dimensions, i.e., motion in a plane and, (3) motion in three dimensions, i.e., motion in space. The above cited examples correspond to three types of motions, respectively. You have studied rectilinear motion in earlier standards. In rectilinear motion the force acting on the object and the velocity of the object both are along one and the same line. The distances are measured along the line only and we can indicate distances along the +ve and -ve axes as being positive and negative, respectively. The study of the motion of an object in a plane or in space becomes much easier and the corresponding equations become more elegant if we use vector quantities. In this Chapter we will first recall basic facts about rectilinear motion. We will use vector notation for this study as it will be useful later when we will study the motion in two dimensions. We will then study the motion in two dimensions which will be restricted to projectile motion only. Circular motion, i.e., the motion of an object around a circular path will be introduced here and will be studied in detail in the next standard.

3.2 Rectilinear Motion:

Consider an object moving along a straight line. Let us assume this line to be along the x -axis. Let $\vec{x}_1$ and $\vec{x}_2$ be the position vectors of the body at times t_1 and t_2 during its motion.

The following quantities can be defined for the motion.

1. **Displacement:** The displacement of the object between t_1 and t_2 is the difference between the position vectors of the object at the two instances. Thus, the displacement is given by

$$\vec{s} = \Delta \vec{x} = \vec{x}_2 - \vec{x}_1 \quad \text{--- (3.1)}$$

Its direction is along the line of motion of the object. Its dimensions are that of length. For example, if an object has travelled through 1 m from time t_1 to t_2 along the +ve x -direction, the magnitude of its displacement is 1 m and its direction is along the +ve x -axis. On the other hand, if the object travelled along the +ve y direction through the same distance in the same time, the magnitude of its displacement is the same as before, i.e., 1 m but the direction of the displacement is along the +ve y -axis.

2. **Path length:** This is the actual distance travelled by the object during its motion. It is a scalar quantity and its dimensions are also that of length. If an object travels along the x -axis from $x = 2$ m to $x = 5$ m then the distance travelled is 3 m. In this case the displacement is also 3 m and its direction is along the +ve x -axis. However, if the object now comes back to $x = 4$, then the distance through which the object has moved increases to $3 + 1 = 4$ m. Its initial position was $x = 2$ m and the final position is now $x = 4$ m and thus, its displacement is $\Delta x = 4 - 2 = 2$ m, i.e., the magnitude of the displacement is 2 m and its direction is along the +ve x -axis. If the object now moves to $x = 1$, then the distance travelled, i.e., the path length increases to $4 + 3 =$

7 m while the magnitude of displacement becomes $2 - 1 = 1$ m and its direction is along the negative x -axis.

- 3. Average velocity:** This is defined as the displacement of the object during the time interval over which average velocity is being calculated, divided by that time interval. As displacement is a vector quantity, the velocity is also a vector quantity. Its dimensions are $[L^1 M^0 T^{-1}]$.

If the position vectors of the object are $\vec{x}_1$ and $\vec{x}_2$ at times t_1 and t_2 respectively, then the average velocity is given by

$$\vec{v}_{av} = \frac{\vec{x}_2 - \vec{x}_1}{(t_2 - t_1)} \quad \text{--- (3.2)}$$

For example, if the positions of an object are $x = +2$ m and $x = +4$ m at times $t = 0$ and $t = 1$ minute respectively, the magnitude of its average velocity during that time is $v_{av} = (4 - 2)/(1 - 0) = 2$ m per minute and its direction will be along the +ve x -axis, and we write $\vec{v}_{av} = 2\hat{i}$ m/min where $\hat{i}$ is a unit vector along x -axis.

- 4. Average speed:** This is defined as the total path length travelled during the time interval over which average speed is being calculated, divided by that time interval.

Average speed = v_{av} = path length/time interval. It is a scalar quantity and has the same dimensions as that of velocity, i.e., $[L^1 M^0 T^{-1}]$.

If the rectilinear motion of the object is only in one direction along a line, then the magnitude of its displacement will be equal to the distance travelled and so the magnitude of average velocity will be equal to the average speed. However if the object reverses its direction (the motion remaining along the same line) then the magnitude of displacement will be smaller than the path length and the average speed will be larger than the magnitude of average velocity.

- 5. Instantaneous velocity:** Instantaneous velocity of an object is its velocity at a

given instant of time. It is defined as the limiting value of the average velocity of the object over a small time interval (Δt) around t when the value of the time interval (Δt) goes to zero.

$$\vec{v} = \lim_{\Delta t \rightarrow 0} \left(\frac{\Delta \vec{x}}{\Delta t} \right) = \frac{d\vec{x}}{dt}, \quad \text{--- (3.3)}$$

$\frac{d\vec{x}}{dt}$ being the derivative of $\vec{x}$ with respect

to t (see Chapter 2).

- 6. Instantaneous speed:** Instantaneous speed is the speed of an object at a given instant of time t . It is the limiting value of the average speed of the object taken over a small time interval (Δt) around t when the time interval goes to zero. In such a limit, the path length will be equal to the magnitude of the displacement and so the instantaneous speed will always be equal to the magnitude of the instantaneous velocity of the object.

Always Remember:

For uniform rectilinear motion, i.e., for an object moving with constant velocity along a straight line

1. The average and instantaneous velocities are equal.
2. The average and instantaneous speeds are the same and are equal to the magnitude of the velocity.

For nonuniform rectilinear motion

1. The average and instantaneous velocities are different.
2. The average and instantaneous speeds are different.
3. The average speed will be different from the magnitude of average velocity.

Example 3.1: A person walks from point P to point Q along a straight road in 10 minutes, then turns back and returns to point R which is midway between P and Q after further 4 minutes. If PQ is 1 km, find the average speed

and velocity of the person in going from P to R.

Solution: The path length travelled by the person is 1.5 km while the displacement is the distance between R and P which is 0.5 km. The time taken for the motion is 14 min.

The average speed = $1.5 / 14 = 0.107 \text{ km/min} = 6.42 \text{ km/hr}$.

The magnitude of the average velocity = $0.5/14 = 0.0357 \text{ km/min} = 2.142 \text{ km/hr}$.

Graphical Study of Motion

We can study the motion of an object by using graphs showing its position as a function of time. Figure 3.1 shows the graphs of position as a function of time for five different types of motion of an object. Figure 3.1(a) shows an object at rest, for which the x - t graph is a horizontal straight line. Since the position is not changing, displacement of the object zero. Velocity is displacement (which is zero) divided by time interval or the derivative of displacement with respect to time. It can be obtained from the slope of the line plotted in the figure which is zero.

Figure 3.1(b) shows x - t graph for an object moving with constant velocity along the +ve x -axis. Since velocity is constant, displacement is proportional to elapsed time. The slope of the straight line is +ve, showing that the velocity is along the +ve x -axis. As the motion is uniform, the average velocity is same as the instantaneous velocity at all times. Also, the speed is equal to the magnitude of the velocity.

Figure 3.1(c) shows the x - t graph for a body moving with uniform velocity but along the -ve x -axis, the slope of the line being -ve. Figure 3.1(d) shows the x - t graph of an object having oscillatory motion with constant speed. The direction of velocity changes from +ve to -ve and vice versa over fixed intervals of time.

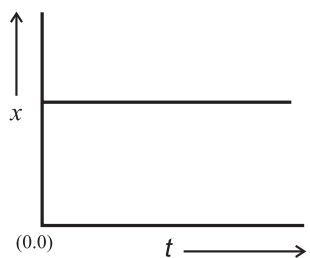


Fig 3.1 (a): Object at rest.

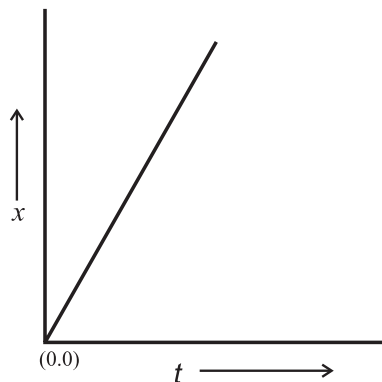


Fig 3.1 (b): Object with uniform velocity along +ve x -axis.

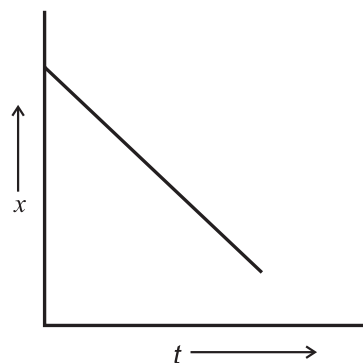


Fig 3.1 (c): Object with uniform velocity along -ve x -axis.

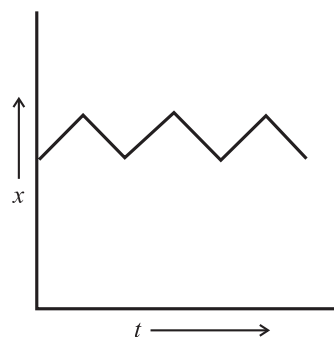


Fig 3.1 (d): Object performing oscillatory motion.

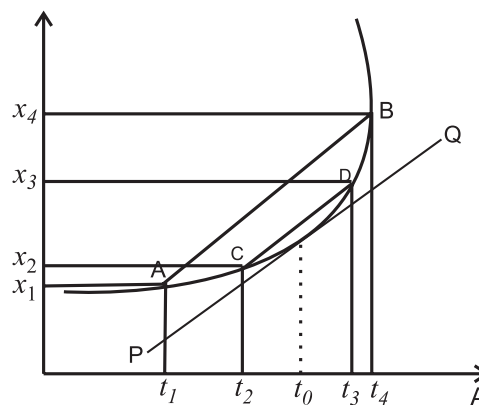


Fig.3.1 (e): Object in nonuniform motion.

Figure 3.1(e) shows the motion of an

object with nonuniform velocity. Its velocity changes with time and, therefore, the average and instantaneous velocities are different. Figure shows the average velocity over time interval from t_1 to t_2 around time t_0 , which can be seen from Eq. (3.2) to be the slope of line AB. For a smaller time interval from t_2 to t_3 , the average velocity is the slope of the line CD. If we keep reducing the time interval around t_0 , we will ultimately come to a limit, when the time interval will go to zero and lines AB, CD... will go over to the tangent to the curve at t_0 . The instantaneous velocity at t_0 will thus be equal to the slope of the tangent PQ at t_0 (see Eq. (3.3)).

7. Acceleration: Acceleration is defined as the rate of change of velocity with time. It is a vector quantity and its dimensions are $[L^1 M^0 T^{-2}]$. The average acceleration of an object having velocities $\vec{v}_1$ and $\vec{v}_2$ at times t_1 and t_2 is given by

$$\vec{a} = \frac{(\vec{v}_2 - \vec{v}_1)}{(t_2 - t_1)} \quad \text{--- (3.4)}$$

Instantaneous acceleration is the limiting value of the average acceleration when the time interval goes to zero. It is given by

$$\vec{a} = \lim_{\Delta t \rightarrow 0} \left(\frac{\Delta \vec{v}}{\Delta t} \right) = \frac{d\vec{v}}{dt} \quad \text{--- (3.5)}$$

The instantaneous acceleration at a given time is the slope of the tangent to the velocity versus time curve at that time. Figure 3.2 shows the velocity versus time ($v - t$) graphs for four different cases. Figure 3.2(a) represents the motion of an object with zero acceleration, i.e., constant velocity. The shaded area under the velocity-time graph over some time interval t_1 to t_2 , shown in Figs. 3.2(a) is equal to $v_0(t_2 - t_1)$ which is the magnitude of the displacement of the object from t_1 to t_2 . Figure 3.2(b) is the velocity-time graph for an object moving with constant +ve acceleration (magnitude of velocity uniformly increasing with time). Figure 3.2(c) shows similar motion but the object has -ve acceleration, i.e., the acceleration is opposite to the direction of velocity which, therefore, decreases uniformly with time. The area under both the curves between two instants of time is

the displacement of the object during that time interval (as shown below). Figure 3.2(d) shows the motion of an object having nonuniform acceleration. The average acceleration between t_1 and t_2 around t_0 and the instantaneous accelerations at t_0 for the object are shown by straight lines AB and CD respectively.

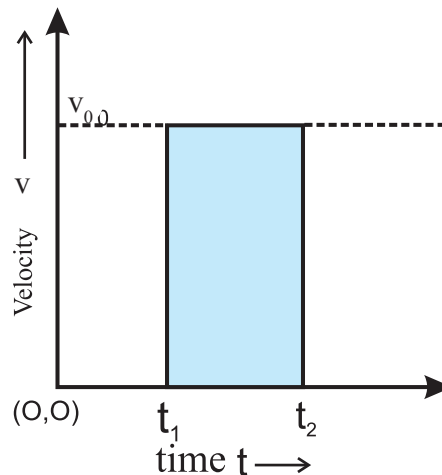


Fig 3.2 (a): Object moving with constant velocity.

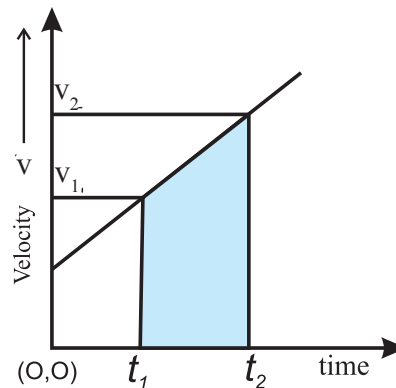


Fig 3.2 (b): Object moving with velocity (v) along +ve x -axis with uniform acceleration along the same direction.

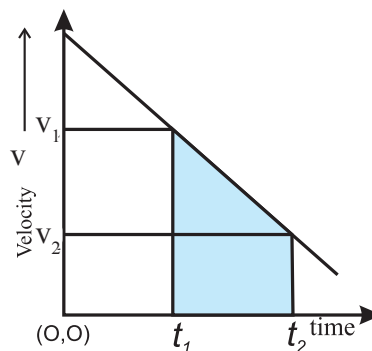


Fig 3.2 (c): Object moving with velocity (v) with negative uniform acceleration.

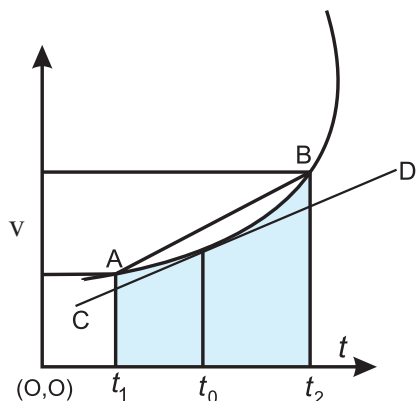


Fig. 3.2 (d): Object moving with nonuniform acceleration.

The area under the velocity-time curves in Figs. 3.2(a) to (d) can be written using the definition of integral given in Chapter 2 as

$$\begin{aligned} \text{Area} &= \int_{t_1}^{t_2} v dt = \int_{t_1}^{t_2} \frac{dx}{dt} dt = \int_{t_1}^{t_2} dx = x(t_2) - x(t_1) \quad \text{--- (3.6)} \\ &= \text{displacement of the object from } t_1 \text{ to } t_2. \end{aligned}$$

Always Remember:

For uniform acceleration, for a rectilinear motion:

1. Velocity-time graph is linear.
2. The area under the velocity-time graph between two instants of time t_1 and t_2 gives the displacement of the object during that time interval.
3. The slope of the velocity-time graph is the acceleration of the object

For nonuniform acceleration in a rectilinear motion:

1. Velocity-time graph is nonlinear.
2. The area under the velocity-time graph between two instants of time t_1 and t_2 gives the displacement of the object during that time interval.
3. The instantaneous acceleration of the object at a given time is equal to the slope of the tangent to the curve at that point.

While using the concept of area under the curve, the origin of the velocity axis (for v-t graph) must be zero.

Equations of Motion for Uniform Acceleration:

We can graphically derive Newton's equations of motion for an object moving with uniform acceleration. Consider an object having position $x = 0$ at $t = 0$. Let the velocity at $t = 0$ be u and at time t be v . The graphical representation of motion is shown in Fig. 3.3. The acceleration is given by the slope of the line AB. Thus,

$$\begin{aligned} \text{Acceleration, } a &= \frac{v - u}{t - 0} = \frac{v - u}{t} \\ \therefore v &= u + at \quad \text{--- (3.7)} \end{aligned}$$

This is the first equation of motion.

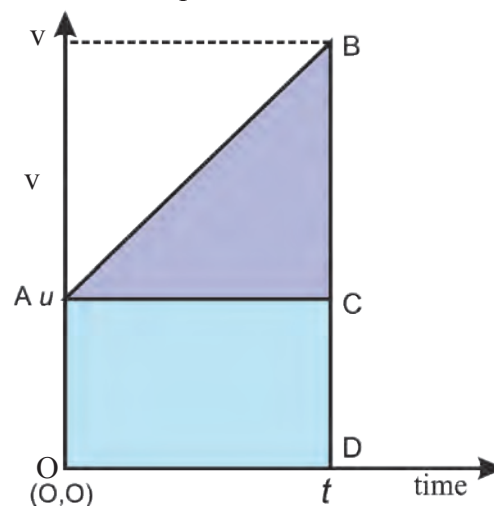


Fig.3.3: Derivation of equation of motion for motion with uniform acceleration.

As we know, the area under the curve in velocity-time graph is the displacement of the object. Thus displacement $s = \text{area of the quadrilateral OABD} = \text{area of triangle ABC} + \text{area of rectangle OACD}$.

$$= \frac{1}{2}(v - u)t + ut$$

$$\text{Using Eq. (3.7), } s = ut + \frac{1}{2}at^2 \quad \text{--- (3.8)}$$

This is the second equation of motion.

As the acceleration is constant, the velocity is increasing linearly with time and we can use average velocity v_{av} , to calculate the displacement using Eq. (3.7) as

$$s = v_{av}t = \left(\frac{v + u}{2} \right) t = \frac{(v + u)(v - u)}{2a}$$

$$\therefore s = (v^2 - u^2) / (2a)$$

$$\therefore v^2 - u^2 = 2as \quad \text{--- (3.9)}$$

This is the third equation of motion. Vector notation was not included here as the motion was rectilinear.

The most common example of uniform rectilinear motion with uniform acceleration of an object in day to day life is a freely falling body. When a body starts with zero velocity at a certain height from the ground and falls under the influence of the gravity of the Earth, it is said to be in free fall. The only other force that acts on it is that of the air resistance or friction. For displacements of a few metres, this force is too small and can be neglected. The acceleration of the body is the acceleration due to gravity which is along the vertical direction and can be assumed to be constant over distances which are small compared to the radius of the Earth. Thus the velocity and acceleration are both along the vertical direction and the motion is a uniform rectilinear motion with uniform acceleration.



Do you know ?

The distance travelled by an object starting from rest and having a uniform acceleration in successive seconds are in the ratio 1:3:5:7... Consider a freely falling object. Let us calculate the distances travelled by it in equal intervals of time t_0 (say). This can be done using the second equation of motion $s = ut_0 + (1/2)gt_0^2$. The initial velocity is zero. Therefore, the distance travelled in the first t_0 interval = $(1/2)gt_0^2$. For simplification let us write $(1/2)g = A$. Then the distance travelled in the first t_0 time interval = $d_1 = At_0^2$. In the time interval $2t_0$, the distance travelled = $A(2t_0)^2$. Hence, the distance travelled in the second t_0 interval is $d_2 = A(4t_0^2 - t_0^2) = 3At_0^2 = 3d_1$. The distance travelled in time interval $3t_0 = A(3t_0)^2$. Thus, the distance travelled in the 3rd t_0 interval = $d_3 = A(9t_0^2 - 4t_0^2) = 5At_0^2 = 5d_1$. Continuing, one can see that the distances $d_1, d_2, d_3 \dots$ are in the ratio 1:3:5:7... This is true for any rectilinear motion, starting from rest, with positive uniform acceleration.

Example 3.2: A stone is thrown vertically upwards from the ground with a velocity 15 m/s. At the same instant a ball is dropped from a point directly above the stone from a height of 30 m. At what height from the ground will the stone and the ball meet and after how much time? (Use $g = 10 \text{ m/s}^2$ for ease of calculation).

Solution: Let us assume that the stone and the ball meet after time t_0 . The distances (not displacements) travelled by the stone and the ball in that time can be obtained from Eq. (3.8) as

$$s_{\text{stone}} = 15t_0 - \frac{1}{2}gt_0^2$$

$$s_{\text{ball}} = \frac{1}{2}gt_0^2$$

When they meet, $s_{\text{stone}} + s_{\text{ball}} = 30$

$$15t_0 - \frac{1}{2}gt_0^2 + \frac{1}{2}gt_0^2 = 30$$

$$t_0 = 30/15 = 2 \text{ s}$$

$$\therefore s_{\text{stone}} = 15(2) - \frac{1}{2}(10)(2)^2 = 30 - 20 = 10 \text{ m}$$

Thus the stone and the ball meet at a height of 10 m.

8. Relative Velocity: You must have often experienced relative motion. The most striking example is when you are going in a train and another train travelling in the same direction along parallel tracks, overtakes you. If you look at that train, it actually seems to be moving much slower than what your train seemed to move and yet it is overtaking you. On the other hand if your train overtakes another train, travelling on a parallel track in the same direction, and you look at that train, you feel that your train has suddenly slowed down. Why does this happen? This is because when you look at the neighbouring train, you are actually experiencing relative motion, i.e., your motion with respect to the other train or the motion of the other train with respect to you. Thus, in the first case as the other train overtakes you what you perceive is the velocity of the train with respect to you, i.e., the difference in the velocities of the two trains which most often is much smaller than the velocity of your train. In the second case, you are moving faster but when you look at that train you only feel your velocity

relative to it and, therefore, your velocity appears to be lower than its actual value. We can define relative velocity of object A with respect to object B as the difference between their velocities, i.e.,

$$v_{AB} = v_A - v_B \quad \text{--- (3.10)}$$

Similarly, the velocity of B with respect to A is given by

$$v_{BA} = v_B - v_A \quad \text{--- (3.11)}$$

We assume that at time $t = 0$, A and B were at the same point $x = 0$. As they are travelling with different velocities, the distance between them will go on increasing with time in direct proportion to the difference in their velocities, i.e., the relative velocity between them.

Example 3.3: An aeroplane A, is travelling in a straight line with a velocity of 300 km/hr with respect to Earth. Another aeroplane B, is travelling in the opposite direction with a velocity of 350 km/hr with respect to Earth. What is the relative velocity of A with respect to B? What should be the velocity of a third aeroplane C moving parallel to A, relative to the Earth if it has a relative velocity of 100 km/hr with respect to A?

Solution: Let v_A , v_B and v_C be the velocities of the three planes relative to the Earth. Relative velocity of A with respect to B = $v_{AB} = v_A - v_B = 300 - (-350) = 650$ km/hr

Relative velocity of C with respect to A = $v_{CA} = v_C - v_A = 100$ km/hr.

Thus, $v_C = v_{CA} + v_A = 400$ km/hr

3.3 Motion in Two Dimensions-Motion in a Plane:

So far we were considering rectilinear motion of an object. The direction of motion of the object was always along one straight line. Now we will consider the motion of an object in two dimensions, i.e., along a plane. Here, the direction of the force acting on an object will not be in the same line as its initial velocity. Thus, the velocity and acceleration will have different directions. For this reason we have to use vector equations. The definitions of various terms given in section 3.2 will remain valid except that the magnitude of the average velocity and

the value of average speed will be different as the magnitude of the displacement need not be equal to the path length. For example, if a particle travels along a circle and comes back to its original position, its displacement will be zero but the path length will be equal to the circumference of the circle.

3.3.1 Average and Instantaneous Velocities:

For studying the motion of an object in two dimensions, for simplicity, we will take the plane to be the x - y plane. To describe the position of an object in this plane we will have to specify, both its x and y coordinates. The definitions of displacement, average and instantaneous velocities, average and instantaneous speeds and acceleration will be the same as those for rectilinear motion except that each of these quantities will now have components along the x and y directions. Let us assume the object to be at point P at time t_1 as shown in Fig. 3.4 (a).

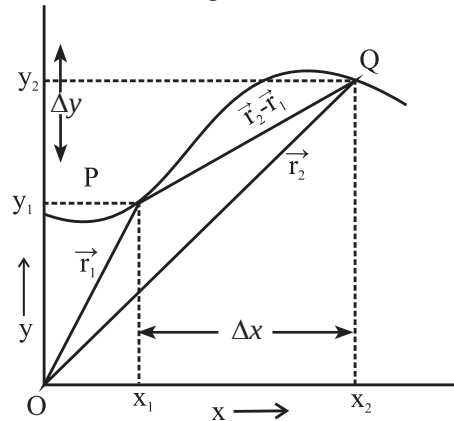


Fig. 3.4 (a) Motion in two dimensions

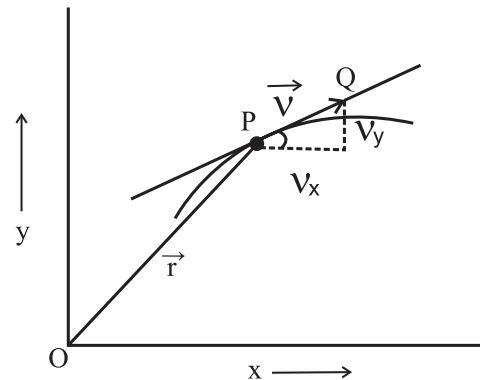


Fig. 3.4 (b) Instantaneous velocity

The position of the object will be described by its position vector $\vec{r}_1$. This can be written in terms of its components along the x and y axes as

$$\vec{r}_1 = x_1 \hat{i} + y_1 \hat{j} \quad \text{--- (3.12)}$$

At time t_2 , let the position of the object be Q and its position vector be $\vec{r}_2$

$$\vec{r}_2 = x_2 \hat{i} + y_2 \hat{j} \quad \text{--- (3.13)}$$

The displacement of the particle from t_1 to t_2 shown by PQ, i.e., in time $t = t_2 - t_1$ is given by

$$\Delta \vec{r} = \vec{r}_2 - \vec{r}_1 = (x_2 - x_1) \hat{i} + (y_2 - y_1) \hat{j} \quad \text{--- (3.14)}$$

We can write the average velocity of the object as

$$\vec{v}_{av} = \frac{\Delta \vec{r}}{\Delta t} = \left(\frac{x_2 - x_1}{t_2 - t_1} \right) \hat{i} + \left(\frac{y_2 - y_1}{t_2 - t_1} \right) \hat{j}$$

$$\vec{v}_{av} = (v_{av})_x \hat{i} + (v_{av})_y \hat{j} \quad \text{--- (3.15)}$$

where, $(v_{av})_x = (x_2 - x_1)/(t_2 - t_1)$ and

$$(v_{av})_y = (y_2 - y_1)/(t_2 - t_1) \quad \text{--- (3.16)}$$

Average velocity is a vector whose direction is along $\Delta \vec{r}$ (see Eq. (3.2)), i.e., along the direction of displacement. In terms of its components, the magnitude (v) and direction (the angle θ that the velocity vector makes with the x-axis) can be written as (see Chapter 2)

$$v_{av} = \sqrt{(v_{av})_x^2 + (v_{av})_y^2} \quad \text{and} \quad \tan \theta = (v_{av})_y / (v_{av})_x \quad \text{--- (3.17)}$$

Figure 3.4(b) shows the trajectory of an object moving in two dimensions. The instantaneous velocity of the object at point P along the trajectory is along the tangent to the curve at P. This is shown by the vector PQ. Its x and y components v_x and v_y are also shown in the figure.

The instantaneous velocity of the object can be written in terms of derivative as (see Eq. 3.3)

$$\vec{v} = \lim_{\Delta t \rightarrow 0} \left(\frac{\Delta \vec{r}}{\Delta t} \right) = \frac{d\vec{r}}{dt} = \left(\frac{dx}{dt} \right) \hat{i} + \left(\frac{dy}{dt} \right) \hat{j} \quad \text{--- (3.18)}$$

The magnitude and direction of the instantaneous velocity are given by

$$v = \sqrt{\left(\frac{dx}{dt} \right)^2 + \left(\frac{dy}{dt} \right)^2}, \quad \text{--- (3.19)}$$

$$\tan \theta = (dy/dt) / (dx/dt) = dy/dx \quad \text{--- (3.20)}$$

which is the slope of the tangent to the curve at the point at which we are calculating the instantaneous velocity.

3.3.2 Average and Instantaneous Acceleration:

Again, the definitions are the same as those for rectilinear motion. Thus, the average acceleration ($\vec{a}_{av}$) of a particle between times t_1 and t_2 can be written as

$$\vec{a}_{av} = \frac{\vec{v}_2 - \vec{v}_1}{t_2 - t_1} = \left(\frac{v_{2x} - v_{1x}}{t_2 - t_1} \right) \hat{i} + \left(\frac{v_{2y} - v_{1y}}{t_2 - t_1} \right) \hat{j} \quad \text{--- (3.21)}$$

where $\vec{v}_2$ and $\vec{v}_1$ are the velocities of the particle at times t_2 and t_1 respectively.

$$\vec{a}_{av} = (a_{av})_x \hat{i} + (a_{av})_y \hat{j} \quad \text{--- (3.22)},$$

$(a_{av})_x$ and $(a_{av})_y$ being the x and y components of the average acceleration.

The magnitude and direction of the acceleration are given by

$$a_{av} = \sqrt{(a_{av})_x^2 + (a_{av})_y^2} \quad \text{--- (3.23)}$$

and

$$\tan \theta = (a_{av})_y / (a_{av})_x \quad \text{--- (3.24)}$$

The instantaneous acceleration is given by (see Eq. (3.5))

$$\begin{aligned} \vec{a} &= \lim_{\Delta t \rightarrow 0} \left(\frac{\Delta \vec{v}}{\Delta t} \right) = \frac{d\vec{v}}{dt} = \left(\frac{dv_x}{dt} \right) \hat{i} + \left(\frac{dv_y}{dt} \right) \hat{j} \quad \text{--- (3.25)} \\ &= \frac{d}{dt} \left(\frac{dx}{dt} \right) \hat{i} + \frac{d}{dt} \left(\frac{dy}{dt} \right) \hat{j} = \left(\frac{d^2x}{dt^2} \right) \hat{i} + \left(\frac{d^2y}{dt^2} \right) \hat{j} \quad \text{--- (3.26)} \end{aligned}$$

Thus, the x and y components of the instantaneous acceleration are respectively given by

$$a_x = d^2x/dt^2 \quad \text{and} \quad a_y = d^2y/dt^2 \quad \text{--- (3.27)}$$

The magnitude and direction of the instantaneous acceleration are given by

$$a = \sqrt{\left(\frac{d^2x}{dt^2} \right)^2 + \left(\frac{d^2y}{dt^2} \right)^2} \quad \text{--- (3.28)},$$

$$\text{and} \quad \tan \theta = (dv_y/dt) / (dv_x/dt) = dv_y/dv_x \quad \text{--- (3.29)}$$

which is the slope of the tangent to the curve in velocity graph, i.e., a plot of v_y versus v_x .

Example 3.4: The position vectors of three particles are given by

$$\vec{x}_1 = (5\hat{i} + 5\hat{j}) \text{ m}, \vec{x}_2 = (5t\hat{i} + 5t\hat{j}) \text{ m} \quad \text{and}$$

$\vec{x}_3 = (5t\hat{i} + 10t^2\hat{j}) \text{ m}$ as a function of time t . Determine the velocity and acceleration for

each, in SI units.

Solution: $\vec{v}_1 = d\vec{x}_1/dt = 0$ as $\vec{x}_1$ does not depend on time t .

Thus, the particle is at rest.

$\vec{v}_2 = d\vec{x}_2/dt = 5\hat{i} + 5\hat{j}$ m/s. $\vec{v}_2$ does not change with time. $\therefore \vec{a}_2 = 0$

$v_2 = \sqrt{5^2 + 5^2} = 5\sqrt{2}$ m/s, $\tan \theta = 5/5 = 1$ or $\theta = 45^\circ$. Thus, the direction of v_2 makes an angle of 45° to the horizontal.

$\vec{v}_3 = d\vec{x}_3/dt = 5\hat{i} + 20t\hat{j}$.

$\therefore v_3 = \sqrt{5^2 + (20t)^2}$ m/s. Its direction is along

$\theta = \tan^{-1}\left(\frac{20t}{5}\right)$ with the horizontal.

$\vec{a}_3 = \frac{d\vec{v}_3}{dt} = 20\hat{j}$ m/s²

Thus, the particle 3 is getting accelerated along the y-axis at 20 m/s².

3.3.3 Equations of Motion for an Object travelling a Plane with Uniform Acceleration:

We have derived equations of motion for an object in rectilinear motion in section 3.2. We will now derive similar equations for a particle moving with uniform acceleration in two dimensions. Let the initial velocity of the object be $\vec{u}$ at $t = 0$ and its velocity at time t be $\vec{v}$. As the acceleration is constant, the average acceleration and the instantaneous acceleration will be equal. By using the definition of acceleration (Eq. (3.21)), we get

$$\vec{a} = (\vec{v} - \vec{u})/(t - 0)$$

$$\text{or } \vec{v} = \vec{u} + \vec{a}t \quad \text{--- (3.30)}$$

which is the same as Eq. (3.7) but is in vector form.

Let the displacement from time $t = 0$ to t be $\vec{s}$. This can be calculated from the average velocity of the object during this time. For

constant acceleration, $\vec{v}_{av} = \frac{\vec{u} + \vec{v}}{2}$

$$\therefore \vec{s} = (\vec{v}_{av})t = \left(\frac{\vec{u} + \vec{v}}{2}\right)t = \left(\frac{\vec{u} + \vec{u} + \vec{a}t}{2}\right)t$$

$$\therefore \vec{s} = \vec{u}t + \frac{1}{2}\vec{a}t^2 \quad \text{--- (3.31),}$$

which is the vector form of Eq. (3.8).

Eq. (3.30) and (3.31) can be resolved into their x and y components so as to get corresponding scalar equations as follows.

$$v_x = u_x + a_x t \quad \text{--- (3.32)}$$

$$\text{and } v_y = u_y + a_y t \quad \text{--- (3.33)}$$

$$s_x = u_x t + \frac{1}{2}a_x t^2 \quad \text{--- (3.34)}$$

$$\text{and } s_y = u_y t + \frac{1}{2}a_y t^2 \quad \text{--- (3.35)}$$

We can see that Eqs. (3.32) and (3.34) involve only the x components of displacement, velocity and acceleration while Eqs. (3.33) and (3.35) involve only the y components of these quantities. Thus the two sets of equations are independent of each other and can be solved independently. We can thus see that the motion along the x direction of an object is completely controlled by the x components of velocity and acceleration while that along the y direction is completely controlled by the y components of these quantities. This makes it easy to study the motion in two dimensions which gets converted to two independent rectilinear motions along two perpendicular directions.

Always Remember:

Motion in two dimensions can be resolved into two independent motions in mutually perpendicular directions.

Example 3.5: The initial velocity of an object is $\vec{u} = 5\hat{i} + 10\hat{j}$ m/s. Its constant acceleration is $\vec{a} = 2\hat{i} + 3\hat{j}$ m/s². Determine the velocity and the displacement after 5 s.

Solution:

$$\vec{v} = \vec{u} + \vec{a}t$$

$$= (5\hat{i} + 10\hat{j}) + (2\hat{i} + 3\hat{j})(5) = 15\hat{i} + 25\hat{j}$$

$$\therefore v = \sqrt{v_x^2 + v_y^2}$$

$$= \sqrt{15^2 + 25^2} = \sqrt{225 + 625} = \sqrt{850}$$

$$= 29.15 \text{ m/s}$$

Direction of $\vec{v}$ with x-axis is $\tan^{-1}\left(\frac{v_y}{v_x}\right) = \tan^{-1}$

$$\left(\frac{25}{15}\right) = \tan^{-1}(1.667) = 59^\circ$$

$$\vec{s} = \vec{u}t + \frac{1}{2}\vec{a}t^2$$

$$= (5\hat{i} + 10\hat{j})(5) + \frac{1}{2}(2\hat{i} + 3\hat{j})5^2$$

$$= 50\hat{i} + (87.5)\hat{j}$$

$$\therefore s = \sqrt{s_x^2 + s_y^2} = \sqrt{50^2 + 87.5^2}$$

$$= \sqrt{2500 + 7656.25}$$

$$= \sqrt{10156.25} = 100.78 \text{ m}$$

$$\text{at } \tan^{-1} \frac{87.5}{50} = 60^\circ 15' \text{ with x-axis.}$$

3.3.4 Relative Velocity:

Relative velocity between two objects moving in a plane can be defined in a way similar to that for objects moving along a straight line. The relative velocity of object A having velocity $\vec{v}_A$, with respect to the object B having velocity $\vec{v}_B$, is given by

$$\vec{v}_{AB} = \vec{v}_A - \vec{v}_B \quad \text{--- (3.36)}$$

Similarly, the relative velocity of object B with respect to object A, is given by

$$\vec{v}_{BA} = \vec{v}_B - \vec{v}_A \quad \text{--- (3.37)}$$

We can see that the magnitudes of the two relative velocities (v_{AB} and v_{BA}) are equal and their directions are opposite.

Consider a number of objects A, B, C, D --- Y, Z, moving with respect to the other. Using the symbol v_{AB} for representing the velocity of A relative to B etc, the velocity of A relative to Z can be written as

$$\vec{v}_{AZ} = \vec{v}_{AB} + \vec{v}_{BC} + \vec{v}_{CD} + \dots + \vec{v}_{XY} + \vec{v}_{YZ}$$

Note the order of subscripts (A $\rightarrow$ B $\rightarrow$ C $\rightarrow$ D --- $\rightarrow$ Z).

Example 3.6: An aeroplane is travelling northward with a speed of 300 km/hr with respect to the Earth, when wind is blowing from east to west at a speed of 100 km/hr. What is the velocity of the aeroplane with respect to the wind?

Solution: Let the velocity of the aeroplane with respect to Earth be $\vec{v}_{AE}$, velocity of wind with respect to Earth be $\vec{v}_{WE}$. The velocity of aeroplane with respect to wind, $\vec{v}_{AW}$ can be determined by the following expression:

$$\vec{v}_{AW} = \vec{v}_{AE} + \vec{v}_{EW} = \vec{v}_{AE} - \vec{v}_{WE} = -100\hat{i} + 300\hat{j},$$

considering north along +y axis.

$$\text{Magnitude of } \vec{v}_{AW} = \sqrt{(10000 + 90000)}$$

$$= 100\sqrt{10} \text{ km/hr, and its direction,}$$

$$\theta = \tan^{-1}\left(\frac{300}{-100}\right) = 71.6^\circ \text{ is towards north of east.}$$

3.3.5 Projectile Motion:

Any object in flight after being thrown with some velocity is called a projectile and its motion is called projectile motion. We often see projectile motion in our day-to-day life. Children throw stones towards trees for getting tamarind pods or mangoes. A bowler bowls a ball towards a batsman in cricket, a basket ball player throws a ball towards the basket, all these are illustrations of projectile motion. In this motion, we have objects (projectiles) with given initial velocity, moving under the influence of the Earth's gravitational field. The projectile has two components of velocity, one in the horizontal, i.e., along x-direction and the other in the vertical, i.e., along the y direction. The acceleration due to gravity acts only along the vertically downward direction. The horizontal component of velocity, therefore, remains unchanged as no force is acting in the horizontal direction, while the vertical component changes in accordance with laws of motion with $\vec{a}_x$ being 0 and $\vec{a}_y (= -\vec{g})$ being the downward acceleration due to gravity (upward is positive). Unless stated otherwise, retarding forces like air resistance, etc., are neglected for the projectile motion.

Let us assume that the initial velocity of the projectile is $\vec{u}$ and its direction makes an angle θ with the horizontal as shown in Fig. 3.5. The projectile is thrown from the ground. We take the x-axis along the ground and y-axis in the vertical direction. The horizontal and vertical components of initial velocity are u

$\cos\theta$ and $u \sin\theta$ respectively. The horizontal component remains unchanged in absence of any force acting in that direction, while the vertical component changes according to (Eq. 3.33) with $a_y = -g$ and $u_y = u \sin\theta$.

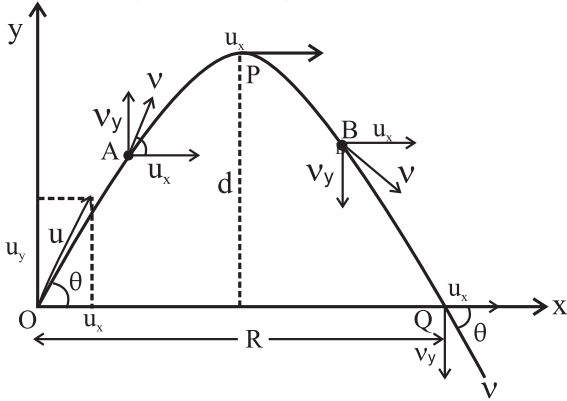


Fig.3.5: Trajectory of a projectile.

Thus, the components of velocity at time t are given by

$$v_x = u_x = u \cos\theta \quad \text{--- (3.38)}$$

$$v_y = u_y - gt = u \sin\theta - gt \quad \text{--- (3.39)}$$

As $0 < \theta < 90^\circ$, the vertical component initially is in the upward direction. Similarly, the displacements of the projectile in the horizontal and vertical directions at time t , according to Eqs. (3.34) and (3.35) are given by

$$s_x = u \cos\theta \cdot t \quad \text{--- (3.40)}$$

$$s_y = u \sin\theta \cdot t - \frac{1}{2} g t^2 \quad \text{--- (3.41)}$$

The direction of motion of the projectile at any time t makes an angle α with the horizontal which is given by

$$\tan \alpha = v_y(t)/v_x(t) \quad \text{--- (3.42)}$$

The vertical velocity keeps on decreasing as the projectile goes up and becomes zero at certain time. At that time the height of the projectile is maximum. The velocity then starts increasing in the downward direction as the particle is now falling under the Earth's gravitational field with a constant horizontal component of velocity. After a while the projectile reaches the ground. The trajectory of the object is shown in Fig. 3.5. The projectile is assumed to start from the origin of the coordinate system, O. The point of maximum height is indicated by P and the point where it falls down to the ground is indicated by Q. The horizontal and vertical components of velocity

are shown at these points as well as at two intermediate points A and B, on the trajectory of the projectile. Note that the horizontal component of velocity remains the same, i.e., u_x , while the vertical component decreases and becomes zero at P. After that it changes its direction, its magnitude increases and becomes equal to u_y again at Q. The horizontal distance covered by the projectile before it falls to the ground is OQ. We can derive the equation of the trajectory of the projectile as follows.

Let the time taken by the projectile to reach the maximum height be t_0 . The trajectory of the object being symmetrical, it can be shown by using equations of motion, that the object will take the same time in going up in air and coming down to the ground. At the highest point P, $t = t_0$ and $v_y = 0$. Using Eq. (3.39),

we get, $0 = u \sin\theta - g t_0$

$$t_0 = (u \sin\theta)/g \quad \text{--- (3.43)}$$

$\therefore$ Total time in air $= T = 2t_0$ is the **time of flight**.

The total horizontal distance travelled by the particle in this time T can be obtained by using Eq. (3.40) as

$$\begin{aligned} R &= u_x \cdot T = u \cos\theta \cdot 2t_0 = u \cos\theta \cdot (2u \sin\theta)/g \\ &= 2 u_x u_y / g = u^2 (2 \sin\theta \cos\theta) / g \\ &= u^2 \sin 2\theta / g \quad \text{--- (3.44)} \end{aligned}$$

This maximum horizontal distance travelled by the projectile is called the **horizontal range** R of the projectile and depends on the magnitude and direction of initial velocity of the projectile as well as the value of acceleration due to gravity at that place.

For maximum horizontal range,

$$\sin 2\theta = 1 \quad \therefore 2\theta = 90^\circ \text{ or } \theta = 45^\circ$$

$$\text{Hence, } R = R_{\max} = \frac{u^2}{g} \text{ for } \theta = 45^\circ$$

The **maximum height** H reached by the projectile, having certain value of θ , is the distance travelled along the vertical (y) direction in time t_0 . This can be calculated by using Eq. (3.41) as

$$\begin{aligned} H &= u \sin\theta \cdot t_0 - \frac{1}{2} g t_0^2 \\ &= u \sin\theta \left(\frac{u \sin\theta}{g} \right) - \frac{1}{2} g \left(\frac{u \sin\theta}{g} \right)^2 \end{aligned}$$

$$= \frac{u^2 \sin^2 \theta}{2g} = \frac{u_y^2}{2g} \quad \text{--- (3.45)}$$



Do you know ?

All the above expressions of T , R , $R_{\max}$ and H are valid if the entire motion is governed only by gravitational acceleration g , i.e., retarding forces like air resistance are absent. However, in reality, it is never so. As a result, time of ascent t_a and time of descent t_d are not equal but $t_a > t_d$. Also, in order to achieve maximum horizontal range for given initial velocity, the angle of projection should be greater than 45° and the range is much less than $\frac{u^2}{g}$.

Example 3.7: A stone is thrown with an initial velocity components of 20 m/s along the vertical, and 15 m/s along the horizontal direction. Determine the position and velocity of the stone after 3 s. Determine the maximum height that it will reach and the total distance travelled along the horizontal on reaching the ground. (Assume $g = 10 \text{ m/s}^2$)

Solution: The initial velocity of the stone in x-direction $= u \cos \theta = 15 \text{ m/s}$ and in y-direction $= u \sin \theta = 20 \text{ m/s}$.

After 3 s, $v_x = u \cos \theta = 15 \text{ m/s}$ and $v_y = u \sin \theta - gt = 20 - 10(3) = -10 \text{ m/s} = 10 \text{ m/s}$ downwards.

$$\begin{aligned} \therefore v &= \sqrt{v_x^2 + v_y^2} = \sqrt{15^2 + 10^2} \\ &= \sqrt{225 + 100} = \sqrt{325} \\ &= 18.03 \text{ m/s} \end{aligned}$$

$$\tan \alpha = v_y / v_x = 10/15 = 2/3$$

$$\therefore \alpha = \tan^{-1}(2/3) = 33^\circ 41' \text{ with the horizontal.}$$

$$s_x = (u \cos \theta) t = 15 \times 3 = 45 \text{ m,}$$

$$s_y = (u \sin \theta) t - \frac{1}{2} gt^2 = 20 \times 3 - 5(3)^2 = 15 \text{ m.}$$

Thus the stone will be at a distance 45 m along horizontal and 15 m along vertical direction from the initial position after time 3 s. The velocity is 18.03 m/s making an angle $33^\circ 41'$ with the horizontal.

The maximum vertical distance travelled is given by $H = (u \sin \theta)^2 / (2g) = 20^2 / (2 \times 10) = 20 \text{ m}$

Maximum horizontal distance travelled

$$R = 2 \cdot u_x \cdot u_y / g = 2(15)(20)/10 = 60 \text{ m}$$

Equation of motion for a projectile

We can derive the equation of motion of the projectile which is the relation between the displacements of the projectile along the vertical and horizontal directions. This can be obtained by eliminating t between the equations giving these displacements, i.e., Eqs. (3.40) and (3.41).

As the projectile starts from $\vec{x} = 0$, we can write $s_x = x$ and $s_y = y$.

$$\therefore s_x = (u \cos \theta) t \quad \therefore t = \frac{s_x}{u \cos \theta} = \frac{x}{u \cos \theta}$$

$$\therefore y = (u \sin \theta) t - \frac{1}{2} gt^2$$

$$= (u \sin \theta) \left(\frac{x}{u \cos \theta} \right) - \frac{1}{2} g \left(\frac{x}{u \cos \theta} \right)^2$$

$$\therefore y = (\tan \theta) x - \frac{1}{2} \left(\frac{g}{u^2 \cos^2 \theta} \right) x^2 \quad \text{--- (3.46)}$$

This is the equation of the trajectory of the projectile. Here, u and θ are constants for the given projectile motion. The above equation is of the form

$$y = Ax + Bx^2 \quad \text{--- (3.47)}$$

which is the equation of a parabola. Thus, the path, i.e., the trajectory of a projectile is a parabola.

3.4 Uniform Circular Motion:

An object moving with *constant speed* along a circular path is said to be in *uniform circular motion* (UCM). Such a motion is only possible if its velocity is *always tangential* to its circular path, *without change in its magnitude*.

To change the direction of velocity, acceleration is a must. However, if the acceleration or its component is in line with the velocity (along or opposite to the velocity), it will *always* change the speed (magnitude of velocity) in which case it will not continue its uniform circular motion. In order to achieve both these requirements, the acceleration must be (i) perpendicular to the tangential velocity, (ii) of constant magnitude and (iii) always directed

towards the centre of the circular trajectory. Such an acceleration is called *centripetal* (centre seeking) acceleration and the force causing this acceleration is *centripetal* force.

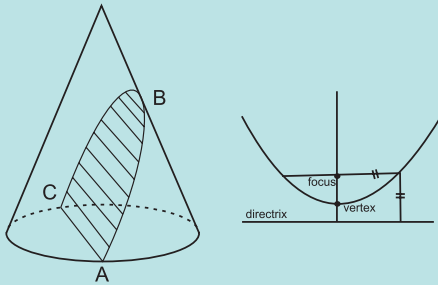
Thus, in order to realize a circular motion, there are two requirements; (i) *tangential velocity* and (ii) *centripetal force of suitable constant magnitude*.

An example is the motion of the moon going around the Earth in an early circular orbit as a result of the constant gravitational attraction of fixed magnitude felt by it towards the Earth.



Do you know ?

A parabola is a symmetrical open curve obtained by the intersection of a cone with a plane which is parallel to its side. Mathematically, the parabola is described with the help of a point called the focus and a straight line called the directrix shown in the accompanying figure. The parabola is the locus of all points which are equidistant from the focus and the directrix. The chord of the parabola which is parallel to the directrix and passes through the focus is called latus rectum of the parabola as shown in the accompanying figure.



3.4.1 Period, Radius Vector and Angular Speed:

Consider an object of mass m , moving with a uniform speed v , along a circle of radius r . Let T be the time period of revolution of the object, i.e., the time taken by the object to complete one revolution or to travel a distance of $2\pi r$.

Thus, $T = 2\pi r/v$

$$\therefore \text{Speed } v = \frac{\text{Distance}}{\text{Time}} = \frac{2\pi r}{T} \quad \text{--- (3.48)}$$

During circular motion of a point object, the position vector of the object from centre of

the circle is the radius vector $\vec{r}$. Its magnitude is radius r and it is directed away from the centre to the particle, i.e., away from the centre of the circle. As the particle performs UCM, this radius vector describes equal angles in equal intervals of time. At this stage we can define a new quantity called angular speed ω which gives the angle described by the radius vector, per unit time. It is analogous to speed which is distance travelled per unit time.

During one complete revolution, the angle described is 2π and the time taken is period T . Hence, the angular speed

$$\omega = \frac{\text{Angle}}{\text{time}} = \frac{2\pi}{T} = \frac{(2\pi)}{\left(\frac{2\pi r}{v}\right)} = \frac{v}{r} \quad \text{--- (3.49)}$$

The unit of ω is radian/sec.

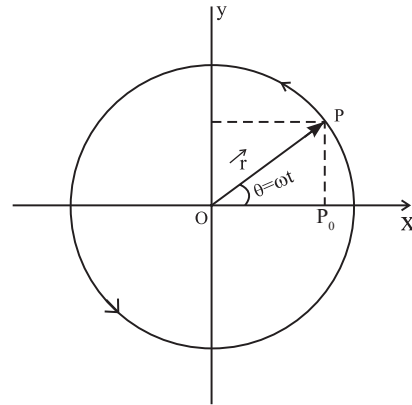


Fig.3.6: Uniform circular motion.

3.4.2 Expression for Centripetal Acceleration:

Figure 3.6 shows a particle P performing a UCM in anticlockwise sense along a circle of radius r with angular speed ω and period T . Let us choose the coordinates such that this motion is in the xy - plane having centre at the origin O . Initially (for simplicity), let the particle be at P_0 on the positive x -axis. At a given instant t , the radius vector of P makes an angle θ with the x -axis.

$$\therefore \theta = \omega t \text{ and so } \frac{d\theta}{dt} = \omega$$

x and y components of the radius vector $\vec{r}$ will then be $r\cos\theta$ and $r\sin\theta$ respectively.

$$\begin{aligned} \therefore \vec{r} &= (r\cos\theta)\hat{i} + (r\sin\theta)\hat{j} \\ &= (r\cos[\omega t])\hat{i} + (r\sin[\omega t])\hat{j} \quad \text{--- (3.50)} \end{aligned}$$

Time derivative of position vector $\vec{r}$ gives

instantaneous velocity $\vec{v}$ and time derivative of velocity $\vec{v}$ gives instantaneous acceleration $\vec{a}$. Magnitudes of r and ω are constants.

$$\begin{aligned}\therefore \vec{v} &= \frac{d\vec{r}}{dt} = r(-\omega \sin[\omega t] \hat{i} + \omega \cos[\omega t] \hat{j}) \\ &= r\omega(-\sin[\omega t] \hat{i} + \cos[\omega t] \hat{j})\end{aligned}\quad \text{--- (3.51)}$$

$$\begin{aligned}\therefore \vec{a} &= \frac{d\vec{v}}{dt} = r\omega(-\omega \cos[\omega t] \hat{i} - \omega \sin[\omega t] \hat{j}) \\ &= -\omega^2(r \cos[\omega t] \hat{i} + r \sin[\omega t] \hat{j}) = -\omega^2 \vec{r}\end{aligned}\quad \text{--- (3.52)}$$

Here minus sign shows that the acceleration is opposite to that of $\vec{r}$, i.e., towards the centre. This is the centripetal acceleration.

The magnitude of acceleration,

$$a = \omega^2 r = \frac{v^2}{r} = \omega v \quad \text{--- (3.53)}$$

The force providing this acceleration should also be along the same direction, hence centripetal.

$$\therefore \vec{F} = m\vec{a} = -m\omega^2 \vec{r} \quad \text{--- (3.54)}$$

$$\text{Magnitude of } F = m\omega^2 r = \frac{mv^2}{r} = m\omega v \quad \text{--- (3.55)}$$

Conical pendulum

In a simple pendulum a mass m is suspended by a string of length l and moves along an arc of a vertical circle. If the mass instead revolves in a horizontal circle and the string which makes a constant angle with the vertical describes a cone whose vertex is the fixed point O , then mass-string system is called a conical pendulum as shown in Fig. 3.7. In the absence of friction, the system will continue indefinitely once started.

As shown in the figure, the forces acting on the bob of mass, m , of the conical pendulum are: (i) Gravitational force, mg , acting vertically downwards, (ii) Force due to tension $\vec{T}$ acting along the string directed towards the support. These are the only two forces acting on the bob.

For the bob to undergo horizontal circular motion, (radius r) the resultant force must be centripetal, (directed towards the centre of the circle). In other words vertical gravitational force must be balanced.

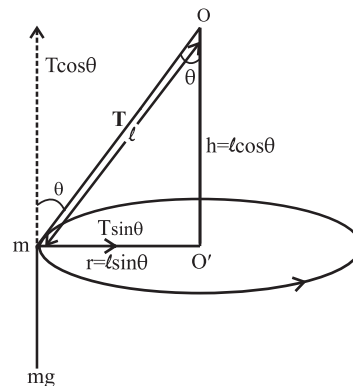


Fig 3.7: Conical pendulum

Thus, we resolve tension $\vec{T}$ into two mutually perpendicular components. Let θ be the angle made by the string with the vertical at any position. The component $T \cos \theta$ is acting vertically upwards. The inclination should be such that $T \cos \theta = mg$, so that there is no net vertical force.

The resultant force on the bob is then $T \sin \theta$ which is radial or centripetal or directed towards centre O' $T \sin \theta = mv^2/r = m\omega^2 r$.

$$\tan \theta = \frac{(mv^2/r)}{mg} = \frac{v^2}{rg}$$

$$\text{Since we know } v = \frac{2\pi r}{T}$$

$$\therefore \tan \theta = \frac{4\pi^2 r^2}{T^2 rg}$$

$$T = 2\pi \sqrt{\frac{r}{g \tan \theta}}$$

$$T = 2\pi \sqrt{\frac{l \sin \theta}{g \tan \theta}} \quad (\because r = l \sin \theta)$$

$$T = 2\pi \sqrt{\frac{l \cos \theta}{g}}$$

$$T = 2\pi \sqrt{\frac{h}{g}} \quad (\because h = l \cos \theta) \quad \text{--- (3.56)}$$

where l is length of the pendulum and h is the vertical distance of the horizontal circle from the fixed point O .

Example 3.8: An object of mass 50 g moves uniformly along a circular orbit with an angular speed of 5 rad/s. If the linear speed of the particle is 25 m/s, what is the radius of the

circle? Calculate the centripetal force acting on the particle.



Do you know ?

1. The centripetal force is not one of the external forces acting on the object. As can be seen from above, the actual forces acting on the bob are T and mg , the resultant of these is the centripetal force. Conversely, if the resultant force is centripetal, motion must be circular.
2. In planetary motion, the gravitational force between Sun and the planets provides the necessary centripetal force for the circular motion.

Solution: The linear speed and angular speed are related by $v = \omega r$

$$\therefore r = v/\omega = 25/5 \text{ m} = 5 \text{ m}.$$

$$\text{Centripetal force acting on the object} = \frac{mv^2}{r} = \frac{0.05 \times 25^2}{5} = 6.25 \text{ N}.$$

Example 3.9: An object is travelling in a horizontal circle with uniform speed. At

$t = 0$, the velocity is given by $\vec{u} = 20\hat{i} + 35\hat{j}$ km/s. After one minute the velocity becomes $\vec{v} = -20\hat{i} - 35\hat{j}$. What is the magnitude of the acceleration?

Solution: Magnitude of initial and final velocities =

$$= u = \sqrt{(20)^2 + (35)^2} \text{ m/s}$$

$$= \sqrt{1625} \text{ m/s}$$

$$= 40.3 \text{ m/s}$$

As the velocity reverses in 1 min, the time period of revolution is 2 min.

$$T = \frac{2\pi r}{u}, \text{ giving } r = \frac{uT}{2\pi}$$

$$a = \frac{u^2}{r} = \frac{u^2 2\pi}{uT} = \frac{2\pi u}{T} = \frac{2 \times 3.14 \times 40.3}{2 \times 60}$$

$$= 2.11 \text{ m/s}^2$$



Internet my friend

1. hyperphysics.phy-astr.gsu.edu/hbase/mot.html#motcon
2. www.college-physics.com/book/mechanics



Exercises

1. Choose the correct option.

- i) An object thrown from a moving bus is an example of
 - (A) Uniform circular motion
 - (B) Rectilinear motion
 - (C) Projectile motion
 - (D) Motion in one dimension
- ii) For a particle having a uniform circular motion, which of the following is constant
 - (A) Speed
 - (B) Acceleration
 - (C) Velocity
 - (D) Displacement
- iii) The bob of a conical pendulum under goes
 - (A) Rectilinear motion in horizontal plane
 - (B) Uniform motion in a horizontal circle
 - (C) Uniform motion in a vertical circle
 - (D) Rectilinear motion in vertical circle
- iv) For uniform acceleration in rectilinear motion which of the following is not correct?
 - (A) Velocity-time graph is linear
 - (B) Acceleration is the slope of velocity time graph
 - (C) The area under the velocity-time graph equals displacement
 - (D) Velocity-time graph is nonlinear
- v) If three particles A, B and C are having velocities $\vec{v}_A$, $\vec{v}_B$ and $\vec{v}_C$ which of the following formula gives the relative velocity of A with respect to B
 - (A) $\vec{v}_A + \vec{v}_B$
 - (B) $\vec{v}_A - \vec{v}_C + \vec{v}_B$

(C) $\vec{v}_A - \vec{v}_B$ (D) $\vec{v}_C - \vec{v}_A$

2. Answer the following questions.

- Separate the following in groups of scalar and vectors: velocity, speed, displacement, work done, force, power, energy, acceleration, electric charge, angular velocity.
- Define average velocity and instantaneous velocity. When are they same?
- Define free fall.
- If the motion of an object is described by $x = f(t)$ write formulae for instantaneous velocity and acceleration.
- Derive equations of motion for a particle moving in a plane and show that the motion can be resolved in two independent motions in mutually perpendicular directions.
- Derive equations of motion graphically for a particle having uniform acceleration, moving along a straight line.
- Derive the formula for the range and maximum height achieved by a projectile thrown from the origin with initial velocity $\vec{u}$ at an angle θ to the horizontal.
- Show that the path of a projectile is a parabola.
- What is a conical pendulum? Show that its time period is given by $2\pi\sqrt{\frac{l\cos\theta}{g}}$, where l is the length of the string, θ is the angle that the string makes with the vertical and g is the acceleration due to gravity.
- Define angular velocity. Show that the centripetal force on a particle undergoing uniform circular motion is $-m\omega^2\vec{r}$.

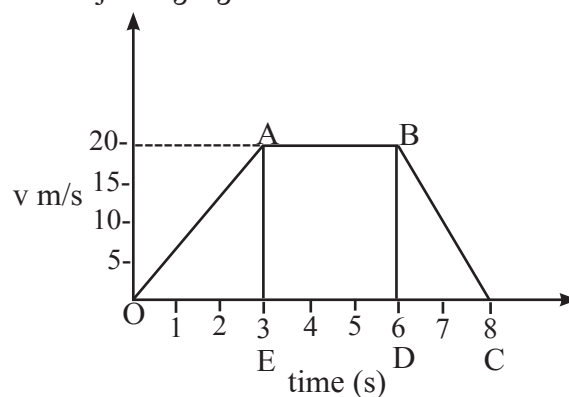
3. Solve the following problems.

- An aeroplane has a run of 500 m to take off from the runway. It starts from rest and moves with constant acceleration to cover the runway in 30 sec. What is the velocity of the aeroplane at the take off?
[Ans: 120 km/hr]

- A car moving along a straight road with a speed of 120 km/hr, is brought to rest by applying brakes. The car covers a distance of 100 m before it stops. Calculate (i) the average retardation of the car (ii) time taken by the car to come to rest.
[Ans: 50/9 m/sec², 6 sec]

- A car travels at a speed of 50 km/hr for 30 minutes, at 30 km/hr for next 15 minutes and then 70 km/hr for next 45 minutes. What is the average speed of the car?
[Ans: 56.66 km/hr]

- A velocity-time graph is shown in the adjoining figure.



Determine:

- initial speed of the car (ii) maximum speed attained by the car (iii) part of the graph showing zero acceleration (iv) part of the graph showing constant retardation (v) distance travelled by the car in first 6 sec.
[Ans: (i) 0 (ii) 20 m/sec (iii) AB (iv) BC (v) 90 m]
- A man throws a ball to maximum horizontal distance of 80 meters. Calculate the maximum height reached.
[Ans: 20 m]
- A particle is projected with speed v_0 at angle θ to the horizontal on an inclined surface making an angle ϕ ($\phi < \theta$) to the horizontal. Find the range of the projectile along the inclined surface.

[Ans: $R = \frac{2v_0^2 \cos\theta \sin(\theta - \phi)}{g \cos^2 \phi}$]

vii) A metro train runs from station A to B to C. It takes 4 minutes in travelling from station A to station B. The train halts at station B for 20 s. Then it starts from station B and reaches station C in next 3 minutes. At the start, the train accelerates for 10 sec to reach the constant speed of 72 km/hr. The train moving at the constant speed is brought to rest in 10 sec. at next station. (i) Plot the velocity- time graph for the train travelling from the station A to B to C. (ii) Calculate the distance between the stations A, B and C.

[Ans: AB = 4.6 km, BC = 3.4 km]

viii) A train is moving eastward at 10 m/sec. A waiter is walking eastward at 1.2 m/sec; and a fly is flying toward the north across the waiter's tray at 2 m/s. What is the velocity of the fly relative to Earth

[Ans: 11.4 m/s, 10° due north of east]

ix) A car moves in a circle at the constant speed of 50 m/s and completes one revolution in 40 s. Determine the magnitude of acceleration of the car.

[Ans: 7.85 m s^{-2}]

x) A particle moves in a circle with constant speed of 15 m/s. The radius of the circle is 2 m. Determine the centripetal acceleration of the particle.

[Ans: 112.5 m s^{-2}]

xi) A projectile is thrown at an angle of 30° to the horizontal. What should be the range of initial velocity (u) so that its range will be between 40 m and 50 m? Assume $g = 10 \text{ m s}^{-2}$.

[Ans: $21.49 \leq u \leq 24.03 \text{ m s}^{-2}$]



Can you recall?

1. What are different types of motions?
2. What do you mean by kinematical equations and what are they?
3. Newton's laws of motion apply to most bodies we come across in our daily lives.
4. All bodies are governed by Newton's law of gravitation. Gravitation of the Earth results into weight of objects.
5. Acceleration is directly proportional to force for fixed mass of an object.
6. Bodies possess potential energy and kinetic energy due to their position and motion respectively which may change. Their total energy is conserved in absence of any external force.

4.1. Introduction:

If an object continuously changes its position, it is said to be in motion. Mechanics is a branch of Physics that deals with motion. There are basically two branches of mechanics (i) Statics, where we deal with objects at rest or in equilibrium under the action of balanced forces and (ii) Kinetics, which deals with actual motion.

Kinetics can be further divided into two branches **(i) Kinematics:** In kinematics, we describe various motions without discussing their cause. Various parameters discussed in kinematics are distance, displacement, speed, velocity and acceleration. **(ii) Dynamics:** In dynamics we describe the motion along with its cause, which is force and/or torque. Parameters discussed in dynamics are momentum, force, energy, power, etc. in addition to those in kinematics.

It must be understood that motion is strictly a relative concept, i.e., it should always be described in context to a reference frame. For example, if you are in a running bus, neither you nor your co-passengers sitting in the bus are in motion in your reference, i.e., moving bus. However, from the ground reference, bus, you and all the passengers are in motion.

If not random, motions in real life may be understood separately as linear, circular or rotational, oscillatory, etc., or some combinations of these. While describing any of these, we need to know the corresponding forces responsible for these motions. Trajectory of any motion is decided by acceleration $\vec{a}$ and the initial velocity $\vec{u}$.

a) Linear motion: Initial velocity may be zero or non-zero. If initial velocity is zero (starting from rest), acceleration in any direction will result into a linear motion.

If initial velocity is not zero, the acceleration must be in line with the initial velocity (along the same or opposite direction to that of the initial velocity) for resultant motion to be linear.

b) Circular motion: If initial velocity is not zero and acceleration is throughout perpendicular to the velocity, the resultant motion will be circular.

c) Parabolic motion: If acceleration is constant and initial velocity is *not* in line with the acceleration, the motion is parabolic, e.g., trajectory of a projectile motion.

d) Other combinations of $\vec{u}$ and $\vec{a}$ will result into different more complicated motion.

4.2. Aristotle's Fallacy:

Aristotle (384BC-322BC) stated that "an external force is required to keep a body in uniform motion". This was probably based on a common experience like a ball rolling on a surface stops after rolling through some distance. Thus, to keep the ball moving with constant velocity, we have to continuously apply a force on it. Similar examples can be found elsewhere, like a paper plane flying through air or a paper boat propelled with some initial velocity.

Correct explanation to Aristotle's fallacy was first given by Galileo (1564-1642), which was later used by Newton (1643-1727) in

formulating *laws of motion*. Galileo showed that all the objects stop moving because of some resistive or opposing forces like friction, viscous drag, etc. In these examples such forces are frictional force for rolling ball, viscous drag or viscous force of air for paper plane and viscous force of water for the boat.

Thus, in reality, for an uninterrupted motion of a body an additional external force is required for overcoming these opposing forces.



Can you tell?

1. Was Aristotle correct?
2. If correct, explain his statement with an illustration.
3. If wrong, give the correct modified version of his statement.

4.3. Newton's Laws of Motion:

First law: Every inanimate object continues to be in its state of rest or of uniform *unaccelerated* motion unless and until it is acted upon by an *external, unbalanced* force.

Second law: Rate of change of linear momentum of a rigid body is directly proportional to the applied force and takes place in the direction of the applied force. On selecting suitable units, it

takes the form $\vec{F} = \frac{d\vec{p}}{dt}$ (where $\vec{F}$ is the force and $\vec{p} = m\vec{v}$ is the linear momentum).

Third law: To every action (force), there is an equal and opposite reaction (force).

Discussion: From Newton's second law of motion, $\vec{F} = \frac{d\vec{p}}{dt} = \frac{d}{dt}(m\vec{v})$. For a *given* body, mass m is constant.

$$\therefore \vec{F} = m \frac{d\vec{v}}{dt} = m\vec{a} \dots \text{(for constant mass)}$$

Thus, if $\vec{F} = 0$, $\vec{v}$ is constant. Hence if there is no force, velocity will not change. This is nothing but Newton's first law of motion.



Can you tell?

What is then special about Newton's first law if it is derivable from Newton's second law?

4.3.1. Importance of Newton's First Law of Motion:

- (i) It shows an equivalence between 'state of rest' and 'state of uniform motion along a straight line' as both need a net unbalanced force to change the state. Both these are referred to as 'state of motion'. The distinction between state of rest and uniform motion lies in the choice of the 'frame of reference'.
- (ii) It defines *force* as an entity (or a physical quantity) that brings about a change in the 'state of motion' of a body, i.e., force is something that initiates a motion or controls a motion. Second law gives its quantitative understanding or its mathematical expression.
- (iii) It defines *inertia* as a fundamental property of every physical object by which the object *resists* any change in its state of motion. Inertia is measured as the mass of the object. More specifically it is called inertial mass, which is the ratio of net force ($|\vec{F}|$) to the corresponding acceleration ($|\vec{a}|$).

4.3.2. Importance of Newton's Second Law of Motion:

- (i) It gives mathematical formulation for quantitative measure of force as rate of change of linear momentum.



Do you know ?

Mathematical expression for force must be

$$\begin{aligned} \text{remembered as } \vec{F} &= \frac{d\vec{p}}{dt} \text{ and } \textbf{not} \text{ as } \vec{F} = m\vec{a} \\ \vec{F} &= \frac{d\vec{p}}{dt} = \frac{d(m\vec{v})}{dt} = \frac{dm}{dt}(\vec{v}) + m\left(\frac{d\vec{v}}{dt}\right) \\ &= \frac{dm}{dt}(\vec{v}) + m\vec{a} \end{aligned}$$

For a *given* body, mass is constant, i.e., $\frac{dm}{dt} = 0$ and only in this case, $\vec{F} = m\vec{a}$
In the case of a rocket, both the terms are needed as both mass and velocity are varying.

- (ii) It defines *momentum* ($\vec{p} = m\vec{v}$) instead of velocity as the fundamental quantity related to motion. What is changed by a force is the momentum and not necessarily the velocity.
- (iii) Aristotle's fallacy is overcome by considering *resultant unbalanced force*.

4.3.3 Importance of Newton's Third Law of Motion:

- (i) It defines *action* and *reaction* as a pair of equal and opposite forces acting along the same line.
- (ii) Action and reaction forces are always on *different* objects.

Consequences:

Action force exerted by a body x on body y , conventionally written as $\vec{F}_{yx}$, is the force experienced by y .

As a result, body y exerts *reaction force* $\vec{F}_{xy}$ on body x .

In this case, body x experiences the force $\vec{F}_{xy}$ only while the body y experiences the force $\vec{F}_{yx}$ only.

Forces $\vec{F}_{xy}$ and $\vec{F}_{yx}$ are equal in magnitude and opposite in their directions, but there is no question of cancellation of these forces as those are experienced by *different* objects.

Forces $\vec{F}_{xy}$ and $\vec{F}_{yx}$ *need not* be contact forces. Repulsive forces between two magnets is a pair of action-reaction forces. In this case the two magnets are not in contact. Gravitational force between Earth and moon or between Earth and Sun are also similar pairs of non-contact action-reaction forces.

Example 4.1: A hose pipe used for gardening is ejecting water horizontally at the rate of 0.5 m/s. Area of the bore of the pipe is 10 cm². Calculate the force to be applied by the gardener to hold the pipe *horizontally* stationary.

Solution: If ejecting water horizontally is considered as action force on the water, the water exerts a backward force (called recoil force) on the pipe as the reaction force.

$$\vec{F} = \frac{d\vec{p}}{dt} = \frac{d(m\vec{v})}{dt} = \frac{dm}{dt}\vec{v} + m\frac{d\vec{v}}{dt}$$

As v , the velocity of ejected water is constant, $F = \frac{dm}{dt}\vec{v}$, where $\frac{dm}{dt}$ is the rate at which mass of water is ejected by the pipe.

As the force is in the direction of velocity (horizontal), we can use scalars. $\therefore F = \frac{dm}{dt}v$

$$\frac{dm}{dt} = \frac{d(V\rho)}{dt} = \frac{d(Al\rho)}{dt} = A\rho\frac{dl}{dt} = A\rho v$$

where V = volume of water ejected

A = area of cross section of bore = 10 cm²

ρ = density of water = 1 g/cc

l = length of the water ejected in time t

$$\frac{dl}{dt} = v = \text{velocity of water ejected} \\ = 0.5 \text{ m/s} = 50 \text{ cm/s}$$

$$F = \frac{dm}{dt}v = (A\rho v)v = A\rho v^2 = 10 \times 1 \times 50^2 \\ = 25000 \text{ dyne} = 0.25 \text{ N}$$

Equal and opposite force must be applied by the gardener.

4.4. Inertial and Non-Inertial Frames of Reference

Consider yourself standing on a railway platform or a bus stand and you see a train or bus moving. According to you, that train or bus is moving or is in motion. As per the experience of the passengers in the train or bus, they are at rest and you are moving (in backward direction). Hence motion itself is a relative concept. To know or describe a motion you need to describe or define some reference. Such a reference is called a *frame of reference*. In the example discussed above, if you consider the platform as the reference, then the passengers and the train are moving. However, if the train is considered as the reference, you and platform, etc. are moving.

Usually a set of coordinates with a suitable origin is enough to describe a frame of reference. If position coordinates of an object are continuously changing with time in a frame of reference, then *that object* is in motion in *that* frame of reference. Any frame of reference in which Newton's first law of motion is applicable is the simplest understanding of an

inertial frame of reference. It means, if there is no net force, there is no acceleration. Thus in an inertial frame, a body will move with constant velocity (which may be zero also) if there is no net force acting upon it. In the absence of a net force, if an object suffers an acceleration, that frame of reference is **not** an inertial frame and is called as non-inertial frame of reference.

Measurements in one inertial frame can be converted to measurements in another inertial frame by a simple transformation, i.e., by simply using some velocity vectors (relative velocity between the two frames of reference).

Illustration: Imagine yourself inside a car with all windows opaque so that you can not see anything outside. Also consider that there is a pendulum tied inside the car and not set into oscillations. If the car just starts its motion (with reference to outside or ground), you will experience a jerk, i.e., acceleration inside the car even though there is no force acting upon you. During this time, the string of the pendulum may be steady, but **not** vertical. During time of acceleration, the car can be considered to be a non-inertial frame of reference. Later on if the car is moving with constant velocity (with reference to the ground), you will not experience any jerky motion within the car and the car can be considered as an inertial frame of reference. In this case, the pendulum string will be vertical, when not oscillating.



Do you know ?

The situations/phenomena that can be explained using Newton's laws of motion fall under Newtonian mechanics. So far as our daily life situations are considered, Newtonian mechanics is perfectly applicable. However, under several extreme conditions we need to use some other theories.

Limitation of Newton's laws of motion

- (i) Newton's laws are applicable only in the inertial frames of reference (discussed later). If the body is in a frame of reference of acceleration (a), we need to use a *pseudo* force ($-m\vec{a}$) in addition to all the other forces while writing the

force equations.

- (ii) Newton's laws are applicable for *point* objects.
- (iii) Newton's laws are applicable to rigid bodies. A body is said to be *rigid* if the relative distances between its particles do not change for any deforming force.
- (iv) For objects moving with speeds comparable to that of light, Newton's laws of motion do not give results that match with the experimental results and Einstein's special theory of relativity has to be used.
- (v) Behaviour and interaction of objects having atomic or molecular sizes cannot be explained using Newton's laws of motion, and quantum mechanics has to be used.

A rocket in intergalactic space (gravity free space between galaxies) with all its engines shut is closest to an ideal inertial frame. However, Earth's acceleration in the reference frame of the Sun is so small that any frame attached to the Earth can be used as an inertial frame for any day-to-day situation or in our laboratories.

4.5 Types of Forces:

4.5.1. Fundamental Forces in Nature:

All the forces in nature are classified into following **four** interactions that are termed as fundamental forces.

- (i) Gravitational force: It is the attractive force between two (point) masses separated by a distance. Magnitude of gravitational force between point masses m_1 and m_2 separated by distance r is given

$$\text{by } F = \frac{Gm_1m_2}{r^2}$$

where $G = 6.67 \times 10^{-11}$ SI units. Between two point masses (particles) separated by a given distance, this is the weakest force having infinite range. This force is always attractive. Structure of the universe is governed by this force.

Common experience of this force for us is gravitational force exerted by Earth on us, which we call as our weight W .

$$\therefore W = \frac{GMm}{R^2} = m \left(\frac{GM}{R^2} \right) = mg$$

where M and R represent respectively mass and radius of the Earth. Distance between ourselves and Earth is taken as radius of the Earth when we are on the surface of the Earth because our size is negligible as compared to radius of the Earth (6.4×10^6 m).

$$\left(\frac{GM}{R^2} \right) = \frac{6.67 \times 10^{-11} \times 6 \times 10^{24}}{(6.4 \times 10^6)^2}$$

$\cong 9.8 \text{ m/s}^2 = g = \text{gravitational acceleration or gravitational field intensity.}$

We **feel** this force only due to normal reaction from the surface of our contact with Earth.

All individual bodies also exert gravitation force on each other but it is too small compared to that by the Earth. For example, mutual gravitational force between two SUMO wrestlers, each of mass 300 kg, assuming the distance between them is 0.5 m, will be

$$F = \frac{6.67 \times 10^{-11} \times 300 \times 300}{0.5^2}$$

$$\cong 2.4 \times 10^{-5} \text{ N} = 24 \text{ } \mu\text{N.}$$

This force is negligibly small in comparison to the weight of each SUMO wrestler $\cong 3000 \text{ N}$.

- (ii) **Electromagnetic (EM) force:** It is an attractive or repulsive force between electrically charged particles. Earlier, electric and magnetic forces were thought to be independent. After the demonstrations by Michael Faraday (1791-1867) and James Clerk Maxwell (1831-1879), electric and magnetic forces were unified through the theory of electromagnetism. These forces are stronger than the gravitational force. Our life is practically governed by these forces. Majority of forces experienced in our daily life, such as force of friction, normal reaction, tension in strings,

collision forces, elastic forces, viscosity (fluid friction), etc. are EM in nature. Under the action of these forces, there is *deformation of objects* that changes intermolecular distances thereby resulting into reaction forces.

- (iii) **Strong (nuclear) force:** This is the strongest force that binds the nucleons together inside a nucleus. Though strongest, it is a short range ($< 10^{-14} \text{ m}$) force. Therefore is very strong attractive force and is charge independent.
- (iv) **Weak (nuclear) force:** This is the interaction between subatomic particles that is responsible for the radioactive decay of atoms, in particular beta emission. The weak nuclear force is not as weak as the gravitational force, but much weaker than the strong nuclear and EM forces. The range of weak nuclear force is exceedingly small, of the order of 10^{-16} m .

Weak interaction force:

The radioactive isotope C^{13} is converted into N^{14} in which a neutron is converted into a proton. This property is used in *carbon dating* to determine the age of a sample.

In radioactive beta decay, the nucleus emits an electron (or positron) and an uncharged particle called neutrino. There are two types of β -decay, β^+ and β^- . During β^+ decay, a proton is converted into a neutron (accompanied by positron emission) and during β^- decay a neutron is converted into a proton (accompanied by electron emission).

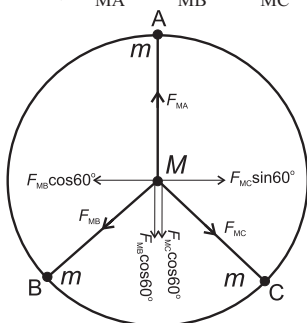
Another most interesting illustration of weak forces is fusion reaction in the core of the Sun. During this, protons are converted into neutrons and a neutrino is emitted due to energy balance. In general, emission of a neutrino is the evidence that there is conversion of a proton into a neutron or a neutron into a proton. This is possible *only* due to weak forces.

Example 4.2. Three identical point masses are fixed symmetrically on the periphery of a circle. Obtain the resultant gravitational force on any point mass M at the centre of the circle. Extend this idea to more than three identical masses symmetrically located on the periphery. How far can you extend this concept?

Solution:

- (i) Figure below shows three identical point masses m on the periphery of a circle of radius r . Mass M is at the centre of the circle. Gravitational forces on M due to these masses are attractive and are directed as shown.

In magnitude, $F_{MA} = F_{MB} = F_{MC} = \frac{GMm}{r^2}$



Forces F_{MB} and F_{MC} are resolved along F_{MA} and perpendicular to F_{MA} as shown. Components perpendicular to F_{MA} cancel each other. Components along F_{MA} are $F_{MB} \cos 60^\circ = F_{MC} \cos 60^\circ = \frac{1}{2} F_{MA}$ each. Magnitude of their resultant is $\frac{1}{2} F_{MA}$ and its direction is opposite to that of F_{MA} . Thus, the resultant force on mass M is zero.

- (ii) For any even number of equal masses, the force due to any mass m is balanced (cancelled) by diametrically opposite mass. For any odd number of masses, as seen for 3, the components perpendicular to one of them cancel each other while the components parallel to one of these add up in such a way that the resultant is zero for any number of identical masses m located symmetrically on the periphery.
- (iii) As the number of masses tends to infinity, their collective shape approaches circumference of the circle, which is nothing but a ring. Thus, the gravitational

force exerted by a ring mass on any other mass at its centre is zero.

In three-dimensions, we can imagine a uniform hollow sphere to be made up of infinite number of such rings with a common diameter. Thus, the gravitational force for any mass kept at the centre of a hollow sphere is zero.



Do you know ?

Unification of forces: Newton unified terrestrial (related to Earth and hence to our daily life) and celestial (related to universe) domains under a common law of gravitation. The experimental discoveries of Oersted (1777-1851) and Faraday showed that electric and magnetic phenomena are in general inseparable leading to what is called 'EM phenomenon'. Electromagnetism and optics were unified by Maxwell with the proposition that light is an EM wave. Einstein attempted to unify gravity and electromagnetism under general relativity but could not succeed. The EM and the weak nuclear force have now been unified as a single 'electro-weak' force.

4.5.2. Contact and Non-Contact Forces:

For some forces, like gravitational force, electrostatic force, magnetostatic force, etc., physical contact is not an essential condition. These forces exist even if the objects are distant or physically separated. Such forces are non-contact forces.

Forces resulting *only* due to contact are called contact forces. All these are EM in nature, arising due to some deformation. Normal reaction, forces occurring during collision, force of friction, etc., are contact forces. There are two common categories of contact forces. Two objects in contact, while exerting mutual force, try to push each other away along their common normal. Quite often we call it as 'normal reaction' force or 'normal' force. While standing on a table, we *push* the table away from us (downward) and the table *pushes* us away from it (upward) both being equal in magnitude and acting along the same 'normal' line.

Force of friction is also a contact force that arises whenever there is a relative motion or tendency of relative motion between surfaces in contact. This is the parallel (or tangential) component of the reaction force. In this case, the molecules of surfaces in contact, which have developed certain equilibrium, are required to be separated.

4.5.3 Real and Pseudo Forces:

Consider ourselves inside a lift (or elevator). When the lift just starts moving up (accelerates upward), we feel a bit heavier as if someone is pushing us down. This is not imaginary or not just a feeling. If we are standing on a weighing scale inside this lift, during this time the weighing scale records an increase in weight. During travelling with uniform upward velocity no such change is recorded. While stopping at some upper floor, the lift undergoes downward acceleration for decreasing the upward velocity. In this case the weighing scale records loss in weight and we also feel lighter. These *extra* upward or downward forces are (i) Measurable, means they are not imaginary, (ii) not accountable as per Newton's second law in the inertial frame and (iii) not among any of the four fundamental forces.

When we are inside a bus such forces are experienced when the bus starts to move (forward acceleration), when the bus is about to stop (backward acceleration) or takes a turn (centripetal acceleration). In all these cases we are inside an accelerated system (which is our frame of reference). If a force measuring device is suitably used – like the weighing scale recording the change in weight – these forces can be recorded and will be found to be always opposite to the acceleration of your frame of reference. They are also exactly equal to $-m\vec{a}$, where m is our mass and $\vec{a}$ is acceleration of the system (frame of reference).

We have already defined or described real forces to be those which obey Newton's laws of motion and are one of the four fundamental forces. Forces in above illustrations do not satisfy this description and cannot be called real forces. Hence these are called **pseudo** forces.

Pseudo in this case does not mean imaginary (because these are measurable with instruments) but just means *non-real*. These forces are measured to be $(-m\vec{a})$. Hence, a term $(-m\vec{a})$ added to resultant force enables us to apply Newton's laws of motion to a non-inertial frame of reference of acceleration $\vec{a}$. Negative sign refers to their direction, which is opposite to that of the acceleration of the reference frame.

As per the illustration of the lift with downward acceleration $\vec{a}$, the weight experienced will be $\vec{W} = m\vec{g} + (-m\vec{a})$

As $\vec{g}$ and $\vec{a}$ are along the same direction in this case, $W = mg - ma$. This explains the *feeling* of a loss in weight.

During upward acceleration, say $\vec{a}_1$, we have, $\vec{W}_1 = m\vec{g} + (+m\vec{a}_1)$

In this case, $\vec{g}$ and $\vec{a}_1$ are oppositely directed. $\therefore W_1 = mg + (+ma_1) = mg + ma_1$ that explains gain in weight or existence of extra downward force.

In mathematics we define a number to be real if its square is zero or positive. Solution set of equations like $x^2 - 6x + 10 = 0$ does not satisfy the criterion to be a real number. Such numbers are *complex* numbers which include $i = \sqrt{-1}$ along with some real part. It means every *non-real* need not be *imaginary* as in literal verbal sense.

Example 4.3: A car of mass 1.5 ton is running at 72 kmph on a straight horizontal road. On turning the engine off, it stops in 20 seconds. While running at the same speed, on the same road, the driver observes an accident 50 m in front of him. He immediately applies the brakes and just manages to stop the car at the accident spot. Calculate the braking force.

Solution: On turning the engine off,

$$u = 20 \text{ m s}^{-1}, v = 0, t = 20 \text{ s}$$

$$\therefore a = \frac{v - u}{t} = -1 \text{ m s}^{-2}$$

This is frictional retardation (negative acceleration).

After seeing the accident,

$$u = 20 \text{ m s}^{-1}, v = 0, s = 50 \text{ m}$$

$$\therefore a_1 = \frac{v^2 - u^2}{2s} = -4 \text{ m s}^{-2}$$

This retardation is the combined effect of braking and friction. Thus, braking retardation $= 4 - 1 = 3 \text{ m s}^{-2}$.

$$\therefore \text{Braking force} = \text{mass} \times \text{braking retardation} = 1500 \times 3 = 4500 \text{ N}$$

4.5.4. Conservative and Non-Conservative Forces and Concept of Potential Energy:

Consider an object lying on the ground is lifted and kept on a table. Neglecting air resistance, the amount of work done is the work done *against* gravitational force and it is *independent* of the actual path chosen (Remember, as there is no air resistance, gravitational force is the only force). Similarly, while keeping the same object back on the ground from the table, the work is done *by* the gravitational force. In either case the amount of work done is the same **and** is independent of the actual path chosen. The work done by force $\vec{F}$ in moving the object through a distance $d\vec{x}$ can be mathematically represented as $dW = \vec{F} \cdot d\vec{x} = -dU$ or $dU = -\vec{F} \cdot d\vec{x}$.

If work done by or against a force is independent of the actual path, the force is said to be a **conservative force**. During the work done by a conservative force, the mechanical energy (sum of kinetic and potential energy) is conserved. In fact, we define the concept of potential at a point or potential energy (in the topic of gravitation) with the help of conservative forces only. The work done by or against conservative forces reflects an equal amount of change in the *potential energy*. The corresponding work done is used in changing the position or in achieving the new position in the gravitational field. Hence, potential energy is often referred to as the energy possessed on account of position.

In the illustration given above, the work done is reflected as increase in the gravitational potential energy when the displacement is

against the (gravitational) force. Same amount of potential energy is decreased when the displacement is in the direction of force. In either case it is independent of the actual path but depends only upon the initial and final positions. This change in the potential energy takes place in such a way that the mechanical energy is conserved.

As discussed above, increase in the potential energy, $dU = -\vec{F} \cdot d\vec{x}$ or $U = -\int \vec{F} \cdot d\vec{x}$ where $\vec{F}$ is a conservative force. This concept, will be described in details in Chapter 5 on Gravitation in context of gravitational potential energy and gravitational potential.

During this process, if friction or air resistance is present, additional work is necessary against the frictional force (for the same displacement). This work is strictly path dependent and not recoverable. Such forces (like friction, air drag, etc.) are called *non-conservative* forces. Work done against these forces appears as heat, sound, light, etc. The work done against non-conservative forces is not recoverable even if the path is exactly reversed.

4.5.5. Work Done by a Variable Force:

The popular formula for calculating work done is $W = \vec{F} \cdot \vec{s} = F s \cos \theta$ where θ is the angle between the applied force $\vec{F}$ and the displacement $\vec{s}$.

This formula is applicable only if both force $\vec{F}$ and displacement $\vec{s}$ are constant and finite. In several real-life situations, the force is not constant. For example, while lifting an object through several thousand kilometres, the gravitational force is not constant. The viscous forces like fluid resistance depend upon the speed, hence, quite often are not constant with time. In order to calculate the work done by such variable forces we use integration.

Illustration: Figure 4.1(a) shows variation of a force $\vec{F}$ plotted against corresponding displacements in its direction $\vec{s}$. As the displacement is in the direction of the applied force, vector nature is not used. We need to calculate the work done by this force during

displacement from s_1 to s_2 . As the force is variable, using $W = F(s_2 - s_1)$ directly is not possible. In order to use integration, let us divide the displacement into a large number of infinitesimal (infinitely small) displacements. One of such displacements is ds . It is so small that the force F is practically constant for this displacement. Practically constant means the change in the force is so small that the change can not be recorded. The shaded strip shows one of such displacements. As the force is constant, the area of this strip $F.ds$ is the work done dW for this displacement. Total work done W for displacement $(s_2 - s_1)$ can then be obtained by using integration.

$$\therefore W = \int_{s_1}^{s_2} F.ds$$

Method of integration is applicable if the exact way of variation in $\vec{F}$ with $\vec{s}$ is known and that function is integrable.

The area under the curve between s_1 and s_2 also gives the work done W (if the force axis necessarily starts with zero), as it consists of all the strips of ds between s_1 and s_2 . In Fig. 4.1(b), the variation in the force is linear. In this case, the area of the trapezium AS_1S_2B gives total work done W .

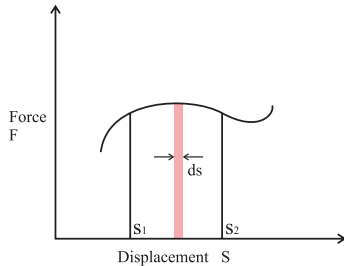


Fig 4.1 (a): Work done by nonlinearly varying force.

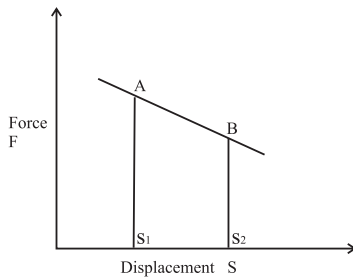


Fig 4.1 (b): Work done by linearly varying force.

Example 4.4: Over a given region, a force (in newton) varies as $F = 3x^2 - 2x + 1$. In this region, an object is displaced from $x_1 = 20$ cm to $x_2 = 40$ cm by the given force. Calculate the amount of work done.

Solution:

$$\begin{aligned} W &= \int_{s_1}^{s_2} F.ds = \int_{0.2}^{0.4} (3x^2 - 2x + 1)dx \\ &= \left[x^3 - x^2 + x \right]_{0.2}^{0.4} \\ &= \left[0.4^3 - 0.4^2 + 0.4 \right] - \left[0.2^3 - 0.2^2 + 0.2 \right] \\ &= 0.304 - 0.168 = 0.136 \text{ J} \end{aligned}$$

4.6. Work Energy Theorem:

If there is a decrease in the potential energy (like a body falling down) due to a conservative force, it is entirely converted into kinetic energy. Work done by the force then appears as kinetic energy. Vice versa if an object is moving against a conservative force its kinetic energy decreases by an amount equal to the work done against the force. This principle is called *work-energy theorem* for conservative forces.

Case I: Consider an object of mass m moving with velocity $\vec{u}$ experiencing a constant opposing force $\vec{F}$ which slows it down to $\vec{v}$ during displacement $\vec{s}$. As $\vec{u}$ and $\vec{F}$ are oppositely directed, the entire motion will be along the same *line*. In this case we need not use the vector form, just $\pm$ signs should be good enough.

If $a = \frac{F}{m}$ is the acceleration, we can write the relevant equation of motion as $v^2 - u^2 = 2(-a)s$ (negative acceleration for opposing force)

Multiplying throughout by $m/2$, we get

$$\frac{1}{2}mu^2 - \frac{1}{2}mv^2 = (ma).s = F.s$$

Left-hand side is decrease in the kinetic energy and the right-hand side is the work done by the force. Thus, change in kinetic energy is equal to work done by the conservative force, which is in accordance with work-energy theorem.

Case II: Accelerating conservative force along

with a retarding non-conservative force:

An object dropped from some point at height h falls down through air. While coming down its potential energy decreases. Equal amount of work is done in this case also. However, this time the work is not entirely converted into kinetic energy but some part of it is used in overcoming the air resistance. This part of energy appears in some other forms such as heat, sound, etc. In this case, the work-energy theorem can mathematically be written as $\Delta PE = \Delta KE + W_{\text{air resistance}}$

(Decrease in the gravitational P.E. = Increase in the kinetic energy + work done against non-conservative forces). Magnitude of air resistance force is not constant but depends upon the speed hence it can be written as $\int \vec{F} \cdot d\vec{s}$ as seen during work done by (or against) a variable force.

4.7. Principle of Conservation of Linear

Momentum:

According to Newton's second law, resultant force is equal to the rate of change of linear momentum or $\vec{F} = \frac{d\vec{p}}{dt}$

In other words, if there is no resultant force, the linear momentum will not change or will remain constant or will be conserved. Mathematically, if $\vec{F} = 0$, $\frac{d\vec{p}}{dt} = 0$ or $\vec{p}$ is constant

Always remember

Isolated system means absence of any external force. A system refers to a set of particles, colliding objects, exploding objects, etc. *Interaction* refers to collision, explosion, etc. During any interaction among such objects the *total* linear momentum of the *entire system* of these particles/objects is constant. Remember, forces during collision or during explosion are *internal* forces for that *entire system*.

During collision of two particles, the two particles exert forces on *each other*. If these particles are discussed independently, these are external forces. However, for the *system of the two particles* together, these forces are internal forces.

(or conserved). This leads us to the *principle of conservation of linear momentum* which can be stated as “**The total momentum of an isolated system is conserved during any interaction.**”

Systems and free body diagrams:

Mathematical approach for application of Newton's second law:

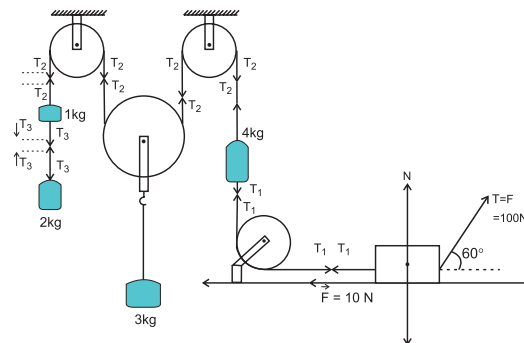


Fig 4.2 (a): System for illustration of free body diagram.

Consider the arrangement shown in Fig.4.2 (a). Pulleys P_1 , P_3 and P_4 are fixed, while P_2 is movable. Force $F = 100$ N, applied at an angle 60° with the horizontal is responsible for the motion, if any. Contact surface of the 5 kg mass offers a constant opposing force $F = 10$ N. Except this, there are no resistive forces anywhere.

Discussion: Until 1 kg mass reaches the pulley P_1 , the motion of 1 kg and 2 kg masses is identical. Thus, these two can be considered to be a *single system* of mass 3 kg except for knowing the tension T_3 . The forces due to tension in the string joining them are internal forces for this system.

All masses *except* the 3 kg mass are travelling same distance in the same time. Thus, their accelerations, if any, have same magnitudes. If the string S connecting 1 kg and 4 kg masses moves by x , the lower string S_1 holding the 3 kg mass moves through $x/2$.

Free body diagrams (FBD): A free body diagram refers to forces acting on only one body at a time, and its acceleration.

Free body diagram of 2 kg mass: Let a be its upward acceleration. According to Newton's second law, it must be due to the resultant vertical force on this mass. To know this force,

we need to know all the individual forces acting on this mass. The agencies exerting forces on this mass are Earth (downward force $1g$) and force due to the tension T_3 .

In this case, the lower half of the string

Practical tip: Easiest way to know the direction of forces due to tension is to put an X-mark on the string. Two halves of this cross indicate the directions of the forces exerted *by the string* on the bodies connected to either parts of the string.

is connected to the 2 kg mass. The direction of T_3 for lower part of the string is upwards as shown in the Fig. 4.2 (b). Upper part of the string is connected to the 1 kg mass. Thus, the direction of T_3 for 1 kg mass will be downwards. However, it will appear only for the free body diagram of the 1 kg mass and will not appear in the free body diagram of 2 kg mass. Hence, the free body diagram of the 2 kg mass will be as follows: Its force equation, according to Newton's second law will then be $T_3 - 2g = 2a$.

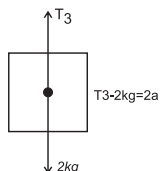


Fig 4.2 (b): Free body diagram for 2 kg mass.

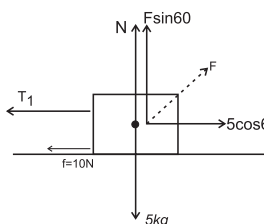


Fig 4.2 (c): Free body diagram for 5 kg mass.

Free body diagram of mass 5 kg: Its horizontal acceleration is also a , but towards right. The force exerting agencies are Earth (force $5g$ downwards), contact surface (normal force N , vertically upwards **and** opposing force $F = 10$ N, towards left), and the two strings on either side (Forces due to their tensions T and T_1). All these are shown in its free body diagrams in Fig. 4.2 (c). On resolving the force F along the vertical and horizontal directions, the free body diagram of 5 kg mass can be drawn as explained below.

As this mass has only horizontal motion,

the vertical forces must cancel. Therefore, along the vertical direction,

$$N + F \sin 60^\circ = 5g$$

Along the horizontal direction,

$$T_1 + 10(\text{opposing force}) = F \cos 60^\circ = T \cos 60^\circ$$

Similar equations can be written for all the masses and also for the movable pulley. On solving these equations simultaneously, we can obtain all the necessary quantities.

Example 4.5: Figure (a) shows a fixed pulley. A massless inextensible string with masses m_1 and $m_2 > m_1$ attached to its two ends is passing over the pulley. Such an arrangement is called an Atwood machine. Calculate accelerations of the masses and force due to the tension along the string assuming axle of the pulley to be frictionless.

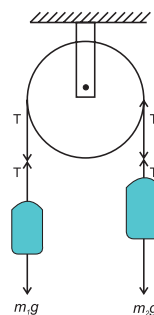


Fig. (a).

Solution: Method I: Direct method: As $m_2 > m_1$, mass m_2 is moving downwards and mass m_1 is moving upwards.

Net downward force

$$= F = (m_2)g - (m_1)g = (m_2 - m_1)g$$

As the string is inextensible, both the masses travel the same distance in the same time. Thus, their accelerations are numerically the same (one upward, other downward). Let it be a .

Thus, total mass in motion, $M = m_2 + m_1$

$$\therefore a = \frac{F}{M} = \left(\frac{m_2 - m_1}{m_2 + m_1} \right) g$$

For mass m_1 , the upward force is the force due to tension T and downward force is mg . It has upward acceleration a . Thus, $T - m_1 g = m a$
 $\therefore T = m_1(g + a)$

Using the expression for a , we get

$$T = \left(\frac{2m_1m_2}{m_1 + m_2} \right) g$$

Method II: (Free body equations)

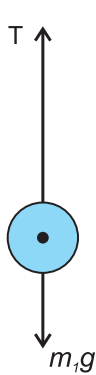


Fig. (b)

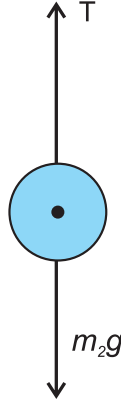


Fig. (b)

Free body diagrams of m_1 and m_2 are as shown in Figs. (b) and (c).

Thus, for the first body, $T - m_1g = m_1a$ --- (I)

For the second body, $m_2g - T = m_2a$ --- (II)

Adding (I) and (II), and solving for a ,

$$\text{we get, } a = \left(\frac{m_2 - m_1}{m_2 + m_1} \right) g \quad \text{--- (III)}$$

Solving Eqs. (II) and (III) for T , we get,

$$T = m_2(g - a) = \left(\frac{2m_1m_2}{m_1 + m_2} \right) g$$

4.8. Collisions:

During collisions a number of objects come together, interact (exert forces on *each other*) and scatter in different directions.

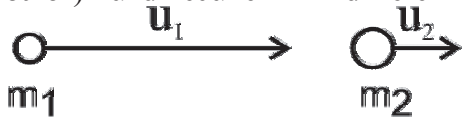


Fig. 4.3 (a): Head on collision-before collision.

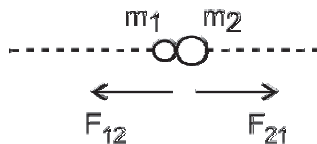


Fig. 4.3 (b): Head on collision-during impact.



Fig. 4.3 (c): Head on collision-after collision.

4.8.1. Elastic and inelastic collisions:

During a collision, the linear momentum of the entire system of particles is always conserved as there is no external force acting on the system of particles. However, the individual momenta of the particles change due to mutual forces, which are internal forces.

$\therefore \sum \vec{p}_{initial} = \sum \vec{p}_{final}$, during **any** collision (or explosion), where $\vec{p}$'s are the linear momenta of the particles.

However, kinetic energy of the entire system may or may not conserve.

Collisions can be of two types: **elastic collisions** and **inelastic collisions**.

Elastic collision: A collision is said to be **elastic** if kinetic energy of the entire system is conserved during the collision (along with the linear momentum). Thus, during an elastic collision,

$$\sum K.E._{initial} = \sum K.E._{final}$$

An elastic collision is *impossible* in daily life. However, in many situations, the interatomic or intermolecular collisions are considered to be elastic (like in kinetic theory of gases, to be discussed in the next standard).

Inelastic Collision: A collision is said to be **inelastic** if there is a loss in the kinetic energy during collision, but linear momentum is conserved. The loss in kinetic energy is either due to internal friction or vibrational motion of atoms causing heating effect. Thus, during an inelastic collision,

$$\sum K.E._{initial} > \sum K.E._{final}$$

During an explosion as energy is supplied internally. Thus,

$$\sum K.E._{final} > \sum K.E._{initial}$$

As stated earlier, $\sum \vec{p}_{initial} = \sum \vec{p}_{final}$ for inelastic collisions or explosion also. In fact, this is always the first equation for discussing these interactions or while solving numerical questions.

4.8.2. Perfectly Inelastic Collision:

This is a special case of inelastic collisions. If colliding bodies join together after collision, it is said to be a **perfectly inelastic collision**.

In other words, the colliding bodies have a common final velocity after a perfectly inelastic collision. Being an inelastic collision, obviously there is a loss in the kinetic energy of the system during a perfectly inelastic collision. In fact, the loss in kinetic energy is maximum in perfectly inelastic collision.

Illustrations:

- (i) Consider a bullet fired towards a block kept on a smooth surface. Collision between bullet and the block will be elastic if the bullet rebounds with exactly the same initial speed and the block remains stationary. If the bullet gets embedded into the block and the two move jointly, it is perfectly inelastic collision. If the bullet rebounds with smaller speed or comes out of the block on the other side with some speed, it is an inelastic (or partially inelastic) collision. Remember, there is nothing called a partially elastic collision. Elastic collisions are always perfectly elastic. An inelastic collision however, may be partially or perfectly inelastic.
- (ii) Visualise a ball dropped from some height on a hard surface, the entire system being in an evacuated space. If the ball rebounds exactly to the same height from where it was dropped, the collision between the ball and the surface (in turn, with the Earth) is elastic. As you know, the ball never reaches the same initial height or a height greater than the initial height. Rebounding to smaller height refers to inelastic collision. Instead of ball, if mud or clay is dropped, it sticks to the surface. This is perfectly inelastic collision.

4.8.3. Coefficient of Restitution e :

For collision of two objects, the negative of ratio of relative velocity of separation to relative velocity of approach is defined as the coefficient of restitution e .

One dimensional or head-on collision: A collision is said to be head-on if the colliding objects move along the same straight line, before and after the collision. Here, we use u_1 , u_2 , v_1 , v_2 as symbols.

Consider such a head-on collision of two bodies of masses m_1 and m_2 with respective initial velocities u_1 and u_2 . As the collision is head on, the colliding masses are along the same line before and after the collision. Hence, vector treatment is not necessary. **(However, velocities must be substituted with proper signs in actual calculation)**. Relative velocity of approach is then $u_a = u_2 - u_1$

Let v_1 and v_2 be their respective velocities after the collision. The relative velocity of recede (or separation) is then $v_s = v_2 - v_1$

$$\begin{aligned} \therefore \text{Coefficient of restitution, } e &= -\frac{v_s}{u_a} \\ &= \frac{-v_2 - v_1}{u_2 - u_1} = \frac{v_1 - v_2}{u_2 - u_1} \\ &= \frac{\text{Relative speed of separation}}{\text{Relative speed of approach}} \quad \text{---(4.1)} \end{aligned}$$

For a perfectly inelastic collision, the colliding bodies move jointly after the collision, i.e., $v_2 = v_1$ or $v_2 - v_1 = 0$. Hence, for a perfectly inelastic collision, $e = 0$. In other words, if $e = 0$, the head-on collision is perfectly inelastic collision.

Coefficient of restitution during a head-on, elastic collision:

Consider the collision described above to be elastic. According to the principle of conservation of linear momentum,

Total initial momentum = Total final momentum.

$$\therefore m_1 u_1 + m_2 u_2 = m_1 v_1 + m_2 v_2 \quad \text{--- (4.2)}$$

$$\therefore m_1 (u_1 - v_1) = m_2 (v_2 - u_2) \quad \text{--- (4.3)}$$

As the collision is elastic, total kinetic energy of the system is also conserved.

$$\therefore \frac{1}{2} m_1 u_1^2 + \frac{1}{2} m_2 u_2^2 = \frac{1}{2} m_1 v_1^2 + \frac{1}{2} m_2 v_2^2 \quad \text{--- (4.4)}$$

$$\therefore m_1 (u_1^2 - v_1^2) = m_2 (v_2^2 - u_2^2)$$

$$\begin{aligned} \therefore m_1 (u_1 + v_1)(u_1 - v_1) \\ = m_2 (u_2 + v_2)(v_2 - u_2) \quad \text{--- (4.5)} \end{aligned}$$

Dividing Eq. (4.5) by Eq. (4.3), we get

$$u_1 + v_1 = u_2 + v_2 \quad \text{--- (4.6)}$$

$$\therefore u_2 - u_1 = v_1 - v_2$$

For an elastic collision,

$$e = \frac{v_1 - v_2}{u_2 - u_1} = 1 \quad \text{--- (4.7)}$$

Thus, for an elastic collision, coefficient of restitution, $e = 1$. For a perfectly inelastic collision, $e = 0$ (by definition). Thus, for any collision, the coefficient of restitution lies between 1 and 0.

Above expressions (Eq. (4.1) Eq. (4.7)) are general. While substituting the values of u_1 , u_2 , v_1 , v_2 , their algebraic values must be used in actual calculation.

For example referring to the Fig. 4.3 (a), (b) and (c)

Eq (4.1) gives

$$e = -\frac{v_s}{u_a}$$

Here $u_a = u_1 - u_2$ since $u_1 > u_2$ and $v_s = v_1 + v_2$ since the objects go in opposite directions.

$$\therefore e = -\frac{v_1 + v_2}{u_1 - u_2} \quad \text{--- (a)}$$

Using Eq (4.6),

$$u_1 + v_1 = u_2 + v_2$$

$\therefore$ According to Fig. 4.3,

$$u_1 - v_1 = u_2 + v_2$$

$$\therefore v_1 + v_2 = u_2 - u_1 \quad \text{--- (b)}$$

By substituting in (a),

$$e = -\frac{(u_2 - u_1)}{(u_1 - u_2)} = 1,$$

which is the case of a perfectly elastic collision.

4.8.4. Expressions for final velocities after a head-on, elastic collision:

From Eq. (4.6), $v_2 = u_1 + v_1 - u_2$

Using this in Eq. (4.2), we get

$$m_1 u_1 + m_2 u_2 = m_1 v_1 + m_2 (u_1 + v_1 - u_2)$$

$$\therefore v_1 = \left(\frac{m_1 - m_2}{m_1 + m_2} \right) u_1 + \left(\frac{2m_2}{m_1 + m_2} \right) u_2 \quad \text{--- (4.8)}$$

Subscripts 1 and 2 were arbitrarily chosen. Thus, just interchanging 1 with 2 gives us v_2 as

$$v_2 = \left(\frac{m_2 - m_1}{m_1 + m_2} \right) u_2 + \left(\frac{2m_1}{m_1 + m_2} \right) u_1 \quad \text{--- (4.9)}$$

Equation (4.9) can also be obtained by substituting v_1 from Eq. (4.8) in Eq. (4.6).

Particular cases:

(i) If the bodies are of equal masses (or identical), $m_1 = m_2$, Eqs. (4.8) and (4.9) give

$$v_1 = u_2 \text{ and } v_2 = u_1.$$

Thus, the bodies just exchange their velocities.

(ii) If colliding body is much heavier and the struck body is initially at rest, i.e.,

$$m_1 \gg m_2 \text{ and } u_2 = 0,$$

we can use

$$m_1 \pm m_2 \cong m_1 \text{ and } \frac{m_2}{m_1 + m_2} \cong 0$$

$\therefore v_1 \cong u_1$ and $v_2 \cong 2u_1$, i.e., the massive striking body is practically unaffected and the tiny body which is struck, travels with double the speed of the massive striking body.

(iii) The body which is struck is much heavier than the colliding body and is initially at rest, i.e., $m_1 \ll m_2$ and $u_2 = 0$.

Using similar approximations, we get, $v_1 \cong -u_1$ and $v_2 \cong 0$, i.e., the tiny (lighter) object rebounds with same speed while the massive object is unaffected. This is as good as dropping an elastic object on hard surface of the Earth.



Do you know ?

Are you aware of elasticity of materials? Is there any connection between elasticity of materials and elastic collisions?

Example 4.6: One marble collides head-on with another identical marble at rest. If the collision is partially inelastic, determine the ratio of their final velocities in terms of coefficient of restitution e .

Solution: According to conservation of momentum, $m_1 u_1 + m_2 u_2 = m_1 v_1 + m_2 v_2$

As $m_1 = m_2$, we get, $u_1 + u_2 = v_1 + v_2$

$\therefore$ If $u_2 = 0$, we get, $v_1 + v_2 = u_1 \dots$ (I)

Coefficient of restitution,

$$e = \frac{v_2 - v_1}{u_1 - u_2} \therefore v_2 - v_1 = eu_1 \dots \text{(II)}$$

Dividing Eq. (I) by Eq. (II),

$$\frac{v_1 + v_2}{v_2 - v_1} = \frac{1}{e}$$

Using componendo and dividendo, we get,

$$\frac{v_2}{v_1} = \frac{1+e}{1-e}$$

4.8.5. Loss in the kinetic energy during a perfectly inelastic head-on collision:

Consider a perfectly inelastic, head on collision of two bodies of masses m_1 and m_2 with respective initial velocities u_1 and u_2 . As the collision is perfectly inelastic, they move jointly after the collision, i.e., their final velocity is the same. Let it be v .

According to conservation of linear momentum, $m_1 u_1 + m_2 u_2 = (m_1 + m_2) v$

$$\therefore v = \frac{m_1 u_1 + m_2 u_2}{m_1 + m_2} \quad \text{--- (4.10)}$$

This is the common velocity after a perfectly inelastic collision

Loss in K.E. = Δ (K.E.)

= Total initial K.E. - Total final K.E.

$$\begin{aligned} \therefore \Delta(\text{K.E.}) &= \frac{1}{2} m_1 u_1^2 + \frac{1}{2} m_2 u_2^2 - \frac{1}{2} (m_1 + m_2) v^2 \\ &= \frac{1}{2} m_1 u_1^2 + \frac{1}{2} m_2 u_2^2 - \frac{1}{2} (m_1 + m_2) \left(\frac{m_1 u_1 + m_2 u_2}{m_1 + m_2} \right)^2 \\ \therefore \Delta(\text{K.E.}) &= \frac{1}{2} m_1 u_1^2 + \frac{1}{2} m_2 u_2^2 - \frac{1}{2} \frac{(m_1 u_1 + m_2 u_2)^2}{(m_1 + m_2)} \end{aligned}$$

On simplifying, we get,

$$\Delta(\text{K.E.}) = \frac{1}{2} \left(\frac{m_1 m_2}{m_1 + m_2} \right) (u_1 - u_2)^2 \quad \text{--- (4.11)}$$

Masses are always positive and $(u_1 - u_2)^2$ is also positive. Hence, there is always a loss in the kinetic energy in a perfectly inelastic collision.

Final velocities and loss in K.E. in an inelastic head-on collision:

If e is the coefficient of restitution, using Eq. (4.2), the expressions for final velocities after an inelastic collision can be derived as

$$\begin{aligned} v_1 &= \left(\frac{m_1 - em_2}{m_1 + m_2} \right) u_1 + \left(\frac{[1+e]m_2}{m_1 + m_2} \right) u_2 \\ &= \frac{em_2(u_2 - u_1) + m_1 u_1 + m_2 u_2}{m_1 + m_2} \quad \text{and} \\ v_2 &= \left(\frac{m_2 - em_1}{m_1 + m_2} \right) u_2 + \left(\frac{[1+e]m_1}{m_1 + m_2} \right) u_1 \\ &= \frac{em_1(u_1 - u_2) + m_1 u_1 + m_2 u_2}{m_1 + m_2} \end{aligned}$$

Loss in the kinetic energy is given by

$$\Delta(\text{K.E.}) = \frac{1}{2} \left(\frac{m_1 m_2}{m_1 + m_2} \right) (u_1 - u_2)^2 (1 - e^2)$$

As $e < 1$, $(1 - e^2)$ is always positive. Thus, there is always a loss of K.E. in an inelastic collision. Also, for a perfectly inelastic collision, $e = 0$. Hence, in this case, the loss is maximum.

Using $e = 1$, these equations lead us to an elastic collision and for $e = 0$ they lead us to a perfectly inelastic collision. Verify that they give the same expressions that are derived earlier.

The quantity $\mu = \frac{m_1 m_2}{m_1 + m_2}$ is the reduced mass of the system.

Impulse or change in momenta of the bodies:

During collision, the linear momentum delivered by first body (particle) to the second body must be equal to change in momentum or *impulse* of the second body, and vice versa.

$$\begin{aligned} \therefore \text{Impulse, } |J| &= |\Delta p_1| = |\Delta p_2| \\ &= |m_1 v_1 - m_1 u_1| = |m_2 v_2 - m_2 u_2| \end{aligned}$$

On substituting the values of v_1 and v_2 and solving, we get

$$\begin{aligned} |J| &= \left(\frac{m_1 m_2}{m_1 + m_2} \right) (1 + e) (|u_1 - u_2|) \\ &= \mu (1 + e) u_{\text{relative}} \end{aligned}$$

$$|u_1 - u_2| = u_{\text{relative}} = \text{velocity of approach}$$

4.8.6. Collision in two dimensions, i.e., a nonhead-on collision:

In this case, the direction of at least one initial velocity is NOT along the line of impact. In order to discuss such collisions mathematically, it is convenient to use two mutually perpendicular directions as shown in Fig. 4.4. One of them is the common tangent at the point of impact, along which there is no force (or along this direction, there is no change in momentum). The other direction is perpendicular to this common tangent through the point of impact, in the two-dimensional plane of initial and final velocities. This is called the line of impact. Internal mutual forces exerted during impact, which are responsible for change in the momenta, are acting along this line. From Fig. (4.4), $\vec{u}_1$ and $\vec{u}_2$, initial velocities make angles α_1 and α_2 respectively with the line of impact while $\vec{v}_1$ and $\vec{v}_2$, final velocities make angles β_1 and β_2 respectively with the line of impact.

According to conservation of linear momentum along the line of contact,

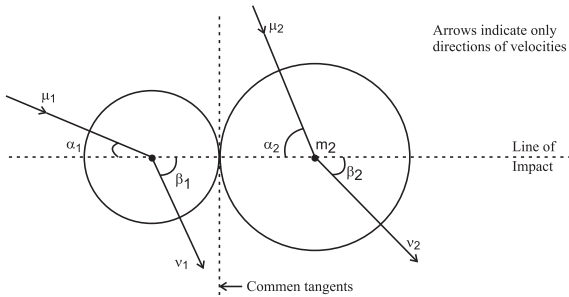


Fig. 4.4: Oblique or non head-on collision.

$$m_1 u_1 \cos \alpha_1 + m_2 u_2 \cos \alpha_2 = m_1 v_1 \cos \beta_1 + m_2 v_2 \cos \beta_2 \quad \text{--- (4.12)}$$

As there is no force along the common tangent (perpendicular to line of impact),

$$m_1 u_1 \sin \alpha_1 = m_1 v_1 \sin \beta_1 \quad \text{--- (4.13)}$$

$$\text{and } m_2 u_2 \sin \alpha_2 = m_2 v_2 \sin \beta_2 \quad \text{--- (4.14)}$$

For coefficient of restitution, along the line of impact,

$$e = - \left(\frac{v_2 \cos \beta_2 - v_1 \cos \beta_1}{u_2 \cos \alpha_2 - u_1 \cos \alpha_1} \right) = \frac{v_2 \cos \beta_2 - v_1 \cos \beta_1}{u_1 \cos \alpha_1 - u_2 \cos \alpha_2} \quad \text{--- (4.15)}$$

Equations (4.12), (4.13), (4.14) and (4.15) are to be solved for the four unknowns v_1 , v_2 , β_1 and β_2 .

Magnitude of the impulse, along the line of impact,

$$|J| = \left(\frac{m_1 m_2}{m_1 + m_2} \right) (1 + e) (|u_1 \cos \alpha_1 - u_2 \cos \alpha_2|) = \mu (1 + e) u_{\text{relative}}$$

along line of impact.

Loss in the kinetic energy = Δ (K.E.)

$$\frac{1}{2} \left(\frac{m_1 m_2}{m_1 + m_2} \right) (u_1 \cos \alpha_1 - u_2 \cos \alpha_2)^2 (1 - e^2) = \frac{1}{2} \mu (u_{\text{relative}}^2) (1 - e^2)$$

Example 4.7: A shell of mass 3 kg is dropped from some height. After falling freely for 2 seconds, it explodes into two fragments of masses 2 kg and 1 kg. Kinetic energy provided by the explosion is 300 J. Using $g = 10 \text{ m/s}^2$, calculate velocities of the fragments. Justify your answer if you have more than one options.

Solution: $m_1 + m_2 = 3 \text{ kg}$.

After falling freely for 2 seconds,

$$v = u + at = 0 + 10(2) = 20 \text{ m s}^{-1} = u_1 = u_2$$

According to conservation of linear momentum, $m_1 u_1 + m_2 u_2 = m_1 v_1 + m_2 v_2$

$$\therefore 3 \times 20 = 2v_1 + 1v_2 \therefore v_2 = 60 - 2v_1 \quad \text{--- (I)}$$

K.E. provided = Final K.E. - Initial

$$\text{K.E.} = \frac{1}{2} m_1 v_1^2 + \frac{1}{2} m_2 v_2^2 - \frac{1}{2} (m_1 + m_2) u^2$$

$$\therefore \frac{1}{2} 2v_1^2 + \frac{1}{2} v_2^2 - \frac{1}{2} 3(20)^2 = 300 \text{ J}$$

$$\text{or } 2v_1^2 + v_2^2 = 1800$$

$$\text{or } 2v_1^2 + (60 - 2v_1)^2 = 1800 \text{ using Eq. (I)}$$

$$\therefore 3600 - 240v_1 + 6v_1^2 = 1800$$

$$\therefore v_1^2 - 40v_1 + 300 = 0$$

$$\therefore v_1 = 30 \text{ m s}^{-1} \text{ or } 10 \text{ m s}^{-1} \text{ and}$$

$$v_2 = 0 \text{ or } 40 \text{ m s}^{-1}$$

There are two possible answers since the positions of two fragments can be different as explained below.

If $v_1 = 30 \text{ m s}^{-1}$ and $v_2 = 0$, lighter fragment 2 should be above. On the other hand, if $v_1 = 10 \text{ m s}^{-1}$ and $v_2 = 40 \text{ m s}^{-1}$, lighter fragment 2 should be below, both moving downwards.

Example 4.8: Bullets of mass 40 g each, are fired from a machine gun at a rate of 5 per second towards a firmly fixed hard surface of area 10 cm^2 . Each bullet hits normal to the surface at 400 m/s and rebounds in such a way that the coefficient of restitution for the collision between bullet and the surface is 0.75. Calculate average force and average pressure experienced by the surface due to this firing.

Solution: For the collision,

$$u_1 = 400 \text{ m s}^{-1}, e = 0.75, v_1 = ?$$

For the firmly fixed hard surface, $u_2 = v_2 = 0$

$$e = 0.75 = \frac{v_1 - v_2}{u_2 - u_1} = \frac{v_1 - 0}{0 - 400} \therefore v_1 = -300 \text{ m/s.}$$

-ve sign indicates that the bullet rebounds in exactly opposite direction.

Change in momentum of each bullet
 $= m(v_1 - u_1)$

Equal and opposite will be the momentum transferred to the surface, per collision.

$\therefore$ Momentum transferred to the surface, per collision

$$p = m(u_1 - v_1) = 0.04(400 - [-300]) = 28 \text{ N s}$$

The rate of collision is same as rate of firing.

$\therefore$ Momentum received by the surface per second, $\frac{dp}{dt} = 28 \times 5 = 140 \text{ N}$

This must be the average force experienced by the surface of area $A = 10 \text{ cm}^2 = 10^{-3} \text{ m}^2$

$\therefore$ Average pressure experienced,

$$P = \frac{F}{A} = \frac{140}{10^{-3}} = 1.4 \times 10^5 \text{ N m}^{-2}$$

≈ 1.4 times the atmospheric pressure.

4.9. Impulse of a force:

According to Newton's first law of motion, any unbalanced force changes linear momentum of the system, i.e., basic effect of an unbalanced force is to change the momentum.

According to Newton's second law of motion, $\vec{F} = \frac{d\vec{p}}{dt}$

$$\therefore d\vec{p} = \vec{F}.dt$$

The quantity 'change in momentum' is separately named as *Impulse of the force* $\vec{J}$.

If the force is constant, and is acting for a finite and measurable time, we can write

The change in momentum in time t

$$\vec{J} = d\vec{p} = \vec{p}_2 - \vec{p}_1 = \vec{F}.t \quad \text{---(4.16)}$$

For a given body of mass m , it becomes

$$\vec{J} = \vec{p}_2 - \vec{p}_1 = m(\vec{v}_2 - \vec{v}_1) = \vec{F}.t \quad \text{---(4.17)}$$

If $\vec{F}$ is not constant but we know how it varies with time, then

$$\vec{J} = \Delta\vec{p} = \int d\vec{p} = \int \vec{F}.dt \quad \text{--- (4.18)}$$

Always remember:

- 1) Colliding objects experience forces along the line of impact which changes their momenta. For their system, these forces are internal forces. These forces form an action-reaction pair, which are equal and opposite, and act on different objects.
- 2) There is no force along the common tangent, i.e., perpendicular to the line of contact.
- 3) In reality, the impact is followed by emission of sound and heat and occasionally light. Thus, in general, part of mechanical energy- kinetic energy - is lost (i.e., converted into some other non-recoverable forms). However, total energy of the system is conserved.
- 4) In reality, velocity of separation (relative final velocity) is less than velocity of approach (relative initial velocity) along the line of impact. Thus coefficient of restitution $e < 1$.
- 5) Only during elastic collisions (atomic and molecular level only, never possible in real life), the kinetic energy is conserved and the velocity of separation is equal to the velocity of approach or the initial relative velocity is equal to the final relative velocity.

4.9.1. Necessity of defining impulse:

As discussed above, if a force is constant over a given interval of time or if we know how it varies with time, we can calculate the corresponding change in momentum directly by multiplying the force and time.

However, in many cases, an appreciable force acts for an extremely small interval of

time (too small to measure the force and the time independently). However, change in the momentum due to this force is noticeable and can be measured. This change is defined as *impulse* of the force.

Real life illustrations: While (i) hitting a ball with a bat, (ii) giving a kick to a foot-ball, (iii) hammering a nail, (iv) bouncing a ball from a hard surface, etc., appreciable amount of force is being exerted. In such cases the time for which these forces act on respective objects is negligibly small, mostly not easily recordable. However, the effect of this force is a recordable change in the momentum of that object. Thus, it is convenient to define the change in momentum itself as a physical quantity.

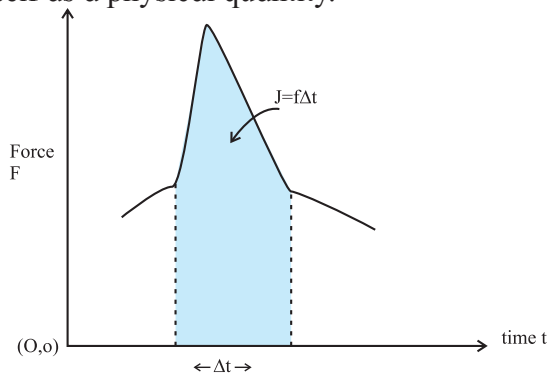


Fig. 4.5: Graphical representation of impulse of a force.

Figure 4.5 shows variation of a force as a function of time e.g., for a collision between bat and ball with the *force axis starting with zero*. The shaded area or the area under the curve gives the product of force against corresponding time (in this case, Δt), hence gives the impulse. For a constant force it is obviously a rectangle. Generally, force is zero before the impact, rises to a maximum and decreases to zero after the impact. For softer tennis ball, the collision time is larger and the maximum force is less. The area under the ($F - t$) graph is the same. Wicket keeper eases off (by increasing the time of collision) while catching a fast ball. As mentioned earlier, it is absolutely necessary that the force axis must start from zero.

Recall from Chapter 3, analogues concepts using area under a curve are (i) Obtaining displacement in a given time interval as area under the curve for $v-t$ graph, with zero origin

for velocity axis. (ii) Obtaining work done by a force as the area under the curve for $F-s$ graph, with zero origin for force axis.

Example 4.9: Mass of an Oxygen molecule is 5.35×10^{-26} kg and that of a Nitrogen molecule is 4.65×10^{-26} kg. During their Brownian motion (random motion) in air, an Oxygen molecule travelling with a velocity of 400 m/s collides elastically with a nitrogen molecule travelling with a velocity of 500 m/s in the exactly opposite direction. Calculate the impulse received by each of them during collision. Assuming that the collision lasts for 1 ms, how much is the average force experienced by each molecule?

Let

$$m_1 = m_O = 5.35 \times 10^{-26} \text{ kg},$$

$$m_2 = m_N = 4.65 \times 10^{-26} \text{ kg}.$$

$$\therefore u_1 = 400 \text{ m s}^{-1} \text{ and } u_2 = -500 \text{ m s}^{-1}$$

taking direction of motion of Oxygen molecule as the positive direction.

For an elastic collision,

$$v_1 = \left(\frac{m_1 - m_2}{m_1 + m_2} \right) u_1 + \left(\frac{2m_2}{m_1 + m_2} \right) u_2 \quad \text{and}$$

$$v_2 = \left(\frac{m_2 - m_1}{m_1 + m_2} \right) u_2 + \left(\frac{2m_1}{m_1 + m_2} \right) u_1$$

$$\therefore v_1 = -437 \text{ m s}^{-1} \text{ and } v_2 = 463 \text{ m s}^{-1}$$

$$\therefore J_O = m_O (v_1 - u_1) = -4.478 \times 10^{-23} \text{ N s},$$

$$J_N = m_N (v_2 - u_2) = +4.478 \times 10^{-23} \text{ N s}$$

As expected, the net impulse or net change in momentum is zero.

$$F_{ON} = \frac{dp_O}{dt} = \frac{J_O}{\Delta t} = \frac{-4.478 \times 10^{-23}}{10^{-3}} \\ = -4.478 \times 10^{-20} \text{ N}$$

$$\text{and } F_{NO} = -F_{ON} = 4.478 \times 10^{-20} \text{ N}$$

4.10. Rotational analogue of a force - moment of a force or torque:

While opening a door fixed to a frame on hinges, we apply the force away from the hinges and perpendicular to the door to open it with ease. In this case we are interested in achieving some angular displacement for the door. If the force is applied near the hinges or

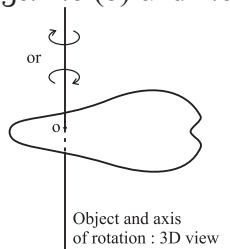
nearly parallel to the door, it is very difficult to open the door. Similarly, if the door is heavier (made up of iron instead of wood or plastic), we need to apply proportionally larger force for the same angular displacement.

It shows that rotational ability of a force not only depends upon the mass (greater force for greater mass), but also upon the point of application of the force (the point should be as away as possible from the axis of rotation) and the angle between direction of the force and the line joining the axis of rotation with the point of application (effect is maximum, if this angle is 90°).

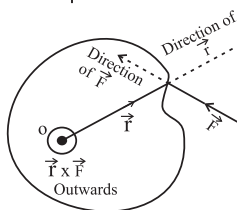
Taking into account all these factors, the quantity **moment of a force** or **torque** is defined as the rotational analogue of a force. As rotation refers to direction (sense of rotation), torque must be a vector quantity. In its mathematical form, torque or moment of a force is given by

$$\vec{\tau} = \vec{r} \times \vec{F} \quad \text{--- (4.17)}$$

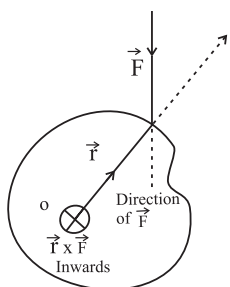
where $\vec{F}$ is the applied force and $\vec{r}$ is the position vector of the point of application of the force from the axis of rotation, as shown in the Figs. 4.6 (b) and 4.6 (c).



Figs. 4.6(a): Illustration of moment of force with object and axis of rotation in 3D view.



Figs. 4.6(b): Top view for moment of force $\vec{F}$ in anticlockwise rotation with $\vec{F}$ and $\vec{r}$ in the plane of paper.



Figs. 4.6(c): Top view of moment of force $\vec{F}$ in clockwise rotation $\vec{F}$ and $\vec{r}$ in the plane of paper.

Figures 4.6(a), 4.6(b) and 4.6(c) illustrate

the directions involved. Figure 4.6(a) is a 3D drawing indicating the laminar (plane or two dimensional) object rotating about a (*fixed*) axis of rotation AOB, the axis being perpendicular to the object and passing through it. Figure 4.6(b) indicates the top view of the object when the rotation is in anticlockwise direction and Fig. 4.6(c) shows the view from the top, if rotation is in clockwise direction. (In fact, Figs. 4.6(b) and 4.6(c) are drawn in such a way that the applied force $\vec{F}$ and position vector $\vec{r}$ of the point of application of the force are in the plane of these figures). Direction of the torque is always perpendicular to the plane containing the vectors $\vec{r}$ and $\vec{F}$ and can be obtained from the rule of cross product or by using the right-hand thumb rule. In Fig. 4.6(b), it is perpendicular to the plane of the figure (in this case, perpendicular to the body) and outwards, i.e., coming out of the paper while in the Fig. 4.6(c), it is inwards, i.e., going into the paper.

In order to indicate the directions which are not in the plane of figure, we use a special convention: $\odot$ for perpendicular to the plane of figure and outwards and $\otimes$ for perpendicular to the plane of figure and inwards.

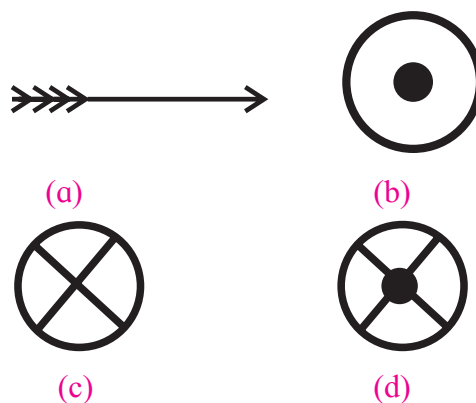


Fig. 4.7: Convention of pictorial representation of vectors as shown in (a) acting in a direction perpendicular to the plane of paper (b) coming out of paper, (c) going in to the paper and (d) perpendicular to the plane of paper.

This convention depends upon a traditional arrow shown in Fig. 4.7 (a). Consider yourself, looking towards the figure from the top. If this arrow approaches you, the tip of the arrow will be prominently seen. Hence circle with a

dot in it [Fig. 4.7 (b)] refers to perpendicular and outwards (or towards you). When you are leaving an arrow, i.e., if an arrow is going away from you, the feathers like a cross will be seen. Hence, a circle with a cross [Fig. 4.6 (c)] indicates perpendicular and inwards (or away from you). Circle with cross and dot indicates a line perpendicular to the plane of figure [Fig. 4.6 (d)].

Magnitude of torque, $\tau = r F \sin \theta$ --- (4.18)
where θ is the smaller angle between the directions of $\vec{r}$ and $\vec{F}$.

Consequences: (i) If r or F is greater, the torque (hence the rotational effect) is greater. Thus, it is recommended to apply the force away from the hinges.

(ii) If $\theta = 90^\circ$, $\tau = \tau_{\max} = rF$. Thus, the force should be applied along normal direction for easy rotation.

(iii) If $\theta = 0^\circ$ or 180° , $\tau = \tau_{\min} = 0$. Thus, if the force is applied parallel or anti-parallel to $\vec{r}$, there is no rotation.

(iv) Moment of a force depends not only on the magnitude and direction of the force, but also on the point where the force acts with respect to the axis of rotation. Same force can have different torque as per its point of application.

4.11. Couple and its torque:

In the discussion of the torque given above, we had considered rotation of the body about a fixed axis and due to a single force. In real life, quite often we apply two equal and opposite forces acting along different lines of action in order to cause rotation. Common illustrations are turning a bicycle handle, turning the steering wheel, opening a common water tap, opening the lid of a bottle (rotation type), etc. Such a pair of forces consisting of two forces of equal magnitude acting in opposite directions along different lines of action is called a couple. It is used to realise a *purely* rotational motion. Moment of a couple or rotational effect of a couple is also called a *torque*.

It may be noted that in the discussion of rotation of a body about a fixed axis due to a single force, there is a reaction force at the fixed axis. Hence, for rotation one always needs

two forces acting in opposite direction along different lines of action.

Torque or Moment of a couple: Figure 4.8 shows a couple consisting of two forces $\vec{F}_1$ and $\vec{F}_2$ of equal magnitudes and opposite directions acting along different lines of action separated by a distance r . Corresponding position vectors should now be defined with reference to the lines of action of forces. Position vector of *any* point on the line of action of force $\vec{F}_1$ from the line of action of force $\vec{F}_2$ is $\vec{r}_{12}$. Similarly, the position vector of *any* point on the line of action of force $\vec{F}_2$ from the line of action of force $\vec{F}_1$ is $\vec{r}_{21}$. Torque or moment of the couple is then given mathematically as

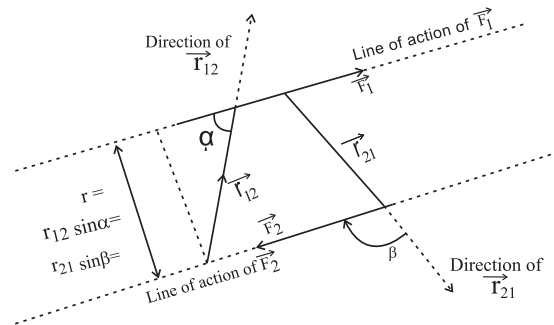


Fig. 4.8: Torque of a couple.

$$\vec{\tau} = \vec{r}_{12} \times \vec{F}_1 = \vec{r}_{21} \times \vec{F}_2 \quad \text{--- (4.19)}$$

From the figure, it is clear that $r_{12} \sin \alpha = r_{21} \sin \beta = r$

If $|\vec{F}_1| = |\vec{F}_2| = F$, the magnitude of torque is given by

$$\tau = r_{12} F_1 \sin \alpha = r_{21} F_2 \sin \beta = r F \quad \text{--- (4.20)}$$

It clearly shows that the torque corresponding to a given couple, i.e., the **moment of a given couple is constant**, i.e., **it is independent of the points of application of forces or the position of the axis of rotation**, but depends only upon magnitude of either force and the separation between their lines of action.

The direction of moment of couple can be obtained by using the vector formula of the torque or by using the right-hand thumb rule. For the couple shown in the Fig. 4.8, it is perpendicular to the plane of the figure and inwards. For a given pair of forces, the direction of the torque is fixed.

In many situations the word *couple* is used synonymous to moment of the couple or its torque, i.e., every time we may not say it as

torque due to the couple, but say that a *couple* is acting.

	Moment of a force	Moment of a couple
1	$\vec{\tau} = \vec{r} \times \vec{f}$	$\vec{\tau} = \vec{r}_{12} \times \vec{F}_1 = \vec{r}_{21} \times \vec{F}_2$
2	$\vec{\tau}$ depends upon the axis of rotation and the point of application of the force.	$\vec{\tau}$ depends only upon the two forces, i.e., it is independent of the axis of rotation or the points of application of forces.
3	It can produce translational acceleration also, if the axis of rotation is not fixed or if friction is not enough.	Does not produce any translational acceleration, but produces only rotational or angular acceleration.
4	Its rotational effect can be balanced by a proper single force or by a proper couple.	Its rotational effect can be balanced only by another couple of equal and opposite torque.

4.11.1. To prove that the moment of a couple is independent of the axis of rotation:

Figure 4.9 shows a rectangular sheet (any object would do) free to rotate only about a fixed axis of rotation, perpendicular to the plane of figure, as shown. A couple consisting of forces $\vec{F}$ and $-\vec{F}$ is acting on the sheet at different locations.

Here we are considering the torque of a couple to be two torques due to individual forces causing rotation about the axis of rotation. In Fig. 4.9(a), the axis of rotation is between the lines of action of the two forces constituting the couple. Perpendicular distances of the axis of rotation from the forces $\vec{F}$ and $-\vec{F}$ are x and y respectively. Rotation due to the pair of forces *in this case* is anticlockwise (from top view), i.e., directions of individual torques due to the two forces are the same.

$$\therefore \tau = \tau_+ + \tau_- = xF + yF = (x + y)F = rF \quad (4.21)$$

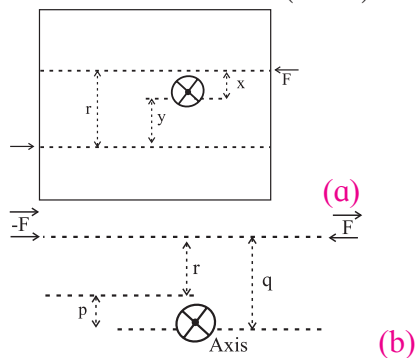


Fig. 4.9: Same couple on same object with fixed axis of rotation at different locations in (a) and (b).

In the Fig. 4.9 (b), lines of action of both the forces are on the same side of the axis of rotation. Thus, *in this case*, the rotation of $+\vec{F}$ is anticlockwise, while that of $-\vec{F}$ is clockwise (from the top view). As a result, their individual torques are oppositely directed. Perpendicular distance of the forces F and $-F$ from the axis of rotation are q and p respectively.

$$\begin{aligned} \therefore \tau &= \tau_+ - \tau_- = qF - pF \\ &= (q - p)F = rF \end{aligned} \quad \text{--- (4.22)}$$

From equations (4.21) and (4.22), it is clear that the torque of a couple is independent of the axis of rotation.

4.12. Mechanical equilibrium:

As a consequence of Newton's second law, the momentum of a system is constant in the absence of an external unbalanced force. This state is called *mechanical equilibrium*.

A particle is said to be in mechanical equilibrium, if no net force is acting upon it. For a system of bodies to be in mechanical equilibrium, the net force acting on *any part* of the system should be zero. In other words, velocity or linear momentum of all parts of the system must be constant or (zero) for the system to be in mechanical equilibrium. Also, there is no acceleration in any part of the system.

Mathematically, $\sum \vec{F} = 0$, for any part of the system for mechanical equilibrium.

4.12.1 Stable, unstable and neutral equilibrium:

Figures 4.10 (a), (b) and (c) show a ball at rest in three situations under the action of balanced forces. In all these cases, it is under equilibrium. However, potential energy-wise, the three cases differ.

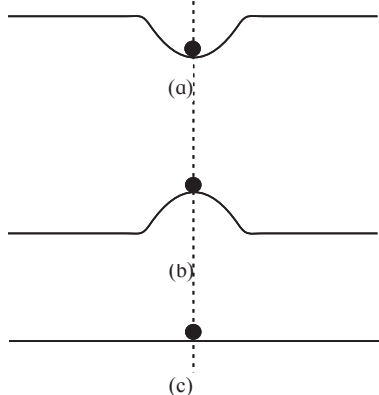


Fig. 4.10: states of mechanical equilibrium (a) stable, (b) unstable and (c) neutral.

Stable equilibrium: In Fig. 4.9(a), the ball is most stable and is said to be in *stable* equilibrium. If it is disturbed slightly from its equilibrium position and released, it tends to recover its position. In this case, potential energy of the system is at its local minimum.

Unstable equilibrium: In Fig. 4.9(b), the ball is said to be in *unstable* equilibrium. If it is slightly disturbed from its equilibrium position, it moves farther from that position. This happens because initially, potential energy of the system is at its local maximum. If disturbed, it tries to achieve the configuration of minimum potential energy.

If potential energy function is known for the system, mathematically, the three equilibria can be explained with the help of derivatives of that function. At any equilibrium position, the first derivative of the potential energy function is zero $\left(\frac{dU}{dx} = 0\right)$.

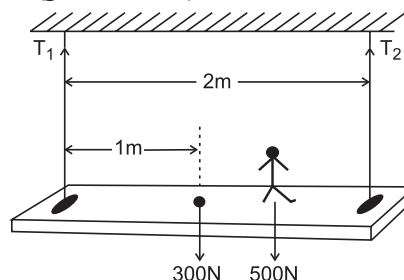
The sign of the second derivative $\left(\frac{d^2U}{dx^2}\right)$

decides the type of equilibrium. It is positive

at stable equilibrium (or vice versa), negative at unstable equilibrium and zero (or does not exist) at neutral equilibrium configuration.

Neutral equilibrium: In Fig. 4.9(c), potential energy of the system is constant over a plane and remains same at any position. Thus, even if the ball is disturbed, it still remains in equilibrium at practically any position. This is described as *neutral* equilibrium.

Example 4.10: A uniform wooden plank of mass 30 kg is supported symmetrically by two light identical cables; each can sustain a tension up to 500 N. After tying, the cables are exactly vertical and are separated by 2 m. A boy of mass 50 kg, standing at the centre of the plank, is interested in walking on the plank. How far can he walk? ($g = 10 \text{ ms}^{-2}$)



Solution: Let T_1 and T_2 be the tensions along the cables, both acting vertically upwards.

Weight of the plank 300 N is acting vertically downwards through the centre, 1 m from either cable. Weight of the boy, 500 N is vertically downwards at the point where he is standing.

$$\therefore T_1 + T_2 = 300 + 500 = 800 \text{ N}$$

Suppose that the boy is able to walk x m towards the right. Obviously, the tension in the right side cable goes on increasing as he walks towards the cable.

Moments of 300 N and 500 N forces about left end A are clockwise, while that of T_2 is anticlockwise.

As the cable can sustain 500 N, $(T_2)_{\max} = 500 \text{ N}$

Thus, for the equilibrium about A, we can write,

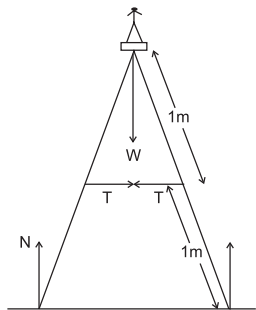
$$300 \times 1 + 500 \times (1 + x) = 500 \times 2 \quad \therefore x = 0.4 \text{ m}$$

Thus, the boy can walk up to 40 cm on either side of the centre.

Example 4.11: A ladder of negligible mass having a cross bar is resting on a frictionless horizontal floor with angle between its legs to be 40° . Each leg is 1 m long. Calculate the

force experienced by the cross bar when a person of mass 50 kg is standing on the ladder. ($g = 10 \text{ m s}^{-2}$)

Solution: Tension T along the cross bar is horizontal. Let L be the length of each leg, which is 1 m.



As there is no friction, there is no horizontal reaction at the floor. Reaction N given by the floor at the base of the ladder will then be only vertical. Thus, along the vertical, two such reactions balance weight $W = mg$ of the person.

$$\therefore N = \frac{mg}{2} = 250 \text{ N}$$

At the left leg, about the upper end, the torque due to N is clockwise and that due to the tension T is anticlockwise. For equilibrium, these two torques should have same magnitude.

$$\therefore N \cdot L \sin 20^\circ = T \cdot \left(\frac{L}{2} \right) \cos 20^\circ$$

$$\therefore T = 2N \tan 20^\circ = 2 \times 250 \times 0.364 = 182 \text{ N}$$

4.13. Centre of mass:

As discussed earlier, Newton's laws of motion and many other laws are applicable for point masses only. However, in real life, we always come across finite objects (objects of measurable sizes). Concept of centre of mass (c.m.) helps us in considering these objects to be point objects at a particular location, thereby allowing us to apply Newton's laws of motion.

4.13.1. Mathematical understanding of centre of mass:

(i) System of n particles: Consider a system of n particles of masses $m_1, m_2 \dots m_n$.

$$\therefore \sum_{i=1}^n m_i = M \text{ the total mass.}$$

Let $\vec{r}_1, \vec{r}_2, \dots, \vec{r}_n$ be their respective position vectors from a given origin O (Fig. 4.11).

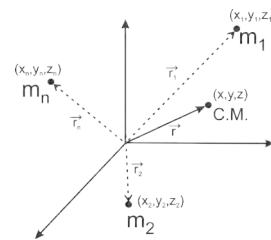


Fig. 4.11: Centre of mass for n particles.

Position vector $\vec{r}$ of their centre of mass from the same origin is then given by

$$\vec{r} = \frac{\sum_{i=1}^n m_i \vec{r}_i}{\sum_{i=1}^n m_i} = \frac{\sum_{i=1}^n m_i \vec{r}_i}{M}$$

If the origin itself is at the centre of mass,

$$\vec{r} = 0 \therefore \sum_{i=1}^n m_i \vec{r}_i = 0, \text{ then}$$

$\sum_{i=1}^n m_i \vec{r}_i$ gives the moment of masses

(similar to moment of force) about the centre of mass.

Thus, centre of mass is a point about which the summation of moments of masses in the system is zero.

If $x_1, x_2, \dots, x_n$ are the respective x -coordinates of $r_1, r_2, \dots, r_n$, the x -coordinate of the centre of mass is given by

$$x = \frac{\sum_{i=1}^n m_i x_i}{\sum_{i=1}^n m_i} = \frac{\sum_{i=1}^n m_i x_i}{M}$$

Similarly, y and z -coordinates of the centre of mass are respectively given by

$$y = \frac{\sum_{i=1}^n m_i y_i}{\sum_{i=1}^n m_i} = \frac{\sum_{i=1}^n m_i y_i}{M}$$

and

$$z = \frac{\sum_{i=1}^n m_i z_i}{\sum_{i=1}^n m_i} = \frac{\sum_{i=1}^n m_i z_i}{M}$$

(i) Continuous mass distribution: For a continuous mass distribution **with uniform density**, we need to use integration instead of summation. In this case, the position vector of the centre of mass is given by

$$\vec{r} = \frac{\int \vec{r} dm}{\int dm} = \frac{\int \vec{r} dm}{M},$$

where $\int dm = M$ is the total mass of the object. Then the Cartesian coordinates of c.m. are

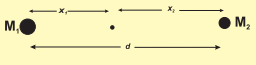
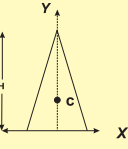
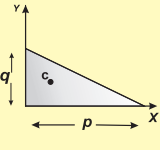
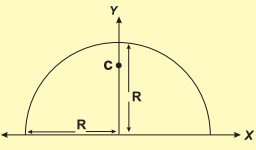
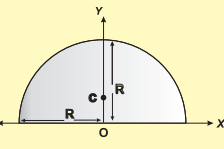
$$x = \frac{\int x dm}{\int dm} = \frac{\int x dm}{M}$$

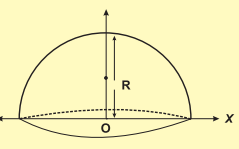
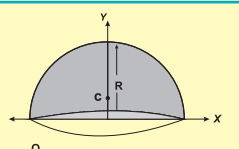
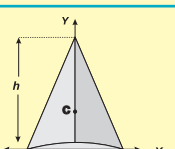
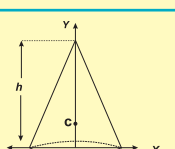
$$y = \frac{\int y dm}{\int dm} = \frac{\int y dm}{M}$$

$$z = \frac{\int z dm}{\int dm} = \frac{\int z dm}{M}$$

Using the expressions given above, the centres of mass of uniform symmetric objects can be obtained. Some of these are listed in the Table 4.1 given below:

Table 4.1: Coordinate of the centre of mass (c.m.) for some symmetrical objects

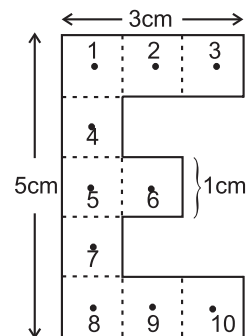
Coordinates of c.m.	Uniform Symmetric Objects
System of two point masses: c.m. divides the distance in inverse proportion of the masses	
Any geometrically symmetric object of uniform density.	Centre of mass at a geometrical centre of the object
Isosceles triangular plate $x_c = 0, y_c = \frac{H}{3}$	
Right angled triangular plate $x_c = \frac{p}{3}, y_c = \frac{q}{3}$	
Thin semicircular ring of radius R $x_c = 0, y_c = \frac{2R}{\pi}$	
Thin semicircular disc of radius R $x_c = 0, y_c = \frac{4R}{3\pi}$	

Hemispherical shell of radius R $x_c = 0, y_c = \frac{R}{2}$	
Solid hemisphere of radius R $x_c = 0, y_c = \frac{3R}{8}$	
Hollow right circular cone of height h $x_c = 0, y_c = \frac{h}{3}$	
Solid right circular cone of height h $x_c = 0, y_c = \frac{h}{4}$	

Example 4.12: A letter 'E' is prepared from a uniform cardboard with shape and dimensions as shown in the figure. Locate its centre of mass.

Solution: As the sheet is uniform, each square can be taken to be equivalent to mass m concentrated at its respective centre. These masses will then be at the points labelled with numbers 1 to 10, as shown in figure. Let us select the origin to be at the left central mass m_5 , as shown and all the co-ordinates to be in cm.

By symmetry, the centre of mass of m_1, m_2 and m_3 will be at m_2 (1, 2) having effective mass $3m$. Similarly, effective mass $3m$ due to m_8, m_9 and m_{10} will be at m_9 (1, -2). Again, by symmetry, the centre of mass of these two ($3m$ each) will have co-ordinates (1, 0). Mass m_6 is also having co-ordinates (1, 0). Thus, the effective mass at (1, 0) is $7m$.



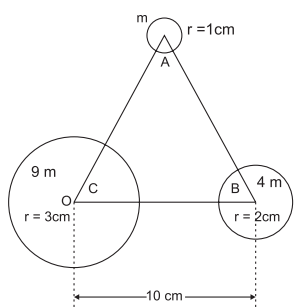
Using symmetry for m_4, m_5 and m_7 , there will be effective mass $3m$ at the origin (0, 0).

Thus, effectively, $3m$ and $7m$ are separated by 1 cm along x -direction. y -coordinate is not required.

$$x_c = \frac{m_1 x_1 + m_2 x_2}{m_1 + m_2} = \frac{3 \times 0 + 7 \times 1}{3 + 7} = 0.7 \text{ cm}$$

Alternately, for two point masses, the centre of mass divides the distance between them in the inverse ratio of their masses. Hence, 1 cm is divided in the ratio 7:3. $\therefore x_c = \frac{7}{7+3} \times 1 = 0.7 \text{ cm}$ from $3m$, i.e., from the origin at m_c

Example 4.13: Three thin walled uniform hollow spheres of radii 1 cm, 2 cm and 3 cm are so located that their centres are on the three vertices of an equilateral triangle ABC having each side 10 cm. Determine centre of mass of the system.



Solution: Mass of a thin walled uniform hollow sphere is proportional to its surface area, (as density is constant) hence proportional to r^2 . Thus, if mass of the sphere at A is $m_A = m$, then $m_B = 4m$ and $m_C = 9m$. By symmetry of the spherical surface, their centres of mass are at their respective centres, i.e., at A, B and C.

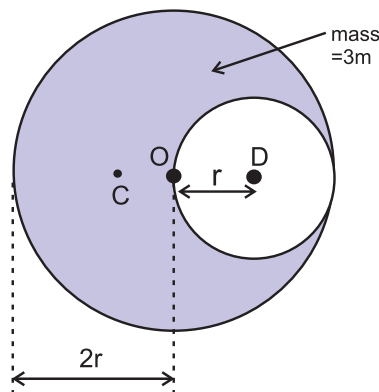
Let us choose the origin to be at C, where the largest mass $9m$ is located and the point B with mass $4m$ on the positive x -axis. With this, the co-ordinates of C are (0, 0) and that of B are (10, 0). (Locating the origin at the larger mass here save our efforts of calculations like multiplications with larger numbers). If A of mass m is taken in the first quadrant, its co-ordinates will be $\left[5, \frac{10\sqrt{3}}{2}\right]$

$$\begin{aligned} x_c &= \frac{m_A x_A + m_B x_B + m_C x_C}{m_A + m_B + m_C} \\ &= \frac{m \times 5 + 4m \times 10 + 9m \times 0}{m + 4m + 9m} = \frac{45}{14} \text{ cm} \end{aligned}$$

$$\begin{aligned} y_c &= \frac{m_A y_A + m_B y_B + m_C y_C}{m_A + m_B + m_C} \\ &= \frac{m \times \frac{10\sqrt{3}}{2} + 4m \times 0 + 9m \times 0}{m + 4m + 9m} = \frac{10\sqrt{3}}{28} \text{ cm} \end{aligned}$$

Example 4.14: A hole of radius r is cut from a uniform disc of radius $2r$. Centre of the hole is at a distance r from centre of the disc. Locate centre of mass of the remaining part of the disc.

Solution: Method I: (Using entire disc): Before cutting the hole, c.m. of the full disc was at its centre. Let this be our origin O. Centre of mass of the cut portion is at its centre D. Thus, it is at a distance $x_1 = r$ from the origin. Let C be the centre of mass of the remaining disc. Obviously, it should be on the extension of the line DO. Let it be at a distance $x_2 = x$ from the origin. As the disc is uniform, mass of any of its part is proportional to the area of that part.



Thus, if m is the mass of the cut disc, mass of the entire disc must be $4m$ and mass of the remaining disc will be $3m$.

$$x_c = \frac{m_1 x_1 + m_2 x_2}{m_1 + m_2} \quad \text{As centre of mass of the full disc is at the origin, we can write,}$$

$$0 = \frac{m \times r + 3m \times (x)}{m + 3m} \quad \therefore x = \frac{-r}{3}$$

Method II: (Using negative mass): Let $\vec{R}$ be the position vector of the centre of mass of the uniform disc of mass M . Mass m is with centre of mass at position vector $\vec{r}$ from the centre of the disc. Position vector of the centre of mass of the remaining disc is then given by

$$\vec{r}_c = \frac{M\vec{R} - m\vec{r}}{M - m} \dots \text{(as if there is a negative mass, i.e., } m_2 = -m)$$

With our description, $M = 4m$, $m = m$,
 $R = 0$ and $r = r \therefore r_c = \frac{-mr}{3m} = \frac{-r}{3}$... Same as method I.

4.13.2. Velocity of centre of mass:

Let $v_1, v_2, \dots v_n$ be the velocities of a system of point masses $m_1, m_2, \dots m_n$. Velocity of the centre of mass of the system is given by

$$\begin{aligned} \vec{v}_{cm} &= \frac{\sum_1^n m_i \vec{v}_i}{\sum_1^n m_i} = \frac{\sum_1^n m_i \vec{v}_i}{M} \\ &= \frac{\text{resultant linear momentum}}{\text{total mass}} \\ &= \text{weighted average of momenta} \end{aligned}$$

x, y and z components of $\vec{v}$ can be obtained similarly.

For continuous distribution, $\vec{v}_{cm} = \frac{\int \vec{v} dm}{M}$

4.13.3. Acceleration of the centre of mass:

Let $a_1, a_2, \dots a_n$ be the accelerations of a system of point masses $m_1, m_2, \dots m_n$. Acceleration of the centre of mass of the system is given by

$$\begin{aligned} \vec{a}_{cm} &= \frac{\sum_1^n m_i \vec{a}_i}{\sum_1^n m_i} = \frac{\sum_1^n m_i \vec{a}_i}{M} \\ &= \frac{\text{resultant force}}{\text{total mass}} \\ &= \text{weighted average of forces} \end{aligned}$$

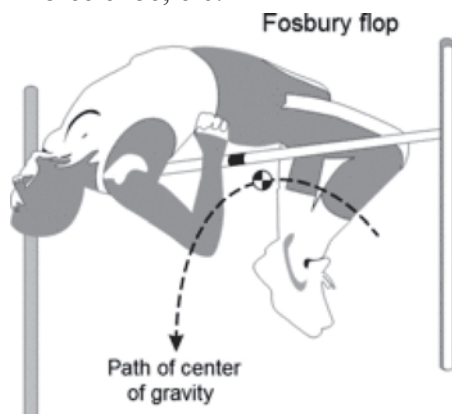
x, y and z components of $\vec{a}$ can be obtained similarly.

For continuous distribution, $\vec{a}_{cm} = \frac{\int \vec{a} dm}{M}$

4.13.4. Characteristics of centre of mass:

- Centre of mass is a hypothetical point at which entire mass of the body can be assumed to be concentrated.
- Centre of mass is a location, and not a physical quantity.
- Centre of mass is particle equivalent of a given object for applying laws of motion.
- Centre of mass is the point at which, if a force is applied, it causes only linear acceleration and not angular acceleration.
- Centre of mass is located at the centroid, for a rigid body **of uniform density**.
- Centre of mass is located at the geometrical centre, for a symmetric rigid body **of uniform density**.
- Location of centre of mass can be changed **only** by an external unbalanced force.
- Internal forces (like during collision or explosion) **never** change the location of centre of mass.
- Position of the centre of mass depends only upon the distribution of mass, however, to describe its location we may use a coordinate system with a *suitable* origin. In statistical terms the centre of mass is decided by the *weighted* average of individual masses. This is obtained by giving proper mass *weightage* to the distance. This should be clear from the mathematical expression for the location of the centre of mass.
- For a system of particles, the centre of mass *need not* coincide with any of the particles.
- While balancing an object on a pivot, the line of action of *weight* must pass through the centre of mass and the pivot. Quite often, this is an unstable equilibrium.
- Centre of mass of a system of only two particles divides the distance between the particles in an inverse ratio of their masses, i.e., it is closer to the heavier mass.
- Centre of mass is a point about which the summation of moments of masses in the system is zero.
- If there is an axial symmetry for a given object, the centre of mass lies on the axis of symmetry.
- If there are multiple axes of symmetry for a given object, the centre of mass is at their point of intersection.
- Centre of mass need not be within the

body (See the photograph given below: **Picture 4.1**). Other examples are a ring, a horse shoe, etc.



Picture 4.1: Courtesy Wikipedia: Estimated center of mass/gravity of a high jumper doing a Fosbury Flop. Note that it is below the bar in this position. This is possible because our head and legs are much heavier than the fleshy part. Increase in the gravitational potential energy of the high jumper depends upon this point.

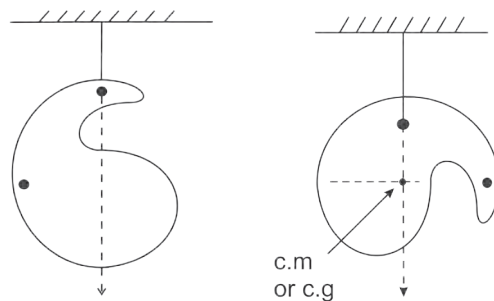
4.14. Centre of gravity

Centre of gravity (c.g.) of a body is the point around which the resultant torque due to force of gravity on the body is zero. Analogous to centre of mass, it is the weighted average of the gravitational forces (weights) on individual particles.

For uniform gravitational field (in simple words, if g is constant), c.g. always coincides with the c.m. Obviously it is true for all the objects *on the Earth* in our daily life. Thus, in common usage, the terms c.g. and c.m. are

used for same purpose. This property can be used to determine the c.g. (or c.m.) of a laminar (laminar means like a leaf – two dimensional) object.

In Fig. 4.12, a laminar object is suspended from a rigid support at two orientations. Lines are to be drawn on the object parallel to the plumb line shown. Plumb line is always vertical, i.e., parallel to the line of action of gravitational force. Intersection of the lines drawn is then the point through which line of action of the gravitational force passes for any orientation. Thus, it gives the location of the c.g. or c.m.



Centre of mass is a fixed property for a given rigid body in spite of any orientation. The centre of gravity may depend upon non-uniformity of the gravitational field, in turn, will depend upon the orientation. For objects on the Earth, this will be possible only if the size of an object is comparable to that of the Earth (size at least few thousand km). In such cases, the c.g. will be slightly lower than the c.m. as on the lower side of an object the gravitational field is stronger. Of course, we shall not come across such an object.



Exercises

1. Choose the correct answer.

- Consider following pair of forces of equal magnitude and opposite directions:
 (P) Gravitational forces exerted on each other by two point masses separated by a distance.
 (Q) Couple of forces used to rotate a water tap.
 (R) Gravitational force and normal force experienced by an object kept on a table.

For which of these pair/pairs the two forces

do NOT cancel each other's translational effect?

- (A) Only P (B) Only P and Q
 (C) Only R (D) Only Q and R

- Consider following forces: (w) Force due to tension along a string, (x) Normal force given by a surface, (y) Force due to air resistance and (z) Buoyant force or upthrust given by a fluid.

Which of these are electromagnetic forces?
 (A) Only w, y and z

- (B) Only w , x and y
 (C) Only x , y and z
 (D) All four.
- iii) At a given instant three point masses m , $2m$ and $3m$ are equidistant from each other. Consider only the gravitational forces between them. Select correct statement/s for this instance only:
 (A) Mass m experiences maximum force.
 (B) Mass $2m$ experiences maximum force.
 (C) Mass $3m$ experiences maximum force.
 (D) All masses experience force of same magnitude.
- iv) The rough surface of a horizontal table offers a definite *maximum* opposing force to initiate the motion of a block along the table, which is proportional to the resultant normal force given by the table. Forces F_1 and F_2 act at the same angle θ with the horizontal and both are just initiating the sliding motion of the block along the table. Force F_1 is a pulling force while the force F_2 is a pushing force. $F_2 > F_1$, because
 (A) Component of F_2 adds up to weight to increase the normal reaction.
 (B) Component of F_1 adds up to weight to increase the normal reaction.
 (C) Component of F_2 adds up to the opposing force.
 (D) Component of F_1 adds up to the opposing force.
- v. A mass $2m$ moving with some speed is *directly* approaching another mass m moving with double speed. After some time, they collide with coefficient of restitution 0.5. Ratio of their respective speeds after collision is
 (A) $2/3$ (B) $3/2$
 (C) 2 (D) $1/2$
- vi. A uniform rod of mass $2m$ is held horizontal by two sturdy, practically inextensible vertical strings tied at its ends. A boy of mass $3m$ hangs himself at one third length of the rod. Ratio of the tension in the string close to the boy to that in the other string is
 (A) 2 (B) 1.5
 (C) $4/3$ (D) $5/3$
- vii. Select WRONG statement about centre of mass:
- (A) Centre of mass of a 'C' shaped uniform rod can never be a point on that rod.
 (B) If the line of action of a force passes through the centre of mass, the moment of that force is zero.
 (C) Centre of mass of our Earth is not at its geometrical centre.
 (D) While balancing an object on a pivot, the line of action of the gravitational force of the earth passes through the centre of mass of the object.
- viii. For which of the following objects will the centre of mass NOT be at their geometrical centre?
 (I) An egg
 (II) a cylindrical box full of rice
 (III) a cubical box containing assorted sweets
 (A) Only (I)
 (B) Only (I) and (II)
 (C) Only (III)
 (D) All, (I), (II) and (III).

2. Answer the following questions.

- i) In the following table, every entry on the left column can match with any number of entries on the right side. Pick up **all** those and write respectively against A, B, C and D.

Name of the force		Type of the force	
A	Force due to tension in a string	P	EM force
B	Normal force	Q	Reaction force
C	Frictional force	R	Conservative force
D	Resistive force offered by air or water for objects moving through it.	S	Non-conservative force

- ii) In real life objects, never travel with uniform velocity, even on a horizontal surface, unless *something* is done? Why is it so? What is to be done?
- iii) For the study of any kind of motion, we never use Newton's first law of motion directly. Why should it be studied?
- iv) Are there any situations in which we cannot apply Newton's laws of motion? Is there any alternative for it?

- v) You are inside a closed capsule from where you are not able to see anything about the outside world. Suddenly you feel that you are pushed towards your right. Can you explain the possible cause (s)? Is it a feeling or a reality? Give at least one more situation like this.
- vi) Among the four fundamental forces, only one force governs your daily life almost entirely. Justify the statement by stating that force.
- vii) Find the odd man out: (i) Force responsible for a string to become taut on stretching (ii) Weight of an object (iii) The force due to which we can hold an object in hand.
- viii) You are sitting next to your friend on ground. Is there any gravitational force of attraction between you two? If so, why are you not coming together naturally? Is any force other than the gravitational force of the earth coming in picture?
- ix) Distinguish between: (A) Real and pseudo forces, (B) Conservative and non-conservative forces, (C) Contact and non-contact forces, (C) Inertial and non-inertial frames of reference.
- x) State the formula for calculating work done by a force. Are there any conditions or limitations in using it directly? If so, state those clearly. Is there any mathematical way out for it? Explain.
- xi) Justify the statement, "Work and energy are the two sides of a coin".
- xii) From the terrace of a building of height H , you dropped a ball of mass m . It reached the ground with speed v . Is the relation $mgH = \frac{1}{2}mv^2$ applicable exactly? If not, how can you account for the difference? Will the ball bounce to the same height from where it was dropped?
- xiii) State the law of conservation of linear momentum. It is a consequence of which law? Given an example from our daily life for conservation of momentum. Does it hold good during burst of a cracker?
- xiv) Define coefficient of restitution and obtain its value for an elastic collision and a perfectly inelastic collision.
- xv) Discuss the following as special cases of elastic collisions and obtain their exact or approximate final velocities in terms of their initial velocities.
- Colliding bodies are identical.
 - A very heavy object collides on a lighter object, initially at rest.
 - A very light object collides on a comparatively much massive object, initially at rest.
- xvi) A bullet of mass m_1 travelling with a velocity u strikes a stationary wooden block of mass m_2 and gets embedded into it. Determine the expression for loss in the kinetic energy of the system. Is this violating the principle of conservation of energy? If not, how can you account for this loss?
- xvii) One of the effects of a force is to change the momentum. Define the quantity related to this and explain it for a variable force. Usually when do we define it instead of using the force?
- xviii) While rotating an object or while opening a door or a water tap we apply a force or forces. Under which conditions is this process easy for us? Why? Define the vector quantity concerned. How does it differ for a single force and for two opposite forces with different lines of action?
- xix) Why is the moment of a couple independent of the axis of rotation even if the axis is fixed?
- xx) Explain balancing or mechanical equilibrium. Linear velocity of a rotating fan as a whole is generally zero. Is it in mechanical equilibrium? Justify your answer.
- xxi) Why do we need to know the centre of mass of an object? For which objects, its position may differ from that of the centre of gravity?

Use $g = 10 \text{ m s}^{-2}$, unless, otherwise stated.

3. Solve the following problems.

- i) A truck of mass 5 ton is travelling on a

horizontal road with 36 km hr^{-1} stops on traveling 1 km after its engine fails suddenly. What fraction of its weight is the frictional force exerted by the road?

If we assume that the story repeats for a car of mass 1 ton i.e., can moving with same speed stops in similar distance same how much will the fraction be?

[Ans: $\frac{1}{200}$ in the both]

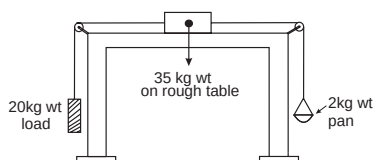
- ii) A lighter object A and a heavier object B are initially at rest. Both are imparted the same linear momentum. Which will start with greater kinetic energy: A or B or both will start with the same energy?

[Ans: A]

- iii) As I was standing on a weighing machine inside a lift it recorded 50 kg wt. Suddenly for few seconds it recorded 45 kg wt. What must have happened during that time? Explain with complete numerical analysis.
[Ans: Lift must be coming down with acceleration $\frac{g}{10} = 1 \text{ m s}^{-2}$]

- iv) Figure below shows a block of mass 35 kg resting on a table. The table is so rough that it offers a self adjusting resistive force 10% of the weight of the block for its sliding motion along the table. A 20 kg wt load is attached to the block and is passed over a pulley to hang freely on the left side. On the right side there is a 2 kg wt pan attached to the block and hung freely. Weights of 1 kg wt each, can be added to the pan. Minimum how many and maximum how many such weights can be added into the pan so that the block does not slide along the table?

[Ans: Min 15, maximum 21].



- v) Power is rate of doing work or the rate at which energy is supplied to the system. A constant force F is applied to a body of mass m . Power delivered by the force at time t from the start is proportional to
(a) t (b) t^2

(c) $\sqrt{t}$

(d) t^0

Derive the expression for power in terms of F , m and t .

[Ans: $p = \frac{F^2 t}{m}, \therefore p \propto t$]

- vi) 40000 litre of oil of density 0.9 g cc is pumped from an oil tanker ship into a storage tank at 10 m higher level than the ship in half an hour. What should be the power of the pump?

[Ans: 2 kW]

- vii) Ten identical masses (m each) are connected one below the other with 10 strings. Holding the topmost string, the system is accelerated upwards with acceleration $g/2$. What is the tension in the 6th string from the top (Topmost string being the first string)?

[Ans: 6 mg]

- viii) Two galaxies of masses 9 billion solar mass and 4 billion solar mass are 5 million light years apart. If, the Sun has to cross the line joining them, without being attracted by either of them, through what point it should pass?

[Ans: 3 million light years from the 9 billion solar mass]

- ix) While decreasing linearly from 5 N to 3 N, a force displaces an object from 3 m to 5 m. Calculate the work done by this force during this displacement.

[Ans: 8 N]

- x) Variation of a force in a certain region is given by $F = 6x^2 - 4x - 8$. It displaces an object from $x = 1 \text{ m}$ to $x = 2 \text{ m}$ in this region. Calculate the amount of work done. [Ans: Zero]

- xi) A ball of mass 100 g dropped on the ground from 5 m bounces repeatedly. During every bounce 64% of the potential energy is converted into kinetic energy. Calculate the following:

- (a) Coefficient of restitution.
(b) Speed with which the ball comes up from the ground after third bounce.
(c) Impulse given by the ball to the ground during this bounce.
(d) Average force exerted by the ground

if this impact lasts for 250 ms.

(e) Average pressure exerted by the ball on the ground during this impact if contact area of the ball is 0.5 cm^2 .

[Ans: 0.8, 5.12 m/s, 1.152 N s, 4.608 N, $9.216 \times 10^4 \text{ N/m}^2$]

- xii) A spring ball of mass 0.5 kg is dropped from some height. On falling freely for 10 s, it explodes into two fragments of mass ratio 1:2. The lighter fragment continues to travel downwards with speed of 60 m/s. Calculate the kinetic energy supplied during explosion.

[Ans: 200 J]

- xiii) A marble of mass $2m$ travelling at 6 cm/s is directly followed by another marble of mass m with double speed. After collision, the heavier one travels with the average initial speed of the two. Calculate the coefficient of restitution.

[Ans: 0.5]

- xiv) A, 2 m long wooden plank of mass 20 kg is pivoted (supported from below) at 0.5 m from either end. A person of mass 40 kg starts walking from one of these pivots to the farther end. How far can the person walk before the plank topples?

[Ans: 1.25 m]

- xv) A 2 m long ladder of mass 10 kg is kept against a wall such that its base is 1.2 m

away from the wall. The wall is smooth but the ground is rough. Roughness of the ground is such that it offers a maximum horizontal resistive force (for sliding motion) half that of normal reaction at the point of contact. A monkey of mass 20 kg starts climbing the ladder. How far can it climb **along** the ladder? How much is the horizontal reaction at the wall?

[Ans: 1.5 m, 15 N]

- xvi) Four uniform solid cubes of edges 10 cm, 20 cm, 30 cm and 40 cm are kept on the ground, touching each other in order. Locate centre of mass of their system.

[Ans: 65 cm,

17.7 cm]

- xvii) A uniform solid sphere of radius R has a hole of radius $R/2$ drilled inside it. One end of the hole is at the centre of the sphere while the other is at the boundary. Locate centre of mass of the remaining sphere.

[Ans: $-R/14$]

- xviii) In the following table, every item on the left side can match with any number of items on the right hand side. Select all those.

Types of collision	Illustrations
(a) Elastic collision	(i) A ball hit by a bat.
(b) Inelastic collision	(ii) Molecular collisions responsible for pressure exerted by a gas.
(c) Perfectly inelastic collision	(iii) A stationary marble A is hit by marble B and the marble B comes to rest.
(d) Head on collision	(iv) A blob of clay dropped on the ground sticks to the ground.
	(v) Out of anger, giving a kick to a wall.
	(vi) A striker hits the boundary of a carrom board in a direction perpendicular to the boundary and rebounds.



Can you recall?

1. When released from certain height why do objects tend to fall vertically downwards?
2. What is the shape of the orbits of planets?
3. What are Kepler's laws?

5.1 Introduction:

All material objects have a natural tendency to get attracted towards the Earth. In many natural phenomena like coconut falling from trees, raindrops falling from the clouds, etc., the same tendency is observed. All bodies are attracted towards the Earth with constant acceleration. This fact was recognized by Italian physicist Galileo. He is said to have demonstrated it by releasing two balls of different masses from top of the leaning tower of Pisa which reached the ground at the same time.

Indian astronomer and mathematician Aryabhatta (476-550 A.D.) studied the motion of the moon, Earth and other planets in the 5th century A.D. In his book 'Aryabhatiya', he concluded that the Earth revolves about its own axis and it moves in a circular orbit around the Sun. Also the moon revolves in a circular orbit around the Earth. Almost a thousand years after Aryabhatta, Tycho Brahe (1546-1601) and Johannes Kepler (1571-1630) studied planetary motion through careful observations. Kepler analysed the huge data meticulously recorded by Tycho Brahe and established three laws of planetary motion. He showed that the motion of planets follow these laws. The reason why planets obey these laws was provided by Newton. He explained that gravitation is the phenomenon responsible for keeping planets in their orbits around the Sun. The moon also revolves around the Earth due to gravitation. Gravitation compels dispersed matter to coalesce, hence the existence of the Earth, the Sun and all material macroscopic objects in the universe.

Every massive object in the universe experiences gravitational force. It is the force of mutual attraction between any two objects by

virtue of their masses. It is always an attractive force with infinite range. It does not depend upon intervening medium. It is much weaker than other fundamental forces. Gravitational force is 10^{-39} times weaker than strong nuclear force.

5.2 Kepler's Laws:

Kepler's laws of planetary motion describe the orbits of the planets around the Sun. He published first two laws in 1609 and the third law in 1619. These laws are the result of the analysis of the data collected by Tycho Brahe through years of observations of the planetary motion.



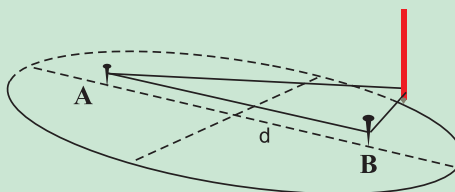
Do you know ?

Drawing an ellipse

An ellipse is the locus of the points in a plane such that the sum of their distances from two fixed points, called the foci, is constant. You can draw an ellipse by the following procedure.

- 1) Insert two tacks or drawing pins, A and B, as shown in the figure into a sheet of drawing paper at a distance 'd' apart.
- 2) Tie the two ends of a piece of thread whose length is greater than '2d' and place the loop around AB as shown in the figure.
- 3) Place a pencil inside the loop of thread, pull the thread taut and move the pencil sidewise, keeping the thread taut.

The pencil will trace an ellipse.



1 Law of orbit

All planets move in elliptical orbits around the Sun with the Sun at one of the foci of the ellipse.

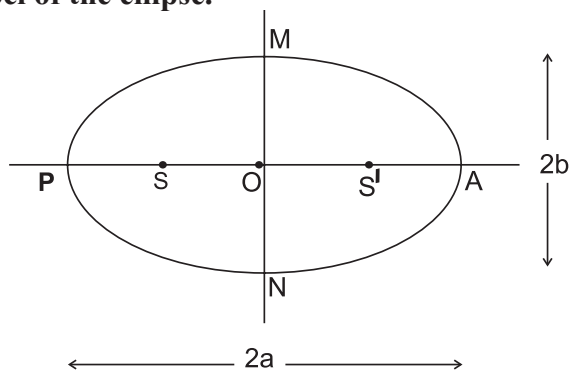


Fig. 5.1: An ellipse traced by a planet with the Sun at the focus.

The orbit of a planet around the Sun is shown in Fig. 5.1.

Here, S and S' are the foci of the ellipse the Sun being at S.

P is the closest point along the orbit from S and is, called 'Perihelion'.

A is the farthest point from S and is, called 'Aphelion'.

PA is the major axis = $2a$.

PO and AO are the semimajor axes = a .

MN is the minor axis = $2b$.

MO and ON are the semiminor axes = b

2. Law of areas

The line that joins a planet and the Sun sweeps equal areas in equal intervals of time.

Kepler observed that planets do not move around the Sun with uniform speed. They move faster when they are nearer to the Sun while they move slower when they are farther from the Sun. This is explained by this law.

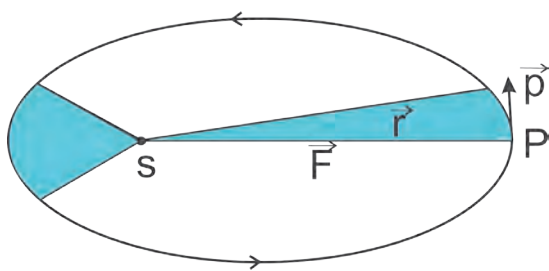


Fig. 5.2: The orbit of a planet P moving around the Sun.

Fig 5.2 shows the orbit of a planet. The shaded areas are the areas swept by SP, the line joining the planet and the Sun, in fixed intervals of time. These are equal according to the second law.

The law of areas can be understood as an outcome of conservation of angular momentum. It is valid for any central force. A central force on an object is a force which is always directed along the line joining the position of object and a fixed point usually taken to be the origin of the coordinate system. The force of gravity due to the Sun on a planet is always along the line joining the Sun and the planet (Fig. 5.2). It is thus a central force. Suppose the Sun is at the origin. The position of planet is denoted by $\vec{r}$ and the perpendicular component of its momentum is denoted by $\vec{p}$ (component $\perp \vec{r}$). The area swept by the planet of mass m in given interval Δt is $\overline{\Delta A}$ which is given by

$$\overline{\Delta A} = \frac{1}{2}(\vec{r} \times \vec{v} \Delta t) \quad \text{--- (5.1)}$$

As for small Δt , $\vec{v}$ is perpendicular to $\vec{r}$ and this is the area of the triangle.

$$\therefore \frac{\overline{\Delta A}}{\Delta t} = \frac{1}{2}(\vec{r} \times \vec{v}) \quad \text{--- (5.2)}$$

Linear momentum ($\vec{p}$) is the product of mass and velocity.

$$\vec{p} = m\vec{v} \quad \text{--- (5.3)}$$

$\therefore$ putting $\vec{v} = \vec{p}/m$ in the above equation, we get

$$\frac{\overline{\Delta A}}{\Delta t} = \frac{1}{2} \left(\vec{r} \times \frac{\vec{p}}{m} \right) \quad \text{--- (5.4)}$$

Angular momentum $\vec{L}$ is the rotational equivalent of linear momentum and is defined as

$$\therefore \vec{L} = \vec{r} \times \vec{p} \quad \text{--- (5.5)}$$

For central force the angular momentum is conserved.

$$\therefore \frac{\overline{\Delta A}}{\Delta t} = \frac{\vec{L}}{2m} = \text{constant} \quad \text{--- (5.6)}$$

$\therefore$ This proves the law of areas. This is a consequence of the gravitational force being a central force.

3. Law of periods

The square of the time period of revolution of a planet around the Sun is proportional to the cube of the semimajor axis of the ellipse traced by the planet.

If r is length of semimajor axis then, this law states that

$$T^2 \propto r^3$$

or $\frac{T^2}{r^3} = \text{constant} \quad \text{--- (5.7)}$

Kepler's laws were based on regular observations of the motion of planets. Kepler did not know why the planets obey these laws, i.e. he had not derived these laws.

Table 5.1 gives data from measurements of planetary motions which confirm Kepler's law of periods.

Table 5.1: Kepler's third law

Planet	Semi-major axis in units of 10^{10} m	Period in years	T^2/r^3 in units of $10^{-34} \text{ y}^2 \text{ m}^{-3}$
Mercury	5.79	0.24	2.95
Venus	10.8	0.615	3.00
Earth	15.0	1	2.96
Mars	22.8	1.88	2.98
Jupiter	77.8	11.9	3.01
Saturn	143	29.5	2.98
Uranus	287	84.	2.98
Neptune	450	165	2.99
Pluto	590	248	2.99

Example 5.1: What would be the average duration of year if the distance between the Sun and the Earth becomes

- (A) thrice the present distance.
- (B) twice the present distance.

Solution:

(A) Let r_1 = Present distance between the Earth and Sun

$$T = 365 \text{ days.}$$

$$\text{If } r_2 = 3r_1, T_2 = ?$$

According to Kepler's law of period

$$T_1^2 \propto r_1^3 \text{ and } T_2^2 \propto r_2^3$$

$$\begin{aligned} \therefore \frac{T_2^2}{T_1^2} &= \frac{r_2^3}{r_1^3} \\ \therefore \frac{T_2}{T_1} &= \left(\frac{r_2}{r_1} \right)^{3/2} \\ &= \left(\frac{3r_1}{r_1} \right)^{3/2} \\ \therefore \frac{T_2}{T_1} &= \sqrt{27} \end{aligned}$$

$$\begin{aligned} T_2 &= T_1 \times \sqrt{27} \\ &= 365 \times \sqrt{27} \\ &= 1897 \text{ days} \end{aligned}$$

(B) If $r_2 = 2r_1$, $T_2 = ?$

$$\begin{aligned} \frac{T_2^2}{T_1^2} &= \frac{r_2^3}{r_1^3} \\ \frac{T_2^2}{T_1^2} &= \left(\frac{2r_1}{r_1} \right)^3 \\ \therefore \frac{T_2}{T_1} &= \sqrt{8} \end{aligned}$$

$$\begin{aligned} T_2 &= T_1 \sqrt{8} \\ &= 365 \sqrt{8} \\ &= 1032 \text{ days.} \end{aligned}$$

5.3 Universal Law of Gravitation:

When objects are released near the surface of the Earth, they always fall down to the ground, i.e., the Earth attracts objects towards itself. Galileo (1564-1642) pointed out that heavy and light objects, when released from the same height, fall towards the Earth at the same speed, i.e., they have the same acceleration. Newton went beyond (the Earth and objects falling on it) and proposed that the force of attraction between masses is universal. Newton stated the universal law of gravitation which led to an explanation of terrestrial gravitation. It also explains Kepler's laws and provides the reason behind the observed motion of planets around the Sun.

In 1665, Newton studied the motion of moon around the Earth. It was known that the moon completes one revolution about the Earth in 27.3 days. The distance from the Earth to

the moon is 3.85×10^5 km. The motion of the moon is in almost a circular orbit around the Earth with constant angular speed ω . As it is a circular motion, the moon must be constantly acted upon by a force directed towards the Earth which is at the centre of the circle. This force is the centripetal force, and is given by

$$F = mr\omega^2 \quad \text{--- (5.8)}$$

where m is the mass of the moon and r is the distance between the centres of the moon and the Earth.

Also we have $F = ma$ from Newton's laws of motion.

$$\therefore ma = mr\omega^2$$

$$\therefore a = r\omega^2 \quad \text{--- (5.9)}$$

As angular velocity in terms of time period is given as

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

we get

$$a = r \left(\frac{2\pi}{T} \right)^2 \quad \text{--- (5.10)}$$

Substituting values of r and T , we get

$$a = \frac{3.85 \times 10^5 \times 10^3 4\pi^2 m}{(27.3 \times 24 \times 60 \times 60)^2 s^2}$$

$$\therefore a = 0.0027 \text{ m/s}^2$$

This is the acceleration of the moon which is towards the centre of the Earth, i.e., centre of orbit in which the moon revolves. What could



Do you know ?

The value of acceleration due to gravity can be assumed to be constant when we are dealing with objects close to the surface of the Earth. This is because the difference in their distances from the centre of the Earth is negligible.

be the force which produces this acceleration?

Newton assumed that the laws of nature are the same for Earthly objects and for celestial bodies. As this acceleration is much smaller than the acceleration felt by bodies near the surface of the Earth (while falling on Earth), he concluded that the acceleration felt

by an object due to the gravitational force of the Earth must be decreasing with distance of the body from the Earth. (Remember that the value of acceleration due to Earth's gravity at the surface is 9.8 m/s^2)

We have,

$$\frac{a_{\text{object}}}{a_{\text{moon}}} = \frac{9.8 \text{ m/s}^2}{0.0027 \text{ m/s}^2} \approx 3600$$

Also,

$$\frac{\text{distance of moon from the Earth's centre}}{\text{distance of object from the Earth's centre}}$$

$$= \frac{3.85 \times 10^5 \text{ km}}{6378 \text{ km}} \approx 60$$

Thus from the above two equations we get

$$\therefore \frac{a_{\text{object}}}{a_{\text{moon}}} = \left[\frac{\text{distance of moon}}{\text{distance of object}} \right]^2 \quad \text{--- (5.11)}$$

Newton therefore concluded that the acceleration of an object towards the Earth is inversely proportional to the square of distance of object from the centre of the Earth.

$$\therefore a \propto \frac{1}{r^2}$$

$$\text{As, } F = ma$$

Therefore, the force exerted by the Earth on an object of mass m at a distance r from it is

$$F \propto \frac{m}{r^2}$$

Similarly an object also exerts a force on the Earth which is

$$F_E \propto \frac{M}{r^2}$$

where M is the mass of the Earth.

According to Newton's third law of motion, the force on a body due to the Earth has to be equal to the force on the Earth due to the object. Hence the force F is also proportional to the mass of the Earth. Hence Newton concluded that the gravitational force between the Earth and an object of mass m is

$$F \propto \frac{Mm}{r^2}$$

He then generalized it to gravitational force between any two objects and stated his

Universal law of gravitation as follows.

Every particle of matter attracts every other particle of matter with a force which is directly proportional to the product of their masses and inversely proportional to the square of the distance between them.

This law is applicable to all material objects in the universe. Hence it is known as the universal law of gravitation.

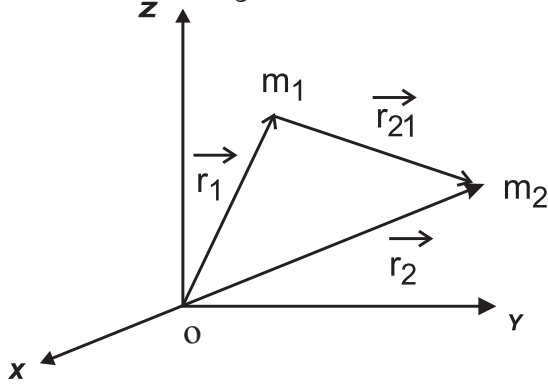


Fig. 5.3: Gravitational force between masses m_1 and m_2 .

If two bodies of masses m_1 and m_2 are separated by a distance r , then the gravitational force of attraction between them can be written as

$$F \propto \frac{m_1 m_2}{r^2}$$

$$\text{or, } F = G \frac{m_1 m_2}{r^2} \quad \text{--- (5.12)}$$

where G is a constant known as the universal gravitational constant. Its value in SI units is given by

$$G = 6.67 \times 10^{-11} \text{ N m}^2/\text{kg}^2$$

and its dimensions are

$$[G] = [\text{L}^3 \text{M}^{-1} \text{T}^2].$$

The gravitational force is an attractive force and it acts along the line joining the two bodies. The forces exerted by two bodies on each other have same magnitude but have opposite directions, they form an action-reaction pair.

Example 5.2: The gravitational force between two bodies is 1 N. If distance between them is doubled, what will be the gravitational force between them?

Solution: Let the masses of the two bodies be m_1 and m_2 and the distance between them be r .

$$\text{The force between them, } F_1 = \frac{Gm_1 m_2}{r^2}$$

When the distance between them is doubled the force becomes, $F_2 = \frac{Gm_1 m_2}{(2r)^2} = \frac{Gm_1 m_2}{4r^2}$

$$\therefore \frac{F_1}{F_2} = \frac{Gm_1 m_2}{r^2} \times \frac{4r^2}{Gm_1 m_2}$$

$$\therefore \frac{F_1}{F_2} = 4$$

$$\therefore F_2 = \frac{1}{4} \text{ N} \quad (\because F_1 = 1 \text{ N})$$

$$F_2 = 0.25 \text{ N}$$

$\therefore$ Force become one fourth (0.25 N)

Figure 5.3 shows two point masses m_1 and m_2 with position vectors $\vec{r}_1$ and $\vec{r}_2$ respectively from origin O . The position vector of m_2 with respect to m_1 is then given by $\vec{r}_{21} = \vec{r}_2 - \vec{r}_1$. Similarly $\vec{r}_{12} = \vec{r}_1 - \vec{r}_2 = -\vec{r}_{21}$ and if $r = |\vec{r}_{12}| = |\vec{r}_{21}|$, the formula for force on m_2 due to m_1 can be expressed in vector form as,

$$\vec{F}_{21} = G \frac{m_1 m_2}{r^2} (-\hat{r}_{21}) \quad \text{--- (5.13)}$$

where $\hat{r}_{21}$ is the unit vector from m_1 to m_2 . The force $\vec{F}_{21}$ is directed from m_2 to m_1 . Similarly, force experienced by m_1 due to m_2 is $\vec{F}_{12}$

$$\vec{F}_{12} = G \frac{m_1 m_2}{r^2} (-\hat{r}_{12}) \quad \text{--- (5.14)}$$

$$\therefore \vec{F}_{12} = -\vec{F}_{21} \quad \text{--- (5.15)}$$

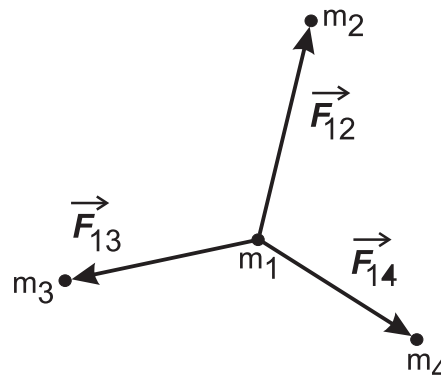


Fig. 5.4: gravitational force due to a collection of masses.

This law refers to two point masses. For a collection of point masses, the force on any one of them is the vector sum of the gravitational forces exerted by all the other point masses. As shown in Fig. 5.4, the resultant force on point mass m_1 is the vector sum of forces $\vec{F}_{12}$, $\vec{F}_{13}$ and $\vec{F}_{14}$ due to point masses m_2 , m_3 and m_4 respectively. Masses m_2 , m_3 and m_4 are also attracted towards mass m_1 and there is also mutual attraction between masses m_2 , m_3 and m_4 but these forces are not shown in the figure.

For n particles, force on i^{th} mass $\vec{F}_i = \sum_{\substack{j=1 \\ j \neq i}}^n \vec{F}_{ij}$

where $\vec{F}_{ij}$ is the force on i^{th} particle due to j^{th} particle.

The gravitational force between an extended object like the Earth and a point mass A can be obtained by obtaining the vector sum of forces on the point mass A due to each of the point mass which make up the extended object. We can consider the following two special cases, for which we can get a simple result. We will state the result here and show how it can be understood qualitatively.

(1) The gravitational force of attraction due to a hollow, thin spherical shell of uniform density, on a point mass situated inside it is zero.

This can be qualitatively understood as follows. First let us consider the case when the point mass A, is at the centre of the hollow thin shell. In this case as every point on the shell is equidistant from A, all points exert force of equal magnitude on A but the directions of these forces are different. Now consider the forces on A due to two diametrically opposite points on the shell. The forces on A due to them will be of equal magnitude but will be in opposite directions and will cancel each other. Thus forces due to all pairs of points diametrically opposite to each other will cancel and there will be no net force on A due to the shell. When the point object is situated elsewhere inside the shell, the situation is not so symmetric. Gravitational force varies directly with mass and inversely with square of the distance. Some

part of the shell may be closer to point A, but its mass is less. Remaining part will then have larger mass but its centre of mass is away from A. However mathematically it can be shown that the net gravitational force on A is still zero, so long as it is inside the shell. In fact, the gravitational force at any point inside any hollow closed object of any shape is zero.

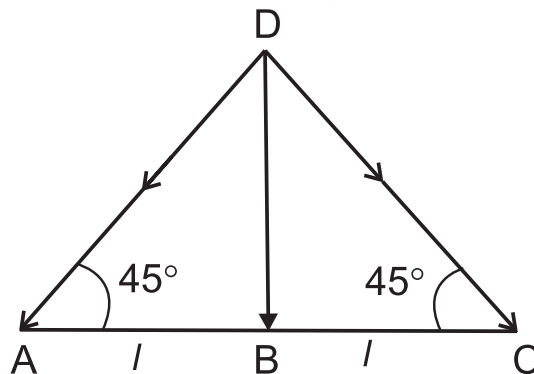
(2) The gravitational force of attraction between a hollow spherical shell or solid sphere of uniform density and a point mass situated outside is just as if the entire mass of the shell or sphere is concentrated at the centre of the shell or sphere.

Gravitational force caused by different regions of shell can be resolved into components along the line joining the point mass to the centre and along a direction perpendicular to this line. The components perpendicular to this line cancel each other and the resultant force remains along the line joining the point to the centre. By mathematical calculations it can be shown to be equal to the force that would have been exerted if the entire mass of the shell was present at the centre of the shell.

It is obvious that case (2) is applicable for any uniform sphere (solid or hollow), so long as the point is outside the sphere.

Example 5.3: Three particles A, B, and C each having mass m are kept along a straight line with $AB = BC = l$. A fourth particle D is kept on the perpendicular bisector of AC at a distance l from B. Determine the gravitational force on D.

Solution : $CD = AD = \sqrt{AB^2 + BD^2} = \sqrt{2} l$
Gravitational force on D = Vector sum of gravitational forces due to A, B and C.



Force due to A = $\frac{Gmm}{(AD)^2} = \frac{Gm^2}{2l^2}$. This will be along $\overline{DA}$

Force due to C = $\frac{Gmm}{(CD)^2} = \frac{Gm^2}{2l^2}$. This is along $\overline{DC}$

Force due to B = $\frac{Gmm}{(BD)^2} = \frac{Gm^2}{l^2}$. This is along $\overline{DB}$

We can resolve the forces along horizontal and vertical directions.

Let the unit vector along horizontal direction $\overline{AC}$ be $\hat{i}$ and along the vertical direction $\overline{BD}$ be $\hat{j}$

Net horizontal force on D

$$= \frac{Gm^2}{2l^2} \cos 45^\circ (-\hat{i}) + \frac{Gm^2}{l^2} \cos 90^\circ (\hat{i}) + \frac{Gm^2}{2l^2} \cos 45^\circ (\hat{i})$$

$$= \frac{-Gm^2}{2\sqrt{2}l^2} + \frac{Gm^2}{2\sqrt{2}l^2} = 0$$

Net vertical force on D

$$= \frac{Gm^2}{2l^2} \cos 45^\circ (-\hat{j}) + \frac{Gm^2}{l^2} (-\hat{j}) + \frac{Gm^2}{2l^2} \cos 45^\circ (-\hat{j})$$

$$= \frac{-Gm^2}{l^2} \left(\frac{1}{\sqrt{2}} + 1 \right) (\hat{j})$$

$$= \frac{Gm^2}{l^2} \left(\frac{1}{\sqrt{2}} + 1 \right) (-\hat{j})$$

$(-\hat{j})$ shows that the net force is directed along DB

5.4 Measurement of the Gravitational

Constant (G):

The magnitude of the gravitational constant G can be found by measuring the force of gravitational attraction between two bodies of masses m_1 and m_2 separated by certain distance ' L '. This can be measured by using the Cavendish balance.

The Cavendish balance consists of a light rigid rod. It is supported at the centre by a fine

vertical metallic fibre about 100 cm long. Two small spheres s_1 and s_2 of lead having equal mass m and diameter about 5 cm are mounted at the ends of the rod and a small mirror M is fastened to the metallic fibre as shown in Fig. 5.5. The mirror can be used to reflect a beam of light onto a scale and thereby measure the angle through which the wire will be twisted.

Two large lead spheres L_1 and L_2 of equal mass M and diameter of about 20 cm are brought close to the small spheres on opposite side as shown in Fig. 5.5. The big spheres attract the nearby small spheres by equal and opposite force. Let $\vec{F}$ be the force of attraction between a big sphere and small sphere near to it. Hence a torque will be generated without exerting any net force on the bar. Due to this torque the bar turns and the suspension wire gets twisted till the restoring torque due to the elastic property of the wire becomes equal to the gravitational torque.

The gravitational force between the spherical balls is the same as if their masses are concentrated at their centres. If r is the initial distance of separation between the centres of the big and the neighbouring small sphere, then the magnitude of the force between them is

$$F = G \frac{mM}{r^2}$$

If length of the rod is L , then the magnitude of the torque arising out of these forces is

$$\tau = FL = G \frac{mM}{r^2} L \quad \text{--- (5.16)}$$

At equilibrium, it is equal and opposite to the restoring torque.

$$\therefore G \frac{mM}{r^2} L = K\theta \quad \text{--- (5.17)}$$

where K is the restoring torque per unit angle and θ is the angle of twist.

By applying a known torque τ_1 and measuring the corresponding angle of twist α , the restoring torque per unit twist can be determined as $K = \tau_1/\alpha$.

Thus, in actual experiment measuring θ and knowing values of τ , m , M and r , the value of G can be calculated from Eq. (5.17). The

gravitational constant measured in this way is found to be

$$G = 6.67 \times 10^{-11} \text{ N m}^2/\text{kg}^2$$

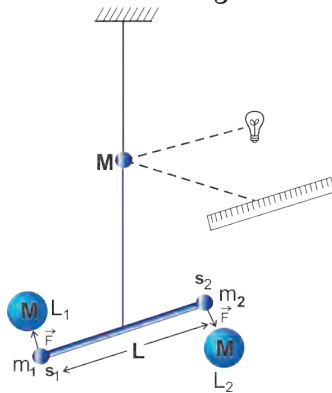


Fig 5.5 : The Cavendish balance.

5.5 Acceleration due to Gravity:

We have seen in section 5.3 that the magnitude of the gravitational force on a point object of mass m due to another point object of mass M at a distance r from it is given by the equation.

$$F = G \frac{mM}{r^2}$$

This formula can be used to calculate the gravitational force on an object due to the Earth. We know that the Earth is an extended object. In many practical applications Earth can be assumed to be a uniform sphere. As seen in section 5.3 its entire mass can be assumed to be concentrated at its centre. Thus if the mass of the Earth is M and that of the point object is m and the distance of the point object from the centre of the Earth is r then the force of attraction between them is given by

$$F = G \frac{Mm}{r^2}$$

If the point object is not acted upon by any other force, it will be accelerated towards the centre of the Earth under the action of this force. Its acceleration can be calculated by using Newton's second law $F = ma$.

Acceleration due to the gravity of the Earth =

$$\begin{aligned} & G \frac{Mm}{r^2} \times \frac{1}{m} \\ &= \frac{GM}{r^2} \end{aligned} \quad \text{--- (5.18)}$$

This is known as the acceleration due to

gravity of the Earth and denoted by g .

If the object is close to the surface of the Earth, $r \cong R$, the radius of the Earth then

$$g_{\text{Earth's surface}} = \frac{GM}{R^2} \quad \text{--- (5.19)}$$

Example 5.4: Calculate mass of the Earth from given data,

Acceleration due to gravity $g = 9.81 \text{ m/s}^2$

Radius of the Earth $R_E = 6.37 \times 10^6 \text{ m}$

$G = 6.67 \times 10^{-11} \text{ N m}^2/\text{kg}^2$

Solution:

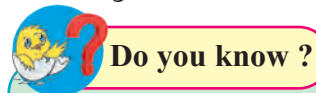
$$g = \frac{GM_E}{R_E^2}$$

$$\therefore M_E = \frac{gR_E^2}{G}$$

$$\therefore M_E = \frac{9.81 \times (6.37 \times 10^6)^2}{6.67 \times 10^{-11}}$$

$$\therefore M_E = 5.97 \times 10^{24} \text{ kg}$$

The value of g depends only on the properties of the Earth and does not depend on the mass of the object. This is exactly what Galileo had found from his experiments of dropping objects with different masses from the same height.



Do you know ?

An object of mass m (much smaller than the mass of the Earth) is attracted towards the Earth and falls on it. The Earth is also attracted by the same force (magnitude) toward the mass m . However, its acceleration towards m will be

$$\begin{aligned} a_{\text{earth}} &= \frac{\left(G \frac{Mm}{r^2} \right)}{M} = \frac{Gm}{r^2} \\ \therefore \frac{a_{\text{earth}}}{g} &= \frac{m}{M} \quad \left(\text{as } g = \frac{GM}{r^2} \right) \end{aligned}$$

As $m \ll M$, $a_{\text{Earth}} \ll g$ and is nearly zero. Thus, practically only the mass m moves towards the Earth.

Example 5.5: Calculate the acceleration due to gravity on the surface of moon if mass of the moon is 1/80 times that of the Earth and diameter of the moon is 1/4 times that of the Earth ($g = 9.8 \text{ m/s}^2$)

Solution:

M_m = Mass of the moon = $M/80$,

where M is mass of the Earth.

R_m = Radius of the moon = $R/4$,

where R is Radius of the Earth.

Acceleration due to gravity on the surface of the Earth, $g = GM/R^2$ --- (1)

Acceleration due to gravity on the surface of the moon, $g_m = GM_m/R_m^2$ --- (2)

∴ From equation (1) and (2)

$$\frac{g_m}{g} = \frac{M_m}{M} \times \left[\frac{R}{R_m} \right]^2$$

$$\frac{g_m}{g} = \frac{1}{80} \times \left[\frac{4}{1} \right]^2$$

$$\therefore \frac{g_m}{g} = \frac{1}{5}$$

$$\therefore g_m = \frac{g}{5} = \frac{9.8}{5}$$

$$\therefore g_m = 1.96 \text{ m/s}^2$$

Example 5.6: Find the acceleration due to gravity on a planet that is 10 times as massive as the Earth and with radius 20 times of the radius of the Earth ($g = 9.8 \text{ m/s}^2$).

Solution : Let mass of the planet be M_p , radius of the Earth and that of the planet be R_E and R_p respectively. Let mass of the Earth be M_E and g_p be acceleration due to gravity on the planet.

$$M_p = 10 M_E$$

$$R_p = 20 R_E, \quad g = 9.8 \text{ m/s}^2$$

$$g_p = ?$$

$$g = \frac{GM_E}{R_E^2}, \quad g_p = \frac{GM_p}{R_p^2}$$

$$\therefore g_p = \frac{G(10M_E)}{(20R_E)^2}$$

$$= \frac{10GM_E}{400R_E^2}$$

$$= \frac{1}{40} g$$

$$= 0.245 \text{ m/s}^2$$

5.6 Variation in the Acceleration due to Gravity with Altitude, Depth, Latitude and Shape:

(A) Variation in g with Altitude:

Consider a body of mass m on the surface of the Earth. The acceleration due to gravity on the Earth's surface is given by,

$$g = \frac{GM}{R^2}$$

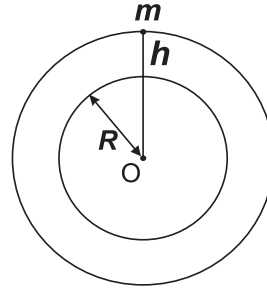


Fig. 5.6 Acceleration due to gravity at height h above the Earth's surface.

When the body is at height h above the surface of the Earth as shown in Fig. 5.6, acceleration due to gravity changes to

$$g_h = \frac{GM}{(R+h)^2}$$

$$\therefore \frac{g_h}{g} = \frac{\frac{GM}{(R+h)^2}}{\frac{GM}{R^2}}$$

$$\therefore \frac{g_h}{g} = \frac{R^2}{(R+h)^2}$$

$$\therefore g_h = \frac{g R^2}{(R+h)^2} \quad \text{--- (5.20)}$$

This equation shows that, the acceleration due to gravity goes on decreasing with increase in altitude of body from the surface of the Earth. We can rewrite Eq. (5.20) as

$$g_h = \frac{g R^2}{R^2 \left(1 + \frac{h}{R} \right)^2}$$

$$\therefore g_h = g \left(1 + \frac{h}{R} \right)^{-2}$$

For small altitude h , i.e., for $\frac{h}{R} \ll 1$,

$$\therefore g_h \approx g \left(1 - \frac{2h}{R} \right) \quad \text{--- (5.21)}$$

(by neglecting higher power terms of $\frac{h}{R}$ as $\frac{h}{R} \ll 1$)

This expression can be used to calculate the value of g at height h above the surface of the Earth as long as $h \ll R$.

Example 5.7 : At what distance above the surface of Earth the acceleration due to gravity decreases by 10% of its value at the surface? (Radius of Earth = 6400 km). Assume the distance above the surface to be small compared to the radius of the Earth.

Solution : $g_h = 90\%$ of g (g decreases by 10% hence it becomes 90%)

$$\text{or, } \frac{g_h}{g} = \frac{90}{100} = 0.9$$

From Eq. (5.21)

$$g_h = g \left[1 - \frac{2h}{R} \right]$$

$$\therefore \frac{g_h}{g} = 1 - \frac{2h}{R}$$

$$\therefore 0.9 = 1 - \frac{2h}{R}$$

$$h = \frac{R}{20}$$

$$h = 320 \text{ km}$$

(B) Variation in g with Depth:

The Earth can be imagined to be a sphere made of large number of concentric uniform spherical shells. The total mass of the Earth is the combined mass of all the shells. When an object is on the surface of the Earth it experiences the gravitational force as if the entire mass of the Earth is concentrated at its centre.

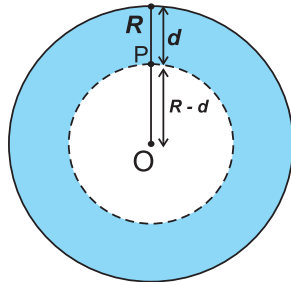


Fig. 5.7 Acceleration due to gravity at depth d below the surface of the Earth.

The acceleration due to gravity according to eq. (5.19) is

$$g = \frac{GM}{R^2}$$

Assuming that the density of the Earth is uniform, it is given by

$$\rho = \frac{\text{Mass}(M)}{\text{Volume}(V)}$$

$$\therefore M = \frac{4}{3} \pi R^3 \rho$$

$$\therefore g = \frac{G \times \frac{4}{3} \pi R^3 \rho}{R^2}$$

$$\therefore g = \frac{4}{3} \pi R \rho G \quad \text{--- (5.22)}$$

Consider a body at a point P at the depth d below the surface of the Earth as shown in Fig. 5.7. Here the force on a body at P due to the material outside the inner sphere shown by shaded region, can be shown to cancel out due to symmetry. The net force on P is only due to the material inside the inner sphere of radius $OP = R - d$. Acceleration due to gravity because of this sphere is

$$g_d = \frac{GM'}{(R-d)^2}$$

where $M' = \text{volume of the inner sphere} \times \text{density}$

$$\therefore M' = \frac{4}{3} \pi (R-d)^3 \times \rho$$

$$\therefore g_d = \frac{G \times \frac{4}{3} \pi (R-d)^3 \rho}{(R-d)^2}$$

$$\therefore g_d = G \times \frac{4}{3} \pi (R-d) \rho \quad \text{--- (5.23)}$$

Dividing Eq. (5.23) by Eq. (5.22) we get,

$$\therefore \frac{g_d}{g} = \frac{R-d}{R}$$

$$\therefore \frac{g_d}{g} = 1 - \frac{d}{R}$$

$$\therefore g_d = g \left[1 - \frac{d}{R} \right] \quad \text{--- (5.24)}$$

This equation gives acceleration due to gravity at depth d below the Earth's surface.

It shows that the acceleration due to gravity decreases with depth.

Special case :

At the centre of the Earth, where $d = R$, Eq. (5.24) gives $g_d = 0$

Hence, a body of mass m if taken to the centre of Earth, will not experience the force of gravity due to the Earth. This can also be understood to be due to symmetry. The case is similar to the force of gravity on an object placed at the centre of a spherical shell as seen in section 5.3.

Thus, the value of acceleration due to gravity is maximum at the surface of the Earth. The value goes on decreasing with

- 1) increase in depth below the Earth's surface. (varies linearly with $(R-d) = r$)
- 2) increase in height above the Earth's surface. (varies inversely with $(R+h)^2 = r^2$)

Graphically the variation of acceleration due to gravity according to depth and height can be expressed as follows. We have plotted the value of g as a function of r , the distance from the centre of the Earth, in Fig. 5.8. For $r < R$ i. e. below the surface of the Earth, we use

Eq. (5.24), according to which $g_d = g \left(1 - \frac{d}{R}\right)$

Writing $R - d = r$, the distance from the centre of the Earth, we get the value of g as a function of r , $g(r) = g \frac{r}{R}$ which is the equation of a straight line with slope g/R and passing through the origin.

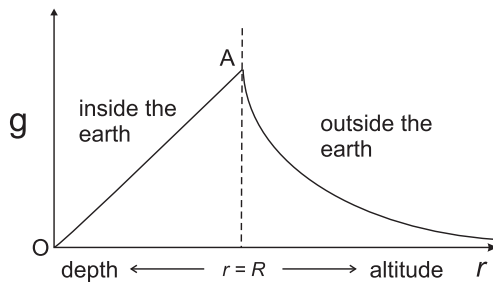


Fig 5.8 - Variation of g due to depth and altitude from the Earth's surface.

For $r > R$ we have to use Eq. (5.20). Writing $R + h = r$ we have

$g(r) = \frac{gR^2}{r^2}$ which is plotted in Fig. 5.8

(C) Variation in g with Latitude and Rotation of the Earth:

Latitude is an angle made by radius vector of any point from centre of the Earth with the equatorial plane. Obviously it ranges from 0° at the equator to 90° at the poles.

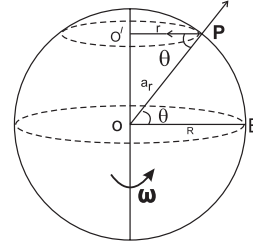


Fig. 5.9 Variation of g with latitude.

The Earth rotates about its polar axis from west to east with uniform angular velocity ω . Hence every point on the surface of the Earth (except the poles) moves in a circle parallel to the equator. The motion of a mass m at point P on the Earth is shown by the dotted circle with centre at O' in Fig. 5.9. Let the latitude of P be θ and radius of the circle be r .

$$PO' = r$$

$\angle EOP = \theta$, E being a point on the equator

$$\therefore \angle OPO' = \theta$$

$$\text{In } \Delta OPO', \cos\theta = \frac{PO'}{PO} = \frac{r}{R}$$

$$\therefore r = R\cos\theta$$

The centripetal acceleration for the mass m , directed along PO' is

$$a = r\omega^2$$

$$a = R\omega^2\cos\theta$$

The component of this centripetal acceleration along PO , i.e., towards the centre of the Earth is

$$a_r = a\cos\theta$$

$$\therefore a_r = R\omega^2\cos\theta.\cos\theta$$

$$a_r = R\omega^2\cos^2\theta$$

Part of the gravitational force of attraction on P acting towards PO is utilized in providing this components of centripetal acceleration.

Thus the effective force of gravitational attraction on m at P can be written as

$$mg' = mg - mR\omega^2\cos^2\theta$$

g' being the effective acceleration due to gravity at P i.e., at latitude θ . This is thus given

$$\text{by } g' = g - R\omega^2 \cos^2 \theta \quad \text{--- (5.25)}$$

As the value of θ increases, $\cos \theta$ decreases. Therefore g' will increase as we move away from equator towards any pole due to the rotation of the Earth.

special case I At equator $\theta = 0$

$$\begin{aligned} \cos \theta &= 1 \\ g' &= g - R\omega^2 \end{aligned}$$

The effective acceleration due to gravity is minimum at equator, as here it is reduced by maximum amount. The reduction here is $g - g' = R\omega^2$

$R = 6.4 \times 10^6 \text{ m}$ --- Radius of the Earth and

ω = Angular velocity of rotation of the Earth

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

$$\therefore \omega = \frac{2\pi}{24 \times 60 \times 60}$$

$$\therefore \omega = 7.275 \times 10^{-5} \text{ s}^{-1}$$

$$\therefore g - g' = R\omega^2$$

$$\therefore g - g' = 0.03386 \text{ m/s}^2$$

Case II At poles $\theta = 90^\circ$

$$\begin{aligned} \cos \theta &= 0 \\ \therefore g' &= g - R\omega^2 \cos \theta \\ &= g - 0 \\ &= g \end{aligned}$$

There is no reduction in acceleration due to gravity at poles, due to the rotation of the Earth as the poles are lying on the axis of rotation and do not revolve.

Variation of g with latitudes at sea level is given in the following table.

Table 5.2: Variation of g with latitude

Latitude ($^\circ$)	g (m/s ²)
0	9.7804
10	9.7819
20	9.7864
30	9.7933
40	9.8017
50	9.8107
60	9.8192
70	9.8261
80	9.8306
90	9.8322

Effect of the shape of the Earth: Quite often we assume the Earth to be a sphere. However, it is actually an ellipsoid; bulged at equator. Hence equatorial radius of Earth (6378 km) is greater than the polar radius (6356 km). Thus, on the equator, there is combined effect of greater radius and rotation in reducing the force of gravity. As a result, the acceleration due to gravity on the equator is $g_E = 9.7804 \text{ m/s}^2$ and on the poles it is $g = 9.8322 \text{ m/s}^2$.

Weight of an object is the force with which the Earth attracts that object. Thus, weight $w = mg$ where m is the mass of the object. As the value of g changes with altitude, depth and latitude, the weight also changes. Weight of an object is minimum at the equator. Similarly, the weight of an object reduces with increasing height above the Earth's surface and with increasing depth below its surface.

5.7 Gravitational Potential and Potential Energy:

In earlier standards, you have studied potential energy as the energy possessed by an object on account of its position or configuration. The word *configuration* corresponds to the distribution of the particles in the object. More specifically, potential energy is the work done *against conservative force* (or forces) *in achieving* a certain position or configuration *of a given system*. It always depends upon the relative positions of the particles in that system. There is a universal principle that states, **Every system always configures itself in order to have minimum potential energy or every system tries to minimize its potential energy.** Obviously, in order to change the configuration, you will have to do work.

Examples:

(I) A spring in its natural state, possesses minimum potential energy. Whenever we stretch it or compress it, we perform work against the conservative force (in this case, the elastic restoring force). Due to this work, the relative distances between the particles of the system change (configuration changes) and its potential energy increases. The spring finally regains its original configuration of minimum potential energy on removal of the applied force.

(II) When an object is lying on the Earth, the system of that object and the Earth has minimum potential energy. This is the gravitational potential energy of the system as these two are bound by the gravitational force. While lifting the object to some height (new position), we do work against the conservative gravitational force in order to achieve the new position.

In its new position, the object is at rest due to balanced forces. If you are holding the object, the force of static friction between the object and your fingers balances the gravitational force. If kept on a surface, the normal reaction force given by the surface balances the gravitational force. However, now, the object has a **capacity** to acquire kinetic energy, when given an opportunity (when allowed to fall). We call this increase in the capacity as the potential energy *gained* by the system. As we raise it more and more, this capacity, and hence potential energy of the system, increases. It falls on the Earth to achieve the configuration of minimum potential energy on dropping it from the new position. Thus, in general, we can write **work done against a conservative force acting on an object = Increase in the potential energy of the system.**

$$\therefore \vec{F} \cdot d\vec{x} = dU$$

Here dU is the change in potential energy while displacing the object through $d\vec{x}$, $\vec{F}$ being the force acting on the object

It should be remembered that potential energy is *always* of the system as a whole. For an object on the Earth, it is of the system of the object and the Earth and not only of that object. There is no meaning to potential energy of an isolated object in the intergalactic (gravity free) space, in the absence of any conservative force acting upon it.

5.7.1 Expression for Gravitational Potential Energy:

Work done against gravitational force $\vec{F}_g$, in displacing an object through a small displacement $d\vec{r}$, appears as increase in the potential energy of the system.

$$\therefore dU = -\vec{F}_g \cdot d\vec{r}$$

Negative sign appears because dU is the

work done by us (external agent) *against* the gravitational force $\vec{F}_g$

For displacement of the object from an initial position $\vec{r}_i$ to the final position $\vec{r}_f$, the change in potential energy ΔU , can be obtained by integrating dU .

$$\therefore \Delta U = \int_{r_i}^{r_f} (dU) = \int_{r_i}^{r_f} (-\vec{F}_g \cdot d\vec{r})$$

Gravitational force of the Earth, $\vec{F}_g = -\frac{GMm}{r^2} \hat{r}$ where $\hat{r}$ is the unit vector in the direction of $\vec{r}$. Negative sign appears here because $\vec{r}$ is from centre of the Earth to the object and $\vec{F}_g$ is directed towards centre of the Earth.

$\therefore$ For 'Earth and mass' system,

$$\begin{aligned} \Delta U &= \int_{r_i}^{r_f} dU = \int_{r_i}^{r_f} -\left(-\frac{GMm}{r^2} \hat{r}\right) \cdot d\vec{r} \\ &= GMm \int_{r_i}^{r_f} \left(\frac{dr}{r^2}\right) \text{ as } d\vec{r} \text{ is along } \hat{r} \\ &= GMm \left(-\frac{1}{r}\right)_{r_i}^{r_f} \\ &= GMm \left(\frac{1}{r_i} - \frac{1}{r_f}\right) \quad \text{--- (5.26)} \end{aligned}$$

Change in potential energy corresponds to the work done against conservative forces.

The absolute value of potential energy is not defined. It is logical as well as convenient to choose the point of zero potential energy to be the point of zero force. For gravitational force, such point is taken at $r = \infty$. This point should be chosen as the initial point so that initially the potential energy is zero. $\therefore U(r_i) = 0$ at $r_i = \infty$ Final point r_f is obviously the point where we need to determine the potential energy of the system. $\therefore r_f = r$

$$\begin{aligned} \therefore U(r)_{\text{grav.}} &= GMm \left(\frac{1}{r_i} - \frac{1}{r_f}\right) \\ &= GMm \left(\frac{1}{\infty} - \frac{1}{r}\right) \\ &= -\frac{GMm}{r} \quad \text{--- (5.27)} \end{aligned}$$

This is gravitational potential energy of the system of object of mass m and Earth of mass M having separation r (between their centres of mass).

Example 5.8: What will be the change in potential energy of a body of mass m when it is raised from height R_E above the Earth's surface to $5/2 R_E$ above the Earth's surface? R_E and M_E are the radius and mass of the Earth respectively.

Solution:

$$\begin{aligned}\Delta U &= GmM_E \left(\frac{1}{r_i} - \frac{1}{r_f} \right) \\ &= GmM_E \left(\frac{1}{2R_E} - \frac{1}{3.5R_E} \right) \\ &= \frac{GmM_E}{R_E} \times \frac{1.5}{2 \times 3.5} = \frac{2.14 GmM_E}{R_E}\end{aligned}$$

5.7.2 Connection of potential energy formula with mgh :

If the object is on the surface of Earth, $r = R$

$$\therefore U_1 = -\frac{GMm}{R}$$

If the object is lifted to height h above the surface of Earth, the potential energy becomes

$$U_2 = -\frac{GMm}{R+h}$$

Increase in the potential energy is given by

$$\begin{aligned}\Delta U &= U_2 - U_1 \\ &= GMm \left(-\frac{1}{R+h} - \left[-\frac{1}{R} \right] \right) \\ &= GMm \left(\frac{h}{R(R+h)} \right) \\ &= \frac{GMmh}{R(R+h)}\end{aligned}$$

If g is acceleration due to the Earth on the surface of Earth, $GM = gR^2$

$$\therefore \Delta U = mgh \left(\frac{R}{R+h} \right) \quad \text{--- (5.28)}$$

Eq. (5.28) gives the work to be done (or energy to be supplied) to raise an object of mass m to a height h , above the surface of the Earth.

If $h \ll R$, we can use $R+h \cong R$. Only in this case $\Delta U = mgh$ --- (5.29)

Thus, mgh is increase in the gravitational potential energy of the Earth-mass system if an object of mass m is lifted to a height h , provided h is negligible compared to radius of the Earth (up to a few kilometers).

5.7.3 Concept of Potential:

From eq. (5.27), the gravitational potential energy of the system of Earth and **any** mass m at a distance r from the centre of the Earth is given by

$$\begin{aligned}U &= -\frac{GMm}{r} \\ &= \left(-\frac{GM}{r} \right) m \\ &= \left[(V_E)_r \right] m \quad \text{--- (5.30)}\end{aligned}$$

The factor $-\frac{GM}{r} = (V_E)_r$ depends only upon

mass of Earth and the location. Thus, it is the same for any mass m bound to the Earth. Conveniently, this is defined as the gravitational potential of Earth at distance r from its centre. In terms of potential, we can write the potential energy of the Earth-mass system as

Gravitational potential energy, U = Gravitational potential $V_r \times$ mass m or Gravitational potential is Gravitational potential energy per unit mass, i.e., $V_r = \frac{U}{m}$. The concept of potential can be defined on similar lines for any conservative force field.

Gravitational potential difference between any two points in gravitational field can be written as

$$V_2 - V_1 = \left(\frac{U_2 - U_1}{m} \right) = \frac{dW}{m} \quad \text{--- (5.31)}$$

= Work done (or change in potential energy) per unit mass

In general, for a system of any two masses m_1 and m_2 , separated by r , we can write

Gravitational potential energy,

$$U = -\frac{Gm_1m_2}{r} = (V_1)m_2 = (V_2)m_1 \quad \text{--- (5.32)}$$

Here V_1 and V_2 are gravitational potentials at r due to m_1 and m_2 respectively.

5.7.4 Escape Velocity:

When any object is thrown vertically up, it falls back to the Earth after reaching a certain height. Higher the speed with which the object is thrown up, greater will be the height. If we keep on increasing the velocity, a stage will come when the object will reach heights so large that it will escape the gravitational field of the Earth and will not fall back on the Earth. This initial velocity is called the escape velocity.

Thus, the minimum velocity with which a body should be thrown vertically upwards from the surface of the Earth so that it escapes the Earth's gravitational field, is called the escape velocity (v_e) of the body. Obviously, as the gravitational force due to Earth becomes zero only at infinite distance, the object has to reach infinite distance in order to escape.

Let us consider the kinetic and potential energies of an object thrown vertically upwards with escape velocity v_e , when it is at the surface of the Earth and when it reaches infinite distance. On the surface of the Earth,

$$\begin{aligned} K.E. &= \frac{1}{2}mv_e^2 \\ P.E. &= -\frac{GMm}{R} \\ \text{Total energy} &= P.E. + K.E. \\ &= \frac{1}{2}mv_e^2 - \frac{GMm}{R} \quad \text{--- (5.33)} \end{aligned}$$

The kinetic energy of the object will go on decreasing with time as it is pulled back by Earth's gravitational force. It will become zero when it reaches infinity. Thus at infinite distance from the Earth

$$K.E. = 0$$

$$\text{Also, } P.E. = -\frac{GMm}{\infty} = 0$$

$$\therefore \text{Total energy} = P.E. + K.E. = 0$$

As energy is conserved

$$\begin{aligned} \frac{1}{2}mv_e^2 - \frac{GMm}{R} &= 0 \\ \text{or, } v_e &= \sqrt{\frac{2GM}{R}} \quad \text{--- (5.34)} \end{aligned}$$

Using the numerical values of G , M and R , the escape velocity is 11.2 km/s.

5.8 Earth Satellites:

The objects which revolve around the Earth are called Earth satellites. Moon is the only natural satellite of the Earth. It revolves in almost a circular orbit around the Earth with period of revolution of nearly 27.3 days. Artificial satellites have been launched by several countries including India. These satellites have different periods of revolution according to their practical use like navigation, surveillance, communication, looking into space and monitoring the weather.

Communication Satellites: These are geostationary satellites. They revolve around the Earth in equatorial plane. They have same sense of rotation as that of the Earth and the same period of rotation as that of the Earth, i. e., one day or 24 hours. Due to this, they appear stationary from the Earth's surface. Hence they are called geostationary satellites or geosynchronous satellites. These are used for communication, television transmission, telephones and radiowave signal transmission, e.g., INSAT group of satellites launched by India.

Polar Satellites: These satellites are placed in lower polar orbits. They are at low altitude 500 km to 800 km. Polar satellites are used for weather forecasting and meteorological purpose. They are also used for astronomical observations and study of Solar radiations.

Period of revolution of polar satellite is nearly 85 minutes, so it can orbit the Earth 16 times per day. They go around the poles of the Earth in a north-south direction while the Earth rotates in an east-west direction about its own axis. The polar satellites have cameras fixed on them. The camera can view small stripes of the Earth in one orbit. In entire day the whole Earth can be viewed strip by strip. Polar and equatorial regions at close distances can be viewed by these satellites.

5.8.1 Projection of Satellite:

For the projection of an artificial satellite, it is necessary for the satellite to have a certain velocity and a minimum two stage rocket. A single stage rocket can not achieve this. When the fuel in first stage of rocket is ignited on the surface of the Earth, it raises the satellite

vertically. The velocity of projection of satellite normal to the surface of the Earth is the vertical velocity. If this vertical velocity is less than the escape velocity (v_e), the satellite returns to the Earth's surface. While, if the vertical velocity is greater than or equal to the escape velocity, the satellite will escape from Earth's gravitational influence and go to infinity. Hence launching of a satellite in an orbit round the Earth can not take place by use of single stage rocket. It requires minimum two stage rocket.

With the help of first stage of rocket, satellite can be taken to a desired height above the surface of the Earth. Then the launcher is rotated in horizontal direction i.e. through 90° using remote control and the first stage of the rocket is detached. Then with the help of second stage of rocket, a specific horizontal velocity (v_h) is given to satellite so that it can revolve in a circular path round the Earth. The exact horizontal velocity of projection that must be given to a satellite at a certain height so that it can revolve in a circular orbit round the Earth is called the **critical velocity** or **orbital velocity** (v_c)

A satellite follows different paths depending upon the horizontal velocity provided to it. Four different possible cases are shown in Fig. 5.10.

Case (I) $v_h < v_c$:

If tangential velocity of projection v_h is less than the critical velocity, the orbit of satellite is an ellipse with point of projection as apogee (farthest from the Earth) and Earth at one of the foci.

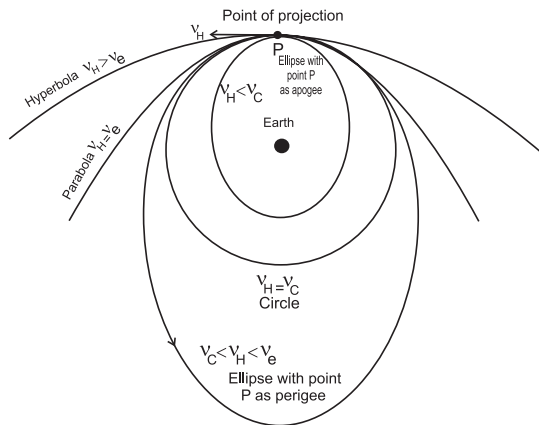


Fig. 5.10: Various possible orbits depending on the value of v_h .

During this elliptical path, if the satellite passes through the Earth's atmosphere, it experiences a nonconservative force of air resistance. As a result it loses energy and spirals down to the Earth.

Case (II) $v_h = v_c$

If the horizontal velocity is exactly equal to the critical velocity, the satellite moves in a stable circular orbit round the Earth.

Case (III) $v_c < v_h < v_e$

If horizontal velocity is greater than the critical velocity and less than the escape velocity at that height, the satellite again moves in an elliptical orbit round the Earth with the point of projection as perigee (point closest to the Earth).

Case (IV) $v_h = v_e$

If horizontal speed of projection is equal to the escape speed at that height, the satellite travels along parabolic path and never returns to the point of projection. Its speed will be zero at infinity.

Case (V) $v_h > v_e$

If horizontal velocity is greater than the escape velocity, the satellite escapes from gravitational influence of Earth transversing a hyperbolic path.

Expression for critical speed

Consider a satellite of mass m revolving round the Earth at height h above its surface.

Let M be the mass of the Earth and R be its radius. If the satellite is moving in a circular orbit of radius $(R+h) = r$, its speed must be the magnitude of critical velocity v_c .

The centripetal force necessary for circular motion of satellite is provided by gravitational force exerted by the Earth on the satellite.

$\therefore$ Centripetal force = Gravitational force

$$\frac{mv_c^2}{r} = \frac{GMm}{r^2}$$

$$\therefore v_c^2 = \frac{GM}{r}$$

$$\therefore v_c = \sqrt{\frac{GM}{r}}$$

$$\therefore v_c = \sqrt{\frac{GM}{(R+h)}} = \sqrt{g_h(R+h)} \quad \text{--- (5.35)}$$

This is the expression for critical speed in the orbit of radius $(R + h)$

It is clear that the critical speed of a satellite is independent of the mass of the satellite. It depends upon the mass of the Earth and the height at which the satellite is revolving or gravitational acceleration at that altitude. The critical speed of a satellite decreases with increase in height of satellite.

Special case

When the satellite is revolving close to the surface of the Earth, the height is very small as compared to the radius of the Earth. Hence the height can be neglected and radius of the orbit is nearly equal to R (i.e. $R \gg h$, $R+h \approx R$)

$$\therefore \text{Critical speed } v_c = \sqrt{\frac{GM}{R}}$$

As G is related to acceleration due to gravity by the relation,

$$g = \frac{GM}{R^2}$$

$$\therefore GM = gR^2$$

$\therefore$ Critical speed in terms of acceleration due to gravity can be obtained as

$$v_c = \sqrt{\frac{gR^2}{R}} = \sqrt{gR}$$

$$= 7.92 \text{ km/s}$$

Obviously, this is the maximum possible critical speed. This is at least 25 times the speed of the fastest passenger aeroplanes.

Example 5.9: Show that the critical velocity of a body revolving in a circular orbit very close to the surface of a planet of radius R and mean

density ρ is $2R\sqrt{\frac{G\pi\rho}{3}}$.

Solution : Since the body is revolving very close to the planet, $h = 0$

$$\text{density } \rho = \frac{M}{V} = \frac{M}{\frac{4}{3}\pi R^3}$$

$$\therefore M = \frac{4}{3}\pi R^3 \rho$$

Critical Velocity

$$v_c = \sqrt{\frac{GM}{R}} = \sqrt{\frac{G \frac{4}{3}\pi R^3 \rho}{R}}$$

$$\therefore v_c = 2R \sqrt{\frac{G\pi\rho}{3}}$$

When a satellite revolves very close to the surface of the Earth, motion of satellite gets affected by the friction produced due to resistance of air. In deriving the above expression the resistance of air is not considered.

5.8.2 Weightlessness in a Satellite:

According to Newton's second law of motion, $F = ma$, where F is the net force acting on an object having acceleration a .

Let us consider the example of a lift or elevator from an inertial frame of reference.

Whether the lift is at rest or in motion, a passenger in it experiences only two forces: (i) Gravitational force mg directed vertically downwards (towards centre of the earth) and (ii) normal reaction force N directed vertically upwards, exerted by the floor of the lift. As these forces are oppositely directed, the net force in the downward direction will be $F = ma - N$.

Though the weight of a body (passenger, in this case) is the gravitational force acting upon it, we experience or *feel* our weight only due to the normal reaction force N exerted by the floor. This, in turn, is equal and opposite to the relative force between the body and the lift. If you are standing on a weighing machine in a lift, the force recorded by the weighing machine is nothing but the normal reaction N .

Case I: Lift having zero acceleration

This happens when the lift is at rest or is moving upwards or downwards with constant velocity:

$$\text{The net force } F = 0 = mg - N \therefore mg = N$$

Hence in this case we feel our normal weight mg .

Case II: Lift having net upward acceleration a_u

This happens when the lift just starts moving upwards or is about to stop at a lower floor during its downward motion (remember, while stopping during downward motion, the acceleration must be upwards).

As the net acceleration is upwards, the upward force must be greater.

$\therefore F = ma_u = N - mg \therefore N = mg + ma_u$, i.e., $N > mg$, hence, we *feel* heavier.

It should also be remembered that this is not an apparent feeling. The weighing machine really records a reading greater than mg .

Case III (a): Lift having net downward acceleration a_d

This happens when the lift just starts moving downwards or is about to stop at a higher floor during its upward motion (remember, while stopping during upward motion, the acceleration must be downwards).

As the net acceleration is downwards, the downward force must be greater.

$\therefore F = ma_d = mg - N \therefore N = mg - ma_d$, i.e., $N < mg$, hence, we *feel* lighter.

It should be remembered that this is not an apparent feeling. The weighing machine really records a reading less than mg .

Case III (b): State of free fall: This will be possible if the cables of the lift are cut. In this case, the downward acceleration $a_d = g$.

If the downward acceleration becomes equal to the gravitational acceleration g , we get, $N = mg - ma_d = 0$.

Thus, there will not be any feeling of weight. This is the state of *total* weightlessness and the weighing machine will record *zero*.

In the case of a revolving satellite, the satellite is performing a circular motion. The acceleration for this motion is centripetal, which is provided by the gravitational acceleration g at the location of the satellite. In this case, $a_d = g$, or the satellite (along with the astronaut) is in the state of free fall. Obviously, the apparent weight will be zero, giving the feeling of *total* weightlessness. Perhaps you might have seen in some videos that the astronauts are floating inside the satellite. It is really difficult for them to change their position.

In spite of free fall, why is the satellite

not falling on the earth? The reason is that the revolving satellite is having a tangential velocity which manages to keep it moving in a circular orbit at that height.

5.8.3 Time Period of a Satellite:

The time taken by a satellite to complete one revolution round the Earth is its time period.

Consider a satellite of mass m projected to height h and provided horizontal velocity equal to the critical velocity. The satellite revolves in a circular orbit of radius $(R+h) = r$.

The distance traced by satellite in one revolution is equal to the circumference of the circular orbit within periodic time T .

$$\therefore \text{Critical speed} = \frac{\text{Circumference of the orbit}}{\text{Time period}}$$

$$v_c = \frac{2\pi r}{T}$$

$$\text{but we have, } v_c = \sqrt{\frac{Gm}{r}}$$

$$\therefore \sqrt{\frac{Gm}{r}} = \frac{2\pi r}{T}$$

$$\text{or, } \frac{GM}{r} = \frac{4\pi^2 r^2}{T^2}$$

$$\therefore T^2 = \frac{4\pi^2 r^3}{GM}$$

As π^2 , G and M are constant, $T^2 \propto r^3$, i.e., the square of period of revolution of satellite is directly proportional to the cube of the radius of orbit.

$$T = 2\pi \sqrt{\frac{r^3}{GM}}$$

$$\therefore T = 2\pi \sqrt{\frac{(R+h)^3}{GM}} \quad \text{--- (5.36)}$$

This is an expression for period of satellite revolving in a *circular orbit* round the Earth. Period of a satellite does not depend on its mass. It depends on mass of the Earth, radius of the Earth and the height of the satellite. If the height of projection is increased, period of the satellite increases. Period of the satellite can also be obtained in terms of acceleration due to

gravity.

$$\text{As } GM = g_h(R+h)^2$$

$$\therefore T = 2\pi \sqrt{\frac{(R+h)^3}{g_h(R+h)^2}}$$

$$\therefore T = 2\pi \sqrt{\frac{R+h}{g_h}}$$

$$\therefore T = 2\pi \sqrt{\frac{r}{g_h}} \quad \text{--- (5.37)}$$

Special case :

When satellite revolves close to the surface of the Earth, $R + h \approx R$ and $g_h \approx g$. Hence the minimum period of revolution is

$$(T)_{\min} = 2\pi \sqrt{\frac{R}{g}} \quad \text{--- (5.38)}$$

Example 5.9: Calculate the period of revolution of a polar satellite orbiting close to the surface of the Earth. Given $R = 6400$ km, $g = 9.8$ m/s².

Solution : h is negligible as satellite is close to the Earth surface.

$$\therefore R + h \approx R$$

$$g_h \approx g$$

$$R = 6400 \text{ km} = 6.4 \times 10^6 \text{ m.}$$

$$T = 2\pi \sqrt{\frac{R}{g}}$$

$$= 2 \times 3.14 \sqrt{\frac{6.4 \times 10^6}{9.8}}$$

$$= 5.075 \times 10^3 \text{ second}$$

$$= 85 \text{ minute (approximately)}$$

Example 5.10: An artificial satellite revolves around a planet in circular orbit close to its surface. Obtain the formula for period of the satellite in terms of density ρ and radius R of planet.

Solution : Period of satellite is given by,

$$T = 2\pi \sqrt{\frac{(R+h)^3}{GM}} \quad \text{--- (1)}$$

Here, the satellite revolves close to the surface of planet, hence h is negligible, hence $R + h \approx R$

$$\text{density } (\rho) = \frac{\text{mass } (M)}{\text{volume } (V)}$$

$$\therefore M = \rho V \quad \text{--- (2)}$$

As planet is spherical in shape, volume of planet is given as

$$V = \frac{4}{3} \pi R^3$$

$$\therefore M = \frac{4}{3} \pi R^3 \rho \quad \text{--- (3)}$$

Substituting the values from eq. (2) and (3) in Eq. (1), we get

$$T = 2\pi \sqrt{\frac{R^3}{G \times \frac{4}{3} \pi R^3 \rho}}$$

$$\therefore T = \sqrt{\frac{3\pi}{G\rho}}$$

5.8.4 Binding Energy of an orbiting satellite:

The minimum energy required by a satellite to escape from Earth's gravitational influence is the binding energy of the satellite.

Expression for Binding Energy of satellite revolving in circular orbit round the Earth

Consider a satellite of mass m revolving at height h above the surface of the Earth in a circular orbit. It possesses potential energy as well as kinetic energy. Let M be the mass of the Earth, R be the Radius of the Earth, v_c be critical velocity of satellite, $r = (R+h)$ be the radius of the orbit.

$\therefore$ Kinetic energy of satellite

$$= \frac{1}{2} m v_c^2$$

$$= \frac{1}{2} \frac{GMm}{r}$$

$$\text{--- (5.39)}$$

The gravitational potential at a distance r

from the centre of the Earth is $-\frac{GM}{r}$

$\therefore$ Potential energy of satellite = Gravitational potential $\times$ mass of satellite

$$= -\frac{GMm}{r}$$

$$\text{--- (5.40)}$$

The total energy of satellite is given as

$$T.E. = K.E. + P.E.$$

$$\begin{aligned}
 &= \frac{1}{2} \frac{GMm}{r} - \frac{GMm}{r} \\
 &= -\frac{1}{2} \frac{GMm}{r} \quad \text{--- (5.41)}
 \end{aligned}$$

Total energy of a circularly orbiting satellite is negative. Negative sign indicates that the satellite is bound to the Earth, due to gravitational force of attraction. For the satellite to be free from the Earth's gravitational

influence its total energy should become non-negative (zero or positive). Hence the minimum energy to be supplied to unbind the satellite is $+\frac{1}{2} \frac{GMm}{r}$. This is the binding energy of a satellite.



Internet my friend

hyperphysics.phy-astr.gsu.edu/hbase/grav.html#grav



Exercises

1. Choose the correct option.

- i) The value of acceleration due to gravity is maximum at _____.
 - (A) the equator of the Earth .
 - (B) the centre of the Earth.
 - (C) the pole of the Earth.
 - (D) slightly above the surface of the Earth.
- ii) The weight of a particle at the centre of the Earth is _____.
 - (A) infinite.
 - (B) zero.
 - (C) same as that at other places.
 - (D) greater than at the poles.
- iii) The gravitational potential due to the Earth is minimum at _____.
 - (A) the centre of the Earth.
 - (B) the surface of the Earth.
 - (C) a points inside the Earth but not at its centre.
 - (D) infinite distance.
- iv) The binding energy of a satellite revolving around planet in a circular orbit is 3×10^9 J. Its kinetic energy is _____.
 - (A) 6×10^9 J
 - (B) -3×10^9 J
 - (C) -6×10^9 J
 - (D) 3×10^9 J

2. Answer the following questions.

- i) State Kepler's law equal of area.
- ii) State Kepler's law of period.
- iii) What are the dimensions of the universal gravitational constant?
- iv) Define binding energy of a satellite.
- v) What do you mean by geostationary satellite?
- vi) State Newton's law of gravitation.
- vii) Define escape velocity of a satellite.
- viii) What is the variation in acceleration due to gravity with altitude?
- ix) On which factors does the escape speed of a body from the surface of Earth depend?
- x) As we go from one planet to another planet, how will the mass and weight of a body change?
- xi) What is periodic time of a geostationary satellite?
- xii) State Newton's law of gravitation and express it in vector form.
- xiii) What do you mean by gravitational constant? State its SI units.
- xiv) Why is a minimum two stage rocket necessary for launching of a satellite?
- xv) State the conditions for various possible orbits of a satellite depending upon the tangential speed of projection.

2. Answer the following questions in detail.

- Derive an expression for critical velocity of a satellite.
- State any four applications of a communication satellite.
- Show that acceleration due to gravity at height h above the Earth's surface is

$$g_h = g \left(\frac{R}{R+h} \right)^2$$
- Draw a labelled diagram to show different trajectories of a satellite depending upon the tangential projection speed.
- Derive an expression for binding energy of a body at rest on the Earth's surface.
- Why do astronauts in an orbiting satellite have a feeling of weightlessness?
- Draw a graph showing the variation of gravitational acceleration due to the depth and altitude from the Earth's surface.
- At which place on the Earth's surface is the gravitational acceleration maximum? Why?
- At which place on the Earth surface the gravitational acceleration minimum? Why?
- Define the binding energy of a satellite. Obtain an expression for binding energy of a satellite revolving around the Earth at certain attitude.
- Obtain the formula for acceleration due to gravity at the depth ' d ' below the Earth's surface.
- State Kepler's three laws of planetary motion.
- State the formula for acceleration due to gravity at depth ' d ' and altitude ' h '

Hence show that their ratio is equal to $\left(\frac{R-d}{R-2h} \right)$ by assuming that the altitude is very small as compared to the radius of the Earth.

- What is critical velocity? Obtain an expression for critical velocity of an orbiting satellite. On what factors does it depend?
- Define escape speed. Derive an expression for the escape speed of an object from the surface of the earth.
- Describe how an artificial satellite using two stage rocket is launched in an orbit around the Earth.

4. Solve the following problems.

- At what distance below the surface of the Earth, the acceleration due to gravity decreases by 10% of its value at the surface, given radius of Earth is 6400 km.

[Ans: 640 km].

- If the Earth were made of wood, the mass of wooden Earth would have been 10% as much as it is now (without change in its diameter). Calculate escape speed from the surface of this Earth.

[Ans: 3.54 km/s]

- Calculate the kinetic energy, potential energy, total energy and binding energy of an artificial satellite of mass 2000 kg orbiting at a height of 3600 km above the surface of the Earth.

Given:- $G = 6.67 \times 10^{-11} \text{ Nm}^2/\text{kg}^2$

$R = 6400 \text{ km}$

$M = 6 \times 10^{24} \text{ kg}$

[Ans: $KE = 40.02 \times 10^9 \text{ J}$,

$PE = -80.09 \times 10^9 \text{ J}$,

$TE = 40.02 \times 10^9 \text{ J}$,

$BE = 40.02 \times 10^9 \text{ J}$]

- Two satellites A and B are revolving around a planet. Their periods of revolution are 1 hour and 8 hours respectively. The radius of orbit of satellite B is $4 \times 10^4 \text{ km}$. find radius of orbit of satellite A.

[Ans: $1 \times 10^4 \text{ km}$]

- v) Find the gravitational force between the Sun and the Earth.
 Given Mass of the Sun = 1.99×10^{30} kg
 Mass of the Earth = 5.98×10^{24} kg
 The average distance between the Earth and the Sun = 1.5×10^{11} m.
 [Ans: 3.5×10^{22} N]
- vi) Calculate the acceleration due to gravity at a height of 300 km from the surface of the Earth. ($M = 5.98 \times 10^{24}$ kg, $R = 6400$ km).
 [Ans :- 8.889 m/s^2]
- vii) Calculate the speed of a satellite in an orbit at a height of 1000 km from the Earth's surface. $M_E = 5.98 \times 10^{24}$ kg, $R = 6.4 \times 10^6$ m.
 [Ans : $7.34 \times 10^3 \text{ m/s}$]
- viii) Calculate the value of acceleration due to gravity on the surface of Mars if the radius of Mars = 3.4×10^3 km and its mass is 6.4×10^{23} kg.
 [Ans : 3.69 m/s^2]
- ix) A planet has mass 6.4×10^{24} kg and radius 3.4×10^6 m. Calculate energy required to remove an object of mass 800 kg from the surface of the planet to infinity.
 [Ans : 5.02×10^{10} J]
- x) Calculate the value of the universal gravitational constant from the given data. Mass of the Earth = 6×10^{24} kg, Radius of the Earth = 6400 km and the acceleration due to gravity on the surface = 9.8 m/s^2
 [Ans : $6.69 \times 10^{-11} \text{ N m}^2/\text{kg}^2$]
- xi) A body weighs 5.6 kg wt on the surface of the Earth. How much will be its weight on a planet whose mass is 7 times the mass of the Earth and radius twice that of the Earth's radius.
 [Ans: 9.8 kg-wt]
- xii). What is the gravitational potential due to the Earth at a point which is at a height of $2R_E$ above the surface of the Earth, Mass of the Earth is 6×10^{24} kg, radius of the Earth = 6400 km and $G = 6.67 \times 10^{-11} \text{ Nm}^2 \text{ kg}^{-2}$.
 [Ans: $2.08 \times 10^7 \text{ J}$]



Can you recall?

1. Can you name a few objects which change their shape and size on application of a force and regain their original shape and size when the force is removed ?
2. Can you name objects which do not regain their original shape and size when the external force is removed?

6.1 Introduction:

Solids are made up of atoms or a group of atoms placed in a definite geometric arrangement. This arrangement is decided by nature so that the resultant force acting on each constituent due to others is zero. This is the equilibrium state of a solid at room temperature. The given equilibrium arrangement does not change with time. It can change only when an external stimulus, like compressive force from all sides, is applied to a solid. The constituents vibrate about their equilibrium positions even at very low temperatures but cannot leave their fixed positions. This fact provides the solids a definite shape and size (allows the solids to maintain a definite shape and size).

If an external force is applied to a solid the constituents are slightly displaced and restoring forces are developed in it. These restoring forces try to bring the constituents back to their equilibrium positions so that the solid can regain its shape. When the deforming forces are removed, the interatomic forces tend to restore the original positions of the molecules and thus the body regains its original shape and size. However, as we will see later, this is possible only within certain limits.

The form of a body is decided by its size and shape, e.g., a tennis ball and a football both are spherical, i.e., they have the same shape. But a tennis ball is smaller in size than a football. When a force is applied to a solid (which is not free to move), the size or shape or both change due to changes in the relative positions of molecules. Such a force is called **deforming force**.

The change in shape or size or both of a body due to an external force is called deformation.

The larger the deforming force on a body,

the larger is its deformation. Deformation could be in the form of change in length of a wire, change in volume of an object or change in shape of a body.

We know that when a deforming force (e.g. stretching) is applied to a rubber band, it gets deformed (elongated) but when the force is removed, it regains its original length. When a similar force is applied to a dough, or clay it also gets deformed but it does not regain its original shape and size after removal of the deforming force. These observations indicate that rubber and clay are different in nature. The property that decides this nature is called **elasticity/plasticity**. We will learn more about these properties of solids in this Chapter .

6.2 Elastic Behavior of Solids:

If a body regains its original shape and size after removal of the deforming force, it is called an elastic body and the property is called elasticity. Here the restoring forces are strong enough to bring the displaced molecules to their original positions. Examples of elastic materials are metals, rubber, quartz, etc.

If a body regains its original shape and size completely and instantaneously upon removal of the deforming force, then it is said to be **perfectly elastic**.

If a body does not regain its original shape and size and retains its altered shape or size upon removal of the deforming force, it is called a **plastic body** and the property is called **plasticity**. Here, the restoring forces are not strong enough to bring the molecules back to their original positions. Examples of plastic materials are clay, putty, plasticine, thick mud, etc. There is no solid which is perfectly elastic or perfectly plastic. The best example of a near ideal elastic solid is quartz fibre and that of a plastic body is putty.

6.3 Stress and Strain:

The elastic properties of a body are described in terms of stress and strain. When a body gets deformed under an applied force, restoring forces are set up internally. They oppose change in shape or size of the body. When body is in equilibrium in its altered shape or size, deforming force and restoring force are equal and opposite.

The internal restoring force per unit area of a body is called stress.

$$\text{stress} = \frac{\text{deforming force}}{\text{area}} = \frac{|\vec{F}|}{A} \quad \text{--- (6.1)}$$

where $\vec{F}$ is internal restoring force (external applied deforming force). SI unit of stress is N m^{-2} or pascal (Pa). The dimensions of a stress are $[L^{-1} M^1 T^{-2}]$.

Strain is a measure of the deformation of a body. When two equal and opposite forces are applied to an elastic body, there is a change in the dimensions of the body, **Strain is defined as the ratio of change in dimensions of the body to its original dimensions.**

$$\text{Strain} = \frac{\text{change in dimensions}}{\text{original dimensions}} \quad \text{--- (6.2)}$$

It is the ratio of two similar quantities. Hence strain is a dimensionless physical quantity. It has no units. There are three types of stress and corresponding strains.

1: Stress produced by a deforming force acting along the length of a body or a rod is called **tensile stress or a longitudinal stress**. The strain produced is called tensile strain.

A) Tensile stress or compressive stress:

Suppose a force $\vec{F}$ is applied along the length of a wire, or perpendicular to its cross section A . This produces an elongation in the wire and the length of the wire increases accordingly, as shown in Fig. 6.1 (a).

$$\text{Tensile stress} = \frac{|\vec{F}|}{A} \quad \text{--- (6.3)}$$

When a rod is pushed at two ends with equal and opposite forces, its length decreases. The restoring force per unit area is called **compressive stress** as shown in Fig. 6.1 (b).

$$\text{Compressive stress} = \frac{|\vec{F}|}{A} \quad \text{--- (6.4)}$$

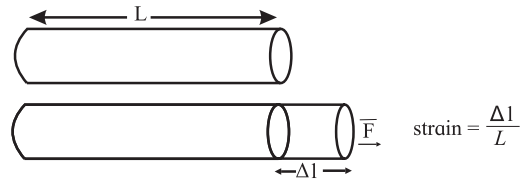


Fig. 6.1 (a): Tensile stress.

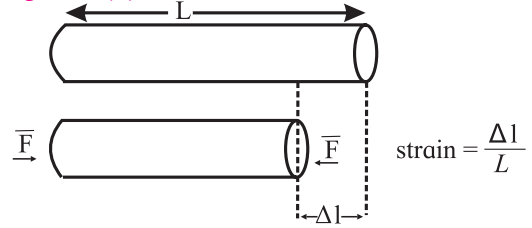


Fig. 6.1 (b): Compressive stress.

B) Tensile strain:

The strain produced by a tensile deforming force is called tensile strain or longitudinal strain or linear strain.

If L is the original length and Δl is the change in length due to the deforming force, then

$$\text{Tensile strain} = \frac{\Delta l}{L} \quad \text{--- (6.5)}$$

2 : When a deforming force acting on a body produces change in its volume, the stress is called volume stress and the strain produced is called **volume strain**.

A) Volume stress or hydraulic stress:

Let $\vec{F}$ be a force acting perpendicular to the entire surface of the body. It acts normally and uniformly all over the surface area A of the body. Such a stress which produces change in size but no change in shape is called volume stress.

$$\text{Volume stress} = \frac{|\vec{F}|}{A} \quad \text{--- (6.6)}$$

Volume stress produces change in size without change in shape of body, it is called hydraulic or hydrostatic volume stress as shown in Fig. 6.2.

B) Volume strain:

A deforming force acting perpendicular to the entire surface of a body produces a volume strain. Let V be the original volume and ΔV be the change in volume due to deforming force, then

$$\text{Volume strain} = \frac{\Delta V}{V} \quad \text{--- (6.7)}$$

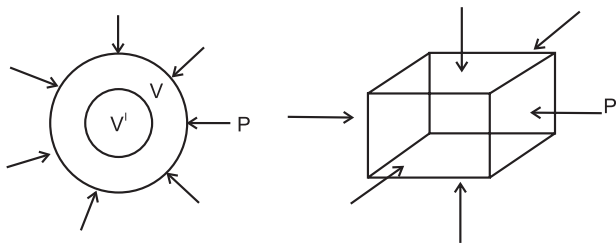


Fig. 6.2 : Volume stress.



Do you know ?

When a balloon is filled with air at high pressure, its walls experience a force from within. This is also volume stress. It tries to expand the balloon and change its size without changing shape. When the volume stress exceeds the limit of bulk elasticity, the balloon explodes. Similarly, a gas cylinder explodes when the pressure inside it exceeds the limit of bulk elasticity of its material.

A submarine when submerged under water is under volume stress.

3 : When a deforming force acting on a body produces change in the shape of a body, shearing stress and shearing strain are produced.

A) Shearing stress:

Let $\vec{F}$ be a tangential force acting on a surface area A . This force produces change in shape of the body without changing its size as shown in Fig. 6.3.

$$\text{Shearing stress} = \frac{\text{Tangential force}}{\text{Area}} \text{--- (6.8)}$$

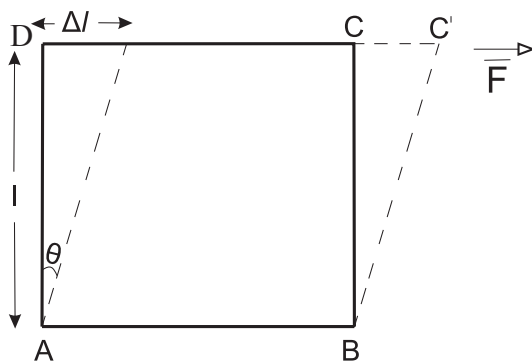


Fig. 6.3 : Tangential force produces shearing stress.

Suppose ABCD is the front face of a cube. A force $\vec{F}$ is applied to the cube so that the bottom of the cube is fixed and only the top surface is slightly displaced. Such force is called

tangential force. **Tangential force is parallel to the top and the bottom surface of the block. The restoring force per unit area developed due to the applied tangential force is called shearing stress or tangential stress.**

B) Shearing strain:

There is a relative displacement, Δl , of the bottom face and the top face of the cube. Such relative displacement of two surfaces is called shear strain. It can be calculated as follows,

$$\text{Shearing strain } \frac{\Delta l}{l} = \tan \theta = \theta \text{ --- (6.9)}$$

when the relative displacement Δl is very small.

6.4 Hooke's Law:

Robert Hooke (1635-1703), an English physicist, studied the tension in a wire and strain produced in it. His study led to a law now known as Hooke's law.

Statement: Within elastic limit, stress is directly proportional to strain.

$$\frac{\text{Stress}}{\text{Strain}} = \text{constant}$$

The constant is called the modulus of elasticity. **The modulus of elasticity of a material is the ratio of stress to the corresponding strain.** It is defined as the slope of the stress-strain curve in the elastic deforming region and depends on the nature of the material. The maximum value of stress up to which stress is directly proportional to strain is called the **elastic limit**. The stress-strain curve within elastic limit is shown in Fig. 6.4

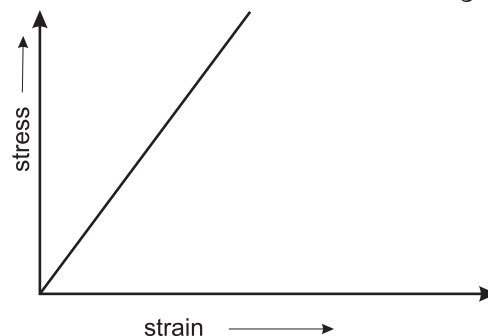


Fig 6.4: Stress versus strain graph within elastic limit for an elastic body.

6.5 Elastic modulus:

There are three types of stress and strain related to change in length, change in volume and change in shape. Hence, we have three moduli of elasticity corresponding to each type

of stress and strain.

6.5.1 Young's modulus (Y):

It is the modulus of elasticity related to change in length of an object like a metal wire, rod, beam, etc., due to the applied deforming force. Hence it is also called as elasticity of length. It is named after the British physicist Thomas Young (1773-1829).

Consider a metal wire of length L having radius r suspended from a rigid support. A load Mg is attached to the free end of the wire. Due to this, deforming force is applied at the free end of the wire in downward direction. In its equilibrium position,

$$\begin{aligned}\text{Longitudinal stress} &= \frac{\text{Applied force}}{\text{Area}} \\ &= \frac{F}{A} \\ &= \frac{Mg}{\pi r^2} \quad \text{--- (6.10)}\end{aligned}$$

It produces a change in length of the wire. If $(L+l)$ is the new length of wire, then l is the extension or elongation in wire.

$$\begin{aligned}\text{Longitudinal strain} &= \frac{\text{change in length}}{\text{original length}} \\ &= \frac{l}{L} \quad \text{--- (6.11)}\end{aligned}$$

Young's modulus is the ratio of longitudinal stress to longitudinal strain.

$$\text{Young's modulus} = \frac{\text{longitudinal stress}}{\text{longitudinal strain}} \quad \text{--- (6.12)}$$

$$\begin{aligned}Y &= \frac{\frac{Mg}{\pi r^2}}{\frac{l}{L}} \\ Y &= \frac{MgL}{\pi r^2 l} \quad \text{--- (6.13)}\end{aligned}$$

SI unit of Young's modulus is N/m^2 . Its dimensions are $[L^{-1} M^1 T^2]$.

Young's modulus indicates the resistance of an elastic solid to elongation or compression. Young's modulus of a material is useful for characterization of an object subjected to compression or tension. Young's modulus is the property of solids only.

Table 6.1: Young's modulus of some familiar materials

Material	Young's modulus $Y \times 10^{10} \text{ Pa (N/m}^2\text{)}$
Lead	1.5
Glass (crown)	6.0
Aluminium	7.0
Silver	7.6
Gold	8.1
Brass	9.0
Copper	11.0
Steel	21.0

Example 6.1: A brass wire of length 4.5m with crosssectional area of $3 \times 10^{-5} \text{ m}^2$ and a copper wire of length 5.0 m with cross sectional area $4 \times 10^{-5} \text{ m}^2$ are stretched by the same load. The same elongation is produced in both the wires. Find the ratio of Young's modulus of brass and copper.

Solution: For brass,

$$L_B = 4.5\text{m}, A_B = 3 \times 10^{-5} \text{ m}^2$$

$$l_B = l, F_B = F$$

$$Y_B = \frac{F_B L_B}{A_B l_B}$$

$$\therefore Y_B = \frac{F \times 4.5}{3 \times 10^{-5} \times l}$$

For copper,

$$L_C = 5\text{m}, A_C = 4 \times 10^{-5} \text{ m}^2$$

$$l_C = l, F_C = F$$

$$Y_C = \frac{F_C L_C}{A_C l_C}$$

$$\therefore Y_C = \frac{F \times 5.0}{4 \times 10^{-5} \times l}$$

$$\frac{Y_B}{Y_C} = \frac{F \times 4.5}{3 \times 10^{-5} \times l} \times \frac{4 \times 10^{-5} \times l}{F \times 5}$$

$$= \frac{18 \times 10^{-5}}{15 \times 10^{-5}} = 1.2$$

Example 6.2: A wire of length 20 m and area of cross section $1.25 \times 10^{-4} \text{ m}^2$ is subjected to a load of 2.5 kg. (1 kgwt = 9.8 N). The elongation produced in wire is $1 \times 10^{-4} \text{ m}$. Calculate Young's modulus of the material.

Solution: Given,

$$L = 20 \text{ m}$$

$$A = 1.25 \times 10^{-4} \text{ m}^2$$

$$F = mg = 2.5 \times 9.8 \text{ N}$$

$$L = 10^{-4} \text{ m}$$

To find: Y

$$Y = \frac{FL}{Al} = \frac{2.5 \times 9.8 \times 20}{1.25 \times 10^{-4} \times 10^{-4}} \\ = 3.92 \times 10^{10} \text{ N m}^{-2}$$

6.5.2 Bulk modulus (K):

It is the modulus of elasticity related to change in volume of an object due to applied deforming force. Hence it is also called as elasticity of volume. Bulk modulus of elasticity is a property of solids, liquids and gases.

If a sphere made from rubber is completely immersed in a liquid, it will be uniformly compressed from all sides. Suppose this compressive force is F . Let the change in pressure on the sphere be dP and let the change in its volume be dV . If the original volume of the sphere is V , then volume strain is defined as

$$\text{Volume strain} = \frac{\text{change in volume}}{\text{original volume}} \\ = -\frac{dV}{V} \quad \text{--- (6.14)}$$

The negative sign indicates that there is a decrease in volume. The magnitude of the volume strain is $\frac{dV}{V}$

Bulk modulus is defined as the ratio of volume stress to volume strain.

$$\text{Bulk modulus} = \frac{\text{volume stress}}{\text{volume strain}} = K \\ K = \frac{dP}{\left(\frac{dV}{V}\right)} = V \frac{dP}{dV} \quad \text{--- (6.17)}$$

SI unit of bulk modulus is N/m^2 . Dimensions of K are $[L^{-1} M^1 T^{-2}]$.

Table 6.2 gives bulk moduli of some familiar materials

Bulk modulus measures the resistance offered by gases, liquids or solids while an attempt is made to change their volume.

The reciprocal of bulk modulus of elasticity is called compressibility of the material.

$$\text{Compressibility} = \frac{1}{\text{Bulk modulus}} \quad \text{--- (6.18)}$$

Compressibility is the fractional decrease in volume, $-\Delta V/V$ per unit increase in pressure. SI unit of compressibility is m^2/N or Pa^{-1} and its dimensions are $[L^1 M^{-1} T^2]$.



Do you know ?

The bulk modulus of water is $2.18 \times 10^8 \text{ Pa}$ and its compressibility is $45.8 \times 10^{-10} \text{ Pa}^{-1}$. Materials with small bulk modulus and large compressibility are easier to compress.

Example 6.3: A metal cube of side 1m is subjected to a force. The force acts normally on the whole surface of cube and its volume changes by $1.5 \times 10^{-5} \text{ m}^3$. The bulk modulus of metal is $6.6 \times 10^{10} \text{ N/m}^2$. Calculate the change in pressure.

Solution: Given,

$$\text{volume of cube} = V = l^3 = (1)^3 = 1 \text{ m}^3$$

$$\text{Change in volume} = dV = 1.5 \times 10^{-5} \text{ m}^3$$

$$\text{Bulk modulus} = K = 6.6 \times 10^{10} \text{ N/m}^2.$$

To find: Change in pressure dP

$$K = V \frac{dP}{dV}$$

$$dP = K \frac{dV}{V}$$

$$dP = \frac{6.6 \times 10^{10} \times 1.5 \times 10^{-5}}{1}$$

$$dP = 9.9 \times 10^5 \text{ N/m}^2.$$

Table 6.2: Bulk modulus of some familiar materials

Material	Bulk modulus K $\times 10^{10} \text{ Pa (N/m}^2\text{)}$
Lead	4.1
Brass	6.0
Glass (crown)	6.0
Aluminium	7.5
Silver	10.0
Copper	14.0
Steel	16.0
Gold	18.0

6.5.3 Modulus of rigidity (η):

The modulus of elasticity related to change in shape of an object is called rigidity modulus. It is the property of solids only as they alone possess a definite shape.

The block shown in Fig. 6.5 is made of a uniform isotropic material. It has a uniform crosssection area A and height l . A cross section of the block is defined as any plane parallel to the top and the bottom surface and cuts the block. Two forces of magnitude ' F ' are applied along top and bottom surface as shown in Fig. (6.5). They constitute a couple. The upper surface is displaced relative to the lower surface by a small distance Δl and corresponding angles change by a small amount $\theta = \Delta l/l$.

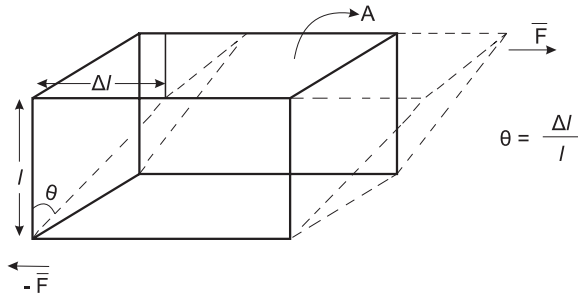


Fig. 6.5: Modulus of rigidity, tangential force F and shear strain θ .

A couple is applied by pushing the top and the bottom surfaces as shown in Fig. 6.5. Similar couple would be applied if the bottom of the block is fixed and only the top is pushed.

The forces $\vec{F}$ and $-\vec{F}$ are parallel to the cross section. This is different than the tensile stress where the force is normal to the cross section.

As a result of the way in which the forces are applied the block is subjected to a shear stress defined by **shear stress** $= F/A$.

The SI unit of shear stress is N/m^2 or Pa. The block is distorted as a result of the shear stress. The top and bottom surface are relatively displaced by a small distance Δl . The corner angle changes by a small amount θ which is called shear strain and is expressed in radian. Shear strain ' θ ' is given by $\theta = \Delta l/l$, (for small Δl).

Shear modulus or modulus of rigidity: It is defined as the ratio of shear stress to shear strain within elastic limits.

$$\eta = \frac{\text{shear stress}}{\text{shear strain}} = \frac{F/A}{\theta} = \frac{F}{A\theta} \quad \text{--- (6.17)}$$

Table 6.3 gives values of rigidity modulus η of some familiar materials.

Table 6.3: Rigidity modulus η of some familiar materials

Material	Rigidity modulus η $\times 10^{10} \text{ Pa (N/m}^2\text{)}$
Lead	0.6
Aluminium	2.5
Glass (crown)	2.5
Silver	2.7
Gold	2.9
Brass	3.5
Copper	4.4
Steel	8.3

Rigidity modulus indicates the resistance offered by a solid to change in its shape.

Example 6.4: Calculate the modulus of rigidity of a metal, if a metal cube of side 40 cm is subjected to a shearing force of 2000 N. The upper surface is displaced through 0.5cm with respect to the bottom. Calculate the modulus of rigidity of the metal.

Solution: Given,

Length of side of cube $= l = 40 \text{ cm} = 0.40 \text{ m}$

Shearing force $= F = 2000 \text{ N} = 2 \times 10^3 \text{ N}$

Displacement of top face $= \Delta l = 0.5 \text{ cm} = 0.005 \text{ m}$

Area $= A = l^2 = 0.16 \text{ m}^2$

To find: modulus of rigidity, η

$$\begin{aligned} \eta &= \frac{F}{A\theta} \\ \theta &= \frac{\Delta l}{l} = \frac{0.005}{0.40} = 0.0125 \\ \eta &= \frac{2.0 \times 10^3 \text{ N}}{(0.16 \text{ m}^2) \cdot (0.0125)} \\ &= 1.0 \times 10^6 \text{ N/m}^2 \end{aligned}$$

6.5.4 Poisson's ratio:

Suppose a wire is fixed at one end and a force is applied at its free end so that the wire gets stretched. Length of the wire increases and at the same time, its diameter decreases, i.e., the wire becomes longer and thinner as shown in Fig. 6.6 (a).

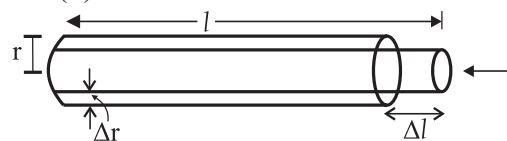


Fig. 6.6 (a): When a wire is stretched its length increases and its diameter decreases.

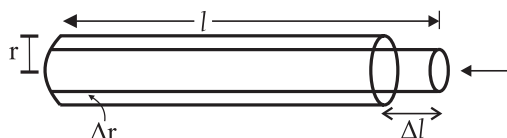


Fig. 6.6 (b): When a wire is compressed its length increases and its diameter increases.

If equal and opposite forces are applied to an object along its length inwards, the object gets compressed (Fig. 6.6 (b)). There is a decrease in dimensions along its length and at the same time there is an increase in its dimensions perpendicular to its length. When length of the wire decreases, its diameter increases.

The ratio of change in dimensions to original dimensions in the direction of the applied force is called **linear strain** while the ratio of change in dimensions to original dimensions in a direction perpendicular to the applied force is called **lateral strain**. Within elastic limit, the ratio of lateral strain to the linear strain is called the **Poisson's ratio**.

If l is the original length of wire, Δl is increase/decrease in length of wire, D is the original diameter and d is corresponding change in diameter of wire then, Poisson's ratio is given by

$$\begin{aligned}\sigma &= \frac{\text{Lateral strain}}{\text{Linear strain}} \\ &= \frac{d/D}{l/L} \\ &= \frac{d.L}{D.l}\end{aligned}\quad \text{--- (6.18)}$$

Poisson's ratio has no unit. It is dimensionless. Table 6.4 gives values of Poisson ratio, σ , of some familiar materials.

Table 6.4: Poisson ratio, σ , of some familiar materials

Material	Poisson ratio σ
Glass (crown)	0.2
Steel	0.28
Aluminium	0.36
Brass	0.37
Copper	0.37
Silver	0.38
Gold	0.42



Do you know ?

For most of the commonly used metals, the value of σ is between 0.25 and 0.35. Many times we assume that volume is constant while stretching a wire. However, in reality, its volume also increases. Using approximations it can be shown that $\sigma_{\text{max}} \approx 0.5$ if volume is unchanged. In practice, it is much less. This shows that volume also increases while stretching.

6.6 Stress-Strain Curve:

Suppose a metal wire is suspended vertically from a rigid support and stretched by applying load to its lower end. The load is gradually increased in small steps until the wire breaks. The elongation produced in the wire is measured during each step. Stress and strain is noted for each load and a graph is drawn by taking tensile strain along x-axis and tensile stress along y-axis. It is a stress-strain curve as shown in Fig. 6.7.

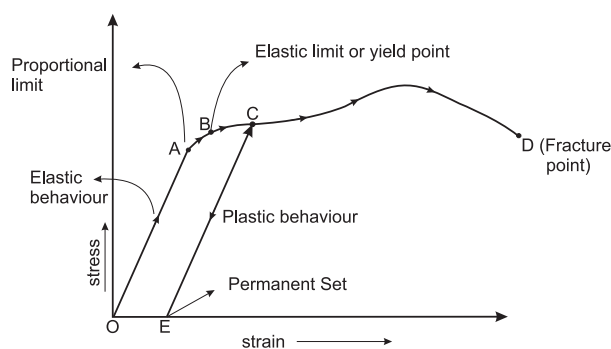


Fig. 6.7 : stress-strain curve.

The initial part of the graph is a straight line OA. This is the region in which Hooke's law is obeyed and stress is directly proportional to strain. The straight line portion ends at A. The stress at this point is called **proportional limit**. If the load is further increased till point B is reached, stress and strain are no longer proportional and Hooke's law is not valid. If the load is gradually removed starting at any point between O and B. The curve is retraced until the wire regains its original length. The change is reversible. The material of the wire shows elastic behaviour in the region OB. Point B is called the **yield point**. The corresponding point is called the **elastic limit**.

When the stress is increased beyond point B, the strain continues to increase. If the load is removed at any point beyond B, C for example, the material does not regain its original length. It follows the line CE. Length of the wire when there is no stress is greater than the original length. The deformation is irreversible and the material has acquired a **permanent set**.

Further increase in load causes a large increase in strain for relatively small increase in stress, until a point D is reached at which **fracture** takes place.

The material shows **plastic flow or plastic deformation** from point B to point D. The material does not regain its original state when the stress is removed. **The deformation is called plastic deformation.**

The curve described above shows all the possibilities for an elastic substance. In particular, many metallic wires (copper, aluminum, silver, etc) exhibit this type of behavior. However, majority of materials in every day life exhibit only some part of it.

Materials such as glass, ceramics, etc., break within the elastic limit. They are called **brittle**.

Metals such as copper, aluminum, wrought iron, etc. have large plastic range of extension. They lengthen considerably and undergo plastic deformation till they break. They are called **ductile**.

Metals such as gold, silver which can be hammered into thin sheets are called **malleable**.

Rubber has large elastic region. It can be stretched so that its length becomes many times its original length, after removal of the stress it returns to its original state but the stress strain curve is not a straight line. A material that can be elastically stretched to a larger value of strain is called an **elastomer**.

In case of some materials like vulcanized rubber, when the stress applied on a body decreases to zero, the strain does not return to zero immediately. The strain lags behind the stress. This lagging of strain behind the stress is called **elastic hysteresis**. Figure 6.8 shows the stress-strain curve for increasing and decreasing load. It encloses a loop. Area of loop gives

the energy dissipated during deformation of a material.

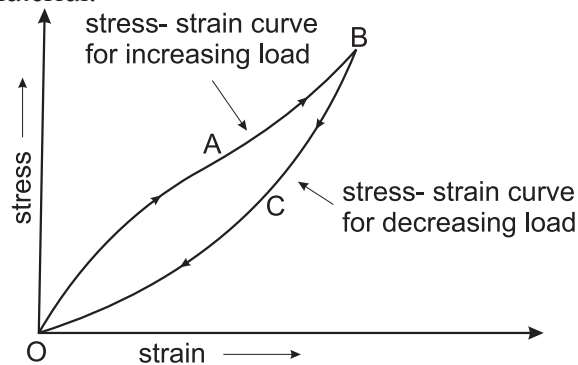


Fig. 6.8: Stress-strain curve for increasing and decreasing load.



Can you tell?

Why does a rubber band become loose after repeated use?

6.7 Strain Energy:

The elastic potential energy gained by a wire during elongation by a stretching force is called as strain energy.

Consider a wire of original length L and cross sectional area A stretched by a force F acting along its length. The wire gets stretched and elongation l is produced in it. The stress and the strain increase proportionately.

$$\text{Longitudinal stress} = \frac{F}{A}$$

$$\text{Longitudinal strain} = \frac{l}{L}$$

$$\text{Young's modulus} = \frac{\text{longitudinal stress}}{\text{longitudinal strain}}$$

$$Y = \frac{\left(\frac{F}{A}\right)}{\left(\frac{l}{L}\right)} = \frac{FL}{Al}$$

$$\therefore F = \frac{YAl}{L} \quad \text{--- (6.19)}$$

The magnitude of stretching force increases from zero to F during elongation of wire. At a certain stage, let ' f ' be the force applied and ' x ' be the corresponding extension. The force at this stage is given by Eq. (6.19) as

$$f = \frac{YAx}{L}$$

For further extension dx in the wire, the work done is given by

$$\text{Work} = (\text{force}) \cdot (\text{displacement}).$$

$$dW = f \, dx$$

$$\therefore dW = \frac{YAx}{L} dx$$

When the wire gets stretched from $x = 0$ to $x = l$, the total work done is given as

$$W = \int_0^l dW$$

$$\therefore W = \int_0^l \frac{YAx}{L} dx$$

$$\therefore W = \frac{YA}{L} \int_0^l x dx$$

$$\therefore W = \frac{YA}{L} \left[\frac{x^2}{2} \right]_0^l$$

$$\therefore W = \frac{YA}{L} \left[\frac{l^2}{2} - \frac{0^2}{2} \right]$$

$$W = \frac{YAl^2}{2L}$$

$$W = \frac{1}{2} \frac{YAl}{L} l$$

$$W = \frac{1}{2} Fl$$

$$\text{Work done} = \frac{1}{2} (\text{load}) \cdot (\text{extension}) \quad \text{--- (6.20)}$$

This work done by stretching force is equal to energy gained by the wire. This energy is strain energy.

$$\text{Strain energy} = \frac{1}{2} (\text{load}) \cdot (\text{extension}) \quad \text{--- (6.21)}$$

Strain energy per unit volume can be obtained by using Eq. (6.20) and various formula of stress, strain and young's modulus.

Work done per unit volume

$$= \frac{\text{work done in stretching wire}}{\text{volume of wire.}}$$

$$= \frac{1}{2} \frac{F \cdot l}{A \cdot L}$$

$$= \frac{1}{2} \left(\frac{F}{A} \right) \left(\frac{l}{L} \right)$$

Work done per unit volume

$$= \frac{1}{2} (\text{stress}) \cdot (\text{strain})$$

Strain energy per unit volume

$$= \frac{1}{2} (\text{stress}) \cdot (\text{strain}) \quad \text{--- (6.22)}$$

$$\text{As } Y = \frac{\text{stress}}{\text{strain}},$$

$$\text{Stress} = Y \cdot (\text{strain}) \text{ and}$$

$$\text{strain} = \frac{\text{stress}}{Y}$$

$\therefore$ Strain energy per unit volume

$$= \frac{1}{2} Y \cdot (\text{strain})^2 \quad \text{--- (6.23)}$$

Also, strain energy per unit volume

$$= \frac{1}{2} \frac{(\text{stress})^2}{Y} \quad \text{--- (6.24)}$$

Thus Eq. (6.22), (6.23) and (6.24) give strain energy per unit volume in various forms.

6.8 Hardness:

Hardness is the property of a material which enables it to resist plastic deformation. Hard materials have little ductility and they are brittle to some extent. **The term hardness also refers to stiffness or resistance to bending, scratching abrasion or cutting.** It is the property of a material which gives it the ability to resist permanent deformation when a load is applied to it. **The greater the hardness, greater is the resistance to deformation.**

The most well-known example of the hard materials is diamond. It is incredibly difficult to scratch a diamond. Metal with very low hardness is aluminium.

Hardness of material is different from its strength and toughness.

If a force is applied to a body it produces deformation in it. Higher is the force required for deformation, the stronger is the material, i.e., the material has more **strength**.

Steel has high strength whereas plasticine clay is not strong because it gets easily deformed even by a small force.

Toughness is the ability of a material to resist fracturing when a force is applied to it. Plasticine clay is relatively tough as it can be stretched and deformed due to applied force without breaking.

A single material may be hard, strong and tough, e.g.,

- 1) Bulletproof glass is hard and tough but not strong.
- 2) Drill bits must be hard, strong and tough for their work.
- 3) Anvils are very tough and strong but they are not hard.

6.9 Friction in Solids:

Whenever the surface of one body slides over another, each body exerts a certain amount of force on the other body. These forces are tangential to the surfaces. The force on each body is opposite to the direction of motion between the two bodies. It prevents or opposes the relative motion between the two bodies. It is a common experience that an object placed on any surface does not move easily when a small force is applied to it. This is because of certain force of opposition acting between the surface of the object and the surface on which it is placed. Even a rolling ball comes to rest after covering a finite distance on playground because of such opposing force. Our foot ware is provided with designs at the bottom of its sole so as to produce force of opposition to avoid slipping. It is difficult to walk without such opposing force. You know what happens when you try to walk fast on polished flooring at home with soap water spread on it. There is a possibility of slipping due to lack of force of opposition. To initiate any motion between a pair of surfaces, we need a certain minimum force. Also after the motion begins, it is constantly opposed by some natural force. This mechanical force between two solid surfaces in contact with each other is called as frictional force. **The property which resists the relative motion between two surfaces in contact is called friction.**

In some cases it is necessary to avoid friction, because friction causes dissipation of energy in machines due to which efficiency of machines decreases. In such cases friction should be reduced by using polished surfaces, lubricants, etc. Relative motion between solids and fluids (i.e. liquids and gases) is also naturally opposed by friction, e.g., a boat on the surface of water experiences opposition to its motion.

In this section we are going to study friction in solids only.

6.9.1 Origin of friction:

If smooth surfaces are observed under powerful microscope, many irregularities and projections are observed. Friction arises due to interlocking of these irregularities between two surfaces in contact. The surfaces can be made extremely smooth by polishing to avoid irregularities but it is noticed that in this case also, friction does not decrease but may increase. Hence the interlocking of irregularities is not the real cause of friction.

According to modern theory, cause of friction is the force of attraction between molecules of two surfaces in actual contact in addition to the force due to the interlocking between the two surfaces. When one body is in contact with another body, the real microscopic area in contact is very small due to irregularities in contact. Figure 6.9 shows the microscopic view of two polished surfaces in contact.

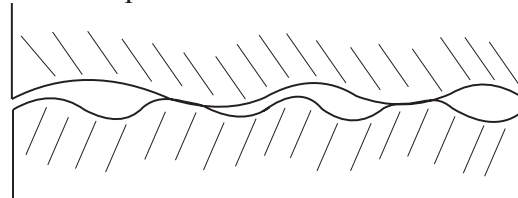


Fig. 6.9: Microscopic view of polished surfaces in contact.

Due to small area, pressure at points of contact is very high. Hence there is a strong force of attraction between the surfaces in contact. If both the surfaces are of the same material the force of attraction is called **cohesive force** while if the surfaces are of different materials the force of attraction is called **adhesive force**. When the surfaces in contact become more and more smooth, the actual area of contact goes on increasing. Due to this, the force of attraction between the molecules increases and hence the friction also increases. Putting some grease or other lubricant (a different material) between the two surfaces reduces the friction.

6.9.2 Types of friction:

1. Static friction:

Suppose a wooden block is placed on a horizontal surface as shown in Fig 6.10. A small horizontal force F is applied to it. The

block does not move with this force as it cannot overcome the frictional force between the block and horizontal surface. In this case, the force of static friction is equal to F and balances it. **The frictional force which balances applied force when the body is static is called force of static friction, F_s . In other words, static friction prevents sliding motion.**

If we keep increasing $\underline{F}$, a stage will come when, for $\underline{F} = \underline{F}_{\max}$, the object will start moving. For $\underline{F} < \underline{F}_{\max}$, the force of static friction is equal to $\underline{F}$. For $\underline{F} \geq \underline{F}_{\max}$, the kinetic friction comes into play. Static friction opposes impending motion i.e. the motion that would take place in absence of frictional force under the applied force.

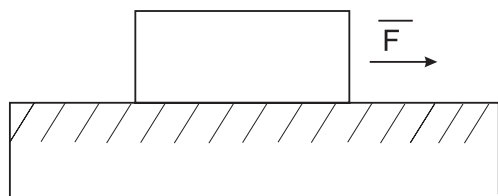


Fig. 6.10: Static friction.

The force of static friction is self adjusting force. When the applied force $\underline{F}$ is very small, the block remains at rest. Here the force of friction is also small. When $\underline{F}$ is increased by a small value, the block remains still at rest as the force of friction is increased to balance the applied force. If applied force is increased, the friction also increases and reaches the maximum value.

Just before the body starts sliding over another body, the value of frictional force is maximum, it is called **limiting force of friction, F_L** . If the direction of applied force is reversed, the direction of static friction is also reversed, i.e., it adjusts its direction also.

Laws of static friction:

- 1] The limiting force of static friction is directly proportional to the normal reaction (N) between the two surfaces in contact.

$$F_L \propto N$$

$$\therefore F_L = \mu_s N \quad \text{--- (6.25)}$$

Where μ_s is constant of proportionality. It is called as **coefficient of static friction**.

$$\therefore \mu_s = \frac{F_L}{N} \quad \text{--- (6.26)}$$

The coefficient of static friction is defined as the ratio of limiting force of friction

to the normal reaction. Table 6.4 gives the coefficient of static friction for some materials.

- 2] The limiting force of friction is independent of the apparent area between the surfaces in contact, so long as the normal reaction remains the same.
- 3] The limiting force of friction depends upon materials in contact and the nature of their surfaces.

Table 6.4: Coefficient of static friction

Material	Coefficient of static friction μ_s
Teflon on Teflon	0.4
Brass on steel	0.51
Copper on steel	0.53
Aluminium on steel	0.61
Steel on steel	0.74
Glass on glass	0.94
Rubber on concrete (dry)	1.0

Example 6.5: The coefficient of static friction between a block of mass 0.25 kg and a horizontal surface is 0.4. Find the horizontal force applied to it.

Solution: Given,

$$\mu_s = 0.4$$

$$m = 0.25 \text{ kg}$$

To find: Force

$$F = \mu_s \cdot N = \mu_s \cdot (mg)$$

$$F = 0.4 \times 0.25 \times 9.8$$

$$F = 0.98 \text{ N}$$

2. Kinetic friction :

Once the sliding of block on the surface starts, the force of friction decreases. The force required to keep the body sliding steadily is thus less than the force required to just start its sliding. The force of friction that comes into play when a body is in steady state of motion over another surface is called force of kinetic friction.

Friction between two surfaces in contact when one body is actually sliding over the other body, is called kinetic friction or dynamic friction.

Laws of kinetic friction :

1. The force of kinetic friction (F_k) is directly proportional to the normal reaction between two surfaces in contact.

$$\therefore F_k \propto N$$

$$\therefore F_k = \mu_k N \quad \text{--- (6.27)}$$

Where μ_k is constant of proportionality. It is called as coefficient of kinetic friction.

$$\therefore \mu_k = \frac{F_k}{N} \quad \text{--- (6.28)}$$

The coefficient of kinetic friction is defined as the ratio of force of kinetic friction to the normal reaction between the two surfaces in contact. Table 6.5 gives the co-efficient of kinetic friction for some materials.

2. Force of kinetic friction is independent of shape and apparent area of the surfaces in contact.
3. Force of kinetic friction depends upon the nature and material of the surfaces in contact.
4. The magnitude of the force of kinetic friction is independent of the relative velocity between the object and the surface provided that the relative velocity is neither too large nor too small.

Table 6.5: Coefficient of kinetic friction

Material	Coefficient of kinetic friction μ_k
Rubber on concrete (dry)	0.25
Glass on glass	0.40
Brass on steel	0.40
Copper on steel	0.44
Aluminium on steel	0.47
Steel on steel	0.57
Teflon on Teflon	0.80

3 Rolling friction :

Motion of a body over a surface is said to be rolling motion if the point of contact of the body with the surface keeps changing continuously.

Friction between two bodies in contact when one body is rolling over the other, is called rolling friction.

For same pair of surfaces, the force of static friction is greater than the force of kinetic

friction while the force of kinetic friction is greater than force of rolling friction. As rolling friction is the minimum, ball bearings are used to reduce friction in parts of machines to increase its efficiency.

Advantages of friction:

Friction is necessary in our daily life.

- We can walk due to friction between ground and feet.
- We can hold object in hand due to static friction.
- Brakes of vehicles work due to friction; hence we can reduce speed or stop vehicles.
- Climbing on a tree is possible due to friction.

Disadvantages of friction

- Friction opposes motion.
- Friction produces heat in different parts of machines. It also produces noise.
- Automobile engines consume more fuel due to friction.

Methods of reducing friction

- Use of lubricants, oil and grease in different parts of a machine.
- Use of ball bearings converts kinetic friction into rolling friction.



Can you tell?

- 1) It is difficult to run fast on sand.
- 2) It is easy to roll than pull a barrel along a road.
- 3) An inflated tyre rolls easily than a flat tyre.
- 4) Friction is a necessary evil.



Internet my friend

1. <https://opentextbc.ca/chapter/friction>.
2. <https://www.livescience.com>
3. <https://www.khanacademy.org/physics>
4. <https://courses.lumenlearning.com/elastiscitychapter/elasticity>
5. <https://www.tooper.com/guides/physics>



Exercises

1. Choose the correct answer:

- i) Change in dimensions is known as.....
(A) deformation (B) formation
(C) contraction (D) strain.
- ii) The point on stress-strain curve at which strain begins to increase even without increase in stress is called....
(A) elastic point (B) yield point
(C) breaking point (D) neck point
- iii) Strain energy of a stretched wire is 18×10^{-3} J and strain energy per unit volume of the same wire and same cross section is 6×10^{-3} J/m³. Its volume will be....
(A) 3cm³ (B) 3 m³
(C) 6 m³ (D) 6 cm³
- iv) ----- is the property of a material which enables it to resist plastic deformation.
(A) elasticity (B) plasticity
(C) hardness (D) ductility
- v) The ability of a material to resist fracturing when a force is applied to it, is called.....
(A) toughness (B) hardness
(C) elasticity (D) plasticity.

2. Answer in one sentence:

- i) Define elasticity.
- ii) What do you mean by deformation?
- iii) State the SI unit and dimensions of stress.
- iv) Define strain.
- v) What is Young's modulus of a rigid body?
- vi) Why bridges are unsafe after a very long use?
- vii) How should be a force applied on a body to produce shearing stress?
- viii) State the conditions under which Hooke's law holds good.
- ix) Define Poisson's ratio.
- x) What is an elastomer?
- xi) What do you mean by elastic hysteresis?
- xii) State the names of the hardest material

and the softest material.

- xiii) Define friction.
- xiv) Why force of static friction is known as 'self-adjusting force'?
- xv) Name two factors on which the coefficient of friction depends.

3. Answer in short:

- i) Distinguish between elasticity and plasticity.
- ii) State any four methods to reduce friction.
- iii) What is rolling friction? How does it arise?
- iv) Explain how lubricants help in reducing friction?
- v) State the laws of static friction.
- vi) State the laws of kinetic friction.
- vii) State advantages of friction.
- viii) State disadvantages of friction.
- ix) What do you mean by a brittle substance? Give any two examples.

4. Long answer type questions:

- i) Distinguish between Young's modulus, bulk modulus and modulus of rigidity.
- ii) Define stress and strain. What are their different types?
- iii) What is Young's modulus? Describe an experiment to find out Young's modulus of material in the form of a long straight wire.
- iv) Derive an expression for strain energy per unit volume of the material of a wire.
- v) What is friction? Define coefficient of static friction and coefficient of kinetic friction. Give the necessary formula for each.
- vi) State Hooke's law. Draw a labeled graph of tensile stress against tensile strain for a metal wire up to the breaking point. In this graph show the region in which Hooke's law is obeyed.

5. Answer the following

- i) Calculate the coefficient of static friction for an object of mass 50 kg placed on horizontal table pulled by attaching a spring balance. The force is increased gradually it is observed that the object just moves when spring balance shows 50N.
[Ans: $\mu_s = 0.102$]
- ii) A block of mass 37 kg rests on a rough horizontal plane having coefficient of static friction 0.3. Find out the least force required to just move the block horizontally.
[Ans: $F_s = 108.8\text{N}$]
- iii) A body of mass 37 kg rests on a rough horizontal surface. The minimum horizontal force required to just start the motion is 68.5 N. In order to keep the body moving with constant velocity, a force of 43 N is needed. What is the value of a) coefficient of static friction? and b) coefficient of kinetic friction?
[Ans: a) $\mu_s = 0.188$
b) $\mu_k = 0.118$]
- iv) A wire gets stretched by 4mm due to a certain load. If the same load is applied to a wire of same material with half the length and double the diameter of the first wire. What will be the change in its length?
[Ans: 0.5mm]
- v) Calculate the work done in stretching a steel wire of length 2m and cross sectional area 0.0225mm^2 when a load of 100 N is slowly applied to its free end. [Young's modulus of steel = $2 \times 10^{11} \text{ N/m}^2$]
[Ans: 2.222J]
- vi) A solid metal sphere of volume 0.31m^3 is dropped in an ocean where water pressure is $2 \times 10^7 \text{ N/m}^2$. Calculate change in volume of the sphere if bulk modulus of the metal is $6.1 \times 10^{10} \text{ N/m}^2$
[Ans: 10^{-4} m^3]
- vii) A wire of mild steel has initial length 1.5 m and diameter 0.60 mm is extended by 6.3 mm when a certain force is applied to it. If Young's modulus of mild steel is $2.1 \times 10^{11} \text{ N/m}^2$, calculate the force applied.
[Ans: 250 N]
- viii) A composite wire is prepared by joining a tungsten wire and steel wire end to end. Both the wires are of the same length and the same area of cross section. If this composite wire is suspended to a rigid support and a force is applied to its free end, it gets extended by 3.25mm. Calculate the increase in length of tungsten wire and steel wire separately.
[Given: $Y_{\text{steel}} = 2 \times 10^{11} \text{ N/m}^2$,
 $Y_{\text{Tungsten}} = 3.40 \times 10^8 \text{ N/m}^2$]
[Ans: extension in tungsten wire = 3.244 mm,
extension in steel wire = 0.0052 mm]
- ix) A steel wire having cross sectional area 1.2 mm^2 is stretched by a force of 120 N. If a lateral strain of 1.455 mm is produced in the wire, calculate the Poisson's ratio.
[Ans: 0.291]
- x) A telephone wire 125m long and 1mm in radius is stretched to a length 125.25m when a force of 800N is applied. What is the value of Young's modulus for material of wire?
[Ans: $1.27 \times 10^{11} \text{ N/m}^2$]
- xi) A rubber band originally 30cm long is stretched to a length of 32cm by certain load. What is the strain produced?
[Ans: 6.667×10^{-2}]
- xii) What is the stress in a wire which is 50m long and 0.01cm^2 in cross section, if the wire bears a load of 100kg?
[Ans: $9.8 \times 10^8 \text{ N/m}^2$]
- xiii) What is the strain in a cable of original length 50m whose length increases by 2.5cm when a load is lifted?
[Ans: 5×10^{-4}]



Can you recall?

1. Temperature of a body determines its hotness while heat energy is its heat content.
2. Pressure is the force exerted per unit area normally on the walls of a container by the gas molecules due to collisions.
3. Solids, liquids and gases expand on heating.
4. Substances change their state from solid to liquid or liquid to gas on heating up to specific temperature.

7.1 Introduction:

In previous lessons, while describing the equilibrium states of a mechanical system or while studying the motion of bodies, only three fundamental physical quantities namely length, mass and time were required. All other physical quantities in mechanics or related to mechanical properties can be expressed in terms of these three fundamental quantities. In this chapter, we will discuss properties or phenomena related to heat. These require a fourth fundamental quantity, the temperature, as mentioned in Chapter 1.

The sensation of hot or cold is a matter of daily experience. A mother feels the temperature of her child by touching its forehead. A cook throws few drops of water on a frying pan to know if it is hot enough to spread the dosa batter. Although not advisable, in our daily lives, we feel hotness or coldness of a body by touching or we dip our fingers in water to check if it is hot enough for taking bath. When we say a body or water is hot, we actually mean that its temperature is more than our hand. However, in this way, we can only compare the hotness or coldness of two objects **qualitatively**. Hot and cold are relative terms. You might recall the example given in your science textbook of VIIIth standard. Lukewarm water seems colder than hot water but hotter than cold water to our hands. We ascribe a property 'temperature' to an object to determine its degree of hotness. The higher the temperature, the hotter is the body. However, the precise temperature of a body can be known only when we have an accurate and easily reproducible way to

quantitatively measure it. Scientific precision requires measurement of a physical quantity in numerical terms. A thermometer is the device to measure the temperature.

In this chapter, we will learn properties of matter and various phenomena that are related to heat. Phenomena or properties having to do with temperature changes and heat exchanges are termed as thermal phenomena or thermal properties. You will understand why the direction of wind near a sea shore changes during day and night, why the metal lid of a glass bottle comes out easily on heating and why two metal vessels locked together can be separated by providing heat to the outer vessel.

7.2 Temperature and Heat:

Heat is energy in transit. When two bodies at different temperatures are brought in contact, they exchange heat. After some time, the heat transfer stops and we say the two bodies are in thermal equilibrium. The property or the deciding factor to determine the state of thermal equilibrium is the temperature of the two bodies. Temperature is a physical quantity that defines the thermodynamic state of a system.

You might have experienced that a glass of ice-cold water when left on a table eventually warms up whereas a cup of hot tea on the same table cools down. It means that when the temperature of a body, ice-cold water or hot tea in the above examples, is different from its surrounding medium, heat transfer takes place between the body and the surrounding medium until the body and the surrounding medium are at the same temperature. We then say that the body and its surroundings have reached

a state of thermal equilibrium and there is no net transfer of heat from one to the other. In fact, whenever two bodies are in contact, there is a transfer of heat owing to their temperature difference.

Matter in any state - solid, liquid or gas - consists of particles (ions, atoms or molecules). In solids, these particles are vibrating about their fixed equilibrium positions and possess kinetic energy due to motion at the given temperature. The particles possess potential energy due to the interatomic forces that hold the particles together at some mean fixed positions. Solids therefore have definite volume and shape. When we heat a solid, we provide energy to the solid. The particles then vibrate with higher energy and we can see that the temperature of the solid increases (except near its melting point). Thus the energy supplied to the solid (does not disappear!) becomes the internal energy in the form of increased kinetic energy of atoms/molecules and raises the temperature of the solid. The temperature is therefore a measure of the average kinetic energy of the atoms/molecules of the body. The greater the kinetic energy is, the faster the molecules will move and higher will be the temperature of the body. If we continue heating till the solid starts to melt, the heat supplied is used to weaken the bonds between the constituent particles. The average kinetic energy of the constituent particles does not change further. The order of magnitude of the average distance between the molecules of the melt remains almost the same as that of solid. Due to weakened bonds liquids do not possess definite shape but have definite volume. The mean distance between the particles and hence the density of liquid is more or less the same as that of the solid. On heating further, the atoms/molecules in liquid gain kinetic energy and temperature of the liquid increases. If we continue heating the liquid further, at the boiling point, the constituents can move freely overcoming the interatomic/molecular forces and the mean distance between the constituents increases so that the particles are farther apart.

As per kinetic theory of gases, for an ideal gas, there are no forces between the molecules

of a gas. Hence gases neither have a definite volume nor shape. Interatomic spacing in solids is $\sim 10^{-10}$ m while the average spacing in liquids is almost twice that in solids. The average inter molecular spacing in gases at normal temperature and pressure (NTP) is $\sim 10^{-9}$ m.

From the above discussion, we understand that **heat supplied to the substance increases the kinetic energy of molecules or atoms of the substance. The average kinetic energy per particle of a substance defines the temperature.** Temperature measures the degree of hotness of an object and not the amount of its thermal energy.

A glass of water, a gas enclosed in a container, a block of copper metal are all examples of a 'system'. We can say that heat in the form of energy is transferred between two (or more) systems or a system and its surroundings by virtue of their temperature difference. SI unit of heat energy is joule (J) and that of temperature is kelvin (K) or celcius ($^{\circ}\text{C}$). The CGS unit of heat energy is erg. ($1\text{J} = 10^7$ erg). The other unit of heat energy, that you have learnt in VIIIth standard, is calorie (cal) and the relation with J is $1\text{ cal} = 4.184\text{ J}$. Heat being energy has dimension $[\text{L}^2\text{M}^1\text{T}^{-2}\text{K}^0]$ while dimension of temperature is $[\text{L}^0\text{M}^0\text{T}^0\text{K}^1]$.

7.3 Measurement of Temperature:

In order to isolate two liquids or gases from each other and from the surroundings, we use containers and partitions made of materials like wood, plastic, glass wool, etc. An ideal wall or partition (not available in practice) separating two systems is one that does not allow any flow or exchange of heat energy from one system to the other. Such a perfect thermal insulator is called an **adiabatic wall** and is generally shown as a **thick cross-shaded** (slanting lines) region. When we wish to allow exchange of heat energy between two systems, we use a partition like a thin sheet of copper. It is termed as a **diathermic wall** and is represented as a thin dark region.

Let us consider two sections of a container separated by an adiabatic wall. Let them contain two different gases. Let us call them system A

and system B. We independently bring systems A and B in thermal equilibrium with a system C. Now if we remove the adiabatic wall separating systems A and B, there will be no transfer of heat from system A to system B or vice versa. This indicates that systems A and B are also in thermal equilibrium. Overall conclusion of this activity can be summarized as follows: If systems A and B are separately in thermal equilibrium with a system C, then A and B are also mutually in thermal equilibrium. When two or more systems/ bodies are in thermal equilibrium, their temperatures are same. This principle is used to measure the temperature of a system by using a thermometer.



Do you know ?

If $T_A = T_B$ and $T_B = T_C$, then $T_A = T_C$ is not a mathematical statement, if T_X represents the temperature of system X. It is the zeroth law of thermodynamics and makes the science of Thermometry possible.

Do you remember that to know the temperature of our body, doctor brings the mercury in the thermometer down to indicate some low temperature. We are then asked to keep the thermometer in our mouth. We have to wait for some time before the thermometer is taken out to know the temperature of our body. There is transfer of heat energy from our body to the thermometer since initially our body is at a higher temperature. When the temperature on the thermometer is same as that of our body, thermal equilibrium is said to be attained and heat transfer stops.

As mentioned above, to precisely know the thermodynamic state of any system, we need to know its temperature. The device used to measure temperature is a thermometer. Thermometry is the science of temperature and its measurement. For measurement of temperature, we need to establish a temperature scale and adopt a set of rules for assigning numbers (with corresponding units).

For the calibration of a thermometer, a standard temperature interval is selected between two easily reproducible fixed

temperatures just as we select the standard of length (metre) to be the distance between two fixed marks. The fact that substances change state from solid to liquid to gas at fixed temperatures is used to define reference temperature called fixed point. The two fixed temperatures selected for this purpose are the melting point of ice or freezing point of water and the boiling point of water. The next step is to sub-divide this standard temperature interval into sub-intervals by utilizing some physical property that changes with temperature and call each sub-interval a degree of temperature. This procedure sets up an empirical scale for temperature.

- * The temperature at which pure water freezes at one standard atmospheric pressure is called **ice point**/ freezing point of water. This is also the melting point of ice.
- * The temperature at which pure water boils and vaporizes into steam at one standard atmospheric pressure is called **steam point**/ boiling point. This is also the temperature at which steam changes to liquid water.

Having decided the fixed point phenomena, it remains to assign numerical values to these fixed points and the number of divisions between them. In 1750, conventions were adopted to assign (i) a temperature at which pure ice melts at one atmosphere pressure (the ice point) to be 0° and (ii) a temperature at which pure water boils at one atmosphere (the steam point) to be 100° so that there are 100 degrees between the fixed points. This was the centigrade scale (*centi* meaning hundred in Latin). This was redefined as celcius scale after the Swedish scientist Anders Celcius (1701-1744). It is a convention to express temperature as degree celcius ($^\circ\text{C}$).

To measure temperature quantitatively, generally two different scales of temperature are used. They are describe below.

- 1) **Celsius scale:-** On this scale, the ice point is marked as 0 and the steam point is marked as 100, both taken at normal atmospheric pressure (10^5 Pa or N/m^2). The interval between these points is divided into 100

equal parts. Each of these is known as degree Celsius and is written as °C.

- 2) **Fahrenheit scale** :- On this scale, the ice point is marked as 32 and the steam point is marked as 212, both taken at normal atmospheric pressure. The interval between these points is divided into 180 equal parts. Each division is known as degree Fahrenheit and is written as °F.

A relationship for conversion between the two scales may be obtained from a graph of fahrenheit temperature (T_F) versus celsius temperature (T_C). The graph is a straight line (Fig. 7.1) whose equation is

$$\frac{T_F - 32}{180} = \frac{T_C}{100} \quad \text{--- (7.1)}$$

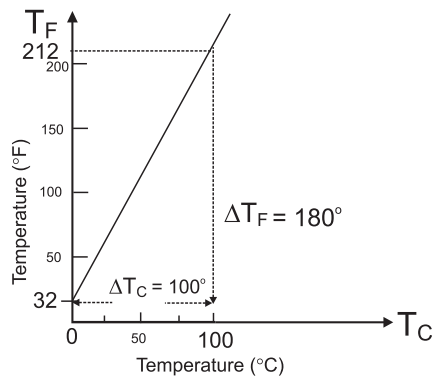


Fig. 7.1: A plot of fahrenheit temperature (T_F) versus celsius temperature (T_C).

Example 7.1: Average room temperature on a normal day is 27 °C. What is the room temperature in °F?

Solution: We have

$$\begin{aligned} \frac{T_F - 32}{180} &= \frac{T_C}{100} \\ \therefore T_F &= \frac{180}{100} T_C + 32 \\ \text{Given } T_C &= 27^\circ\text{C}, \\ T_F &= \frac{180}{100} \times 27 + 32 \\ &= 48.6 + 32 \\ &= 80.6^\circ\text{F} \end{aligned}$$

Example 7.2: Normal human body temperature in feherenheit is 98.4 °F. What is the body temperature in °C?

Solution: We have

$$\begin{aligned} \frac{T_C}{100} &= \frac{T_F - 32}{180} \\ \therefore T_C &= \frac{100}{180} (T_F - 32) \end{aligned}$$

Given $T_F = 98.4^\circ\text{F}$,

$$\begin{aligned} T_C &= \frac{100}{180} (98.4 - 32) \\ &= \frac{100}{180} (66.4) \\ &= 36.89^\circ\text{C} \end{aligned}$$

A device used to measure temperature, is based on the principle of thermal equilibrium. To measure the temperature, we use different measurable properties of materials which change with temperature. Some of them are length of a rod, volume of a liquid, electrical resistance of a metal wire, pressure of a gas at constant volume etc. Such changes in physical properties with temperature are used to design a thermometer. Physical property that is used in the thermometer for measuring the temperature is called the thermometric property and the material employed for the purpose is termed as the thermometric substance. Temperature is measured by exploiting the continuous monotonic variation of the chosen property with temperature. A calibration, however, is required to define the temperature scale.

There are different kinds of thermometers each type being more suitable than others for a certain job. In each type, the physical property used to measure the temperature must vary continuously over a wide range of temperature. It must be accurately measurable with simple apparatus.

An important characteristic of a thermometer is its *sensitivity*, i.e., a change in the thermometric property for a very small change in temperature. Two other characteristics are *accuracy* and *reproducibility*. Also it is important that the system attains thermal equilibrium with the thermometer quickly.

If the values of a thermometric property are P_1 and P_2 at the ice point (0 °C) and steam point (100 °C) respectively and the value of this property is P_T at unknown temperature T , then T is given by the following equation

$$T = \frac{100(P_T - P_1)}{(P_2 - P_1)} \quad \text{--- (7.2)}$$

Ideally, there should be no difference

in temperatures recorded on two different thermometers. This is seen for thermometers based on gases as thermometric substances. In a constant volume gas thermometer, the pressure of a fixed volume of gas (measured by the difference in height) is used as the thermometric property. It is an accurate but bulky instrument.

Liquid-in-glass thermometer depends on the change in volume of the liquid with temperature. The liquid in a glass bulb expands up a capillary tube when the bulb is heated. The liquid must be easily seen and must expand (or contract) rapidly and by a large amount for a small change in temperature over a wide range of temperature. Most commonly used liquids are mercury and alcohol as they remain in liquid state over a wide range. Mercury freezes at -39°C and boils at 357°C ; alcohol freezes at -115°C and boils at 78°C . Thermochromic liquids are ones which change colour with temperature but have a limited range around room temperatures. For example, titanium dioxide and zinc oxide are white at room temperature but when heated change to yellow.

Example 7.3: The length of a mercury column in a mercury-in-glass thermometer is 25 mm at the ice point and 180 mm at the steam point. What is the temperature when the length is 60 mm?

Solution: Here the thermometric property P is the length of the mercury column. Using Eq. (7.2), we get

$$T = \frac{100(60 - 25)}{(180 - 25)} = 22.58^{\circ}\text{C}$$

Resistance thermometer uses the change of electrical resistance of a metal wire with temperature. It measures temperature accurately in the range -2000°C to 1200°C but it is bulky and is best for steady temperatures.

Example 7.4: A resistance thermometer has resistance $95.2\ \Omega$ at the ice point and $138.6\ \Omega$ at the steam point. What resistance would be obtained if the actual temperature is 27°C ?

Solution: Here the thermometric property P is the resistance. Using Eq. (7.2), if R is the resistance at 27°C , we have

$$\begin{aligned} 27 &= \frac{100(R - 95.2)}{(138.6 - 95.2)}, \\ \therefore R &= \frac{27 \times (138.6 - 95.2)}{100} + 95.2 \\ &= 11.72 + 95.2 = 106.92\ \Omega \end{aligned}$$

Normally in research laboratories, a thermocouple is used to measure the temperature. A thermocouple is a junction of two different metals or alloys e.g., copper and iron joined together. When two such junctions at the two ends of two dissimilar metal rods are kept at two different temperatures, an electromotive force is generated that can be calibrated to measure the temperature.

Thermistor is another device used to measure temperature based on the change in resistance of a semiconductor materials i.e., the resistance is the thermometric property. You will learn more about this device in Chapter 14 on Semiconductors.

7.4 Absolute Temperature and Ideal Gas Equation:

7.4.1 Absolute zero and absolute temperature

Experiments carried out with gases at low densities indicate that while pressure is held constant, the volume of a given quantity of gas is directly proportional to temperature (measured in $^{\circ}\text{C}$). Similarly, if the volume of a given quantity of gas is held constant, the pressure of the gas is directly proportional to temperature (measured in $^{\circ}\text{C}$). These relations are graphically shown in Fig. 7.2 (a) and (b). Mathematically, this relationship can be written as $PV \propto T_C$. Thus the volume-temperature or pressure-temperature graphs for a gas are straight lines. They show that gases expand linearly with temperature on a mercury thermometer i.e., equal temperature increase causes equal volume or pressure increase. The similar thermal behavior of all gases suggests that this relationship of gases can be used to measure temperature in a constant-volume gas thermometer in terms of pressure of the gas.

Although actual experimental measurements might differ a little from the ideal linear relationship, the linear relationship

holds over a wide temperature range.

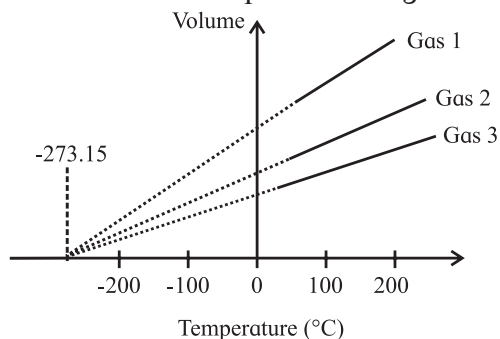


Fig. 7.2 (a): Graph of volume versus temperature (in °C) at constant pressure.

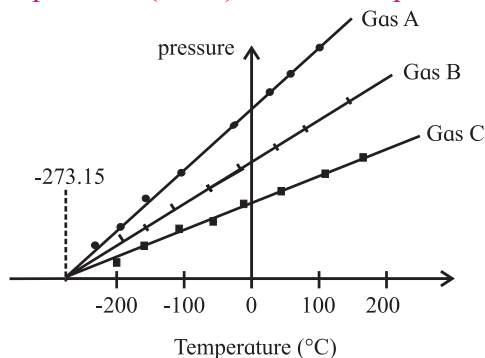


Fig. 7.2 (b): Graph of pressure versus temperature (in °C) at constant volume.

It may be noted that the lines do not pass through the origin i.e., have non-zero intercept along the y-axis. The straight lines have different slopes for different gases. If we assume that the gases do not liquefy even if we lower the temperature, we can extend the straight lines backwards for low temperatures. Is it possible to reach a temperature where the gases stop exerting any pressure i.e., pressure is zero? In a constant pressure thermometer, as the temperature is lowered, the volume decreases. Suppose the gas does not liquefy even at very low temperature, at what temperature, will its volume become zero? Practically it is not possible to keep a material in gaseous state for very low temperature and without exerting any pressure. If we extrapolate the graph of pressure P versus temperature T_C (in °C), the temperature at which the pressure of a gas would be zero is -273.15 °C. It is seen that all the lines for different gases cut the temperature axis at the same point at i.e., -273.15 °C. This point is termed as the **absolute zero of temperature**. It is not possible to attain a temperature lower than this value. Even to achieve absolute zero

temperature is not possible in practice. It may be noted that the point of zero pressure or zero volume does not depend on any specific gas.

The two fixed point scale, described in Section 7.3, had a practical shortcoming for calibrating the scale. It was difficult to precisely control the pressure and identify the fixed points, especially for the boiling point as the boiling temperature is very sensitive to changes in pressure. Hence, a one fixed point scale was adopted in 1954 to define a temperature scale. This scale is called the **absolute scale** or thermodynamic scale. It is named as the kelvin scale after Lord Kelvin (1824-1907).

It is possible for all the three phases - solid, liquid and gas/vapour of a material - to coexist in equilibrium. This is known as the *triple point*. To know the triple point one has to see that three phases coexist in equilibrium and no one phase is dominating. This occurs for each substance at a single unique combination of temperature and pressure. Thus if three phases of water - solid ice, liquid water and water vapour- coexist, the pressure and temperature are automatically fixed. This is termed as the triple point of water and is a single fixed point to define a temperature scale.

The absolute scale of temperature, is so termed since it is based on the properties of an ideal gas and does not depend on the property of any particular substance. The zero of this scale is ideally the lowest temperature possible although it has not been achieved in practice. It is termed as Kelvin scale with its zero at -273.15 °C and temperature intervals same as that on the celsius scale. It is written as K (without °). Internationally, triple point of water has been assigned as 273.16 K at pressure equal to 6.11×10^2 Pa or 6.11×10^{-3} atmosphere, as the standard fixed point for calibration of thermometers. Size of one kelvin is thus $1/273.16$ of the difference between the absolute zero and triple point of water. It is same as one celcius. On celcius scale, the triple point of water is 0.01 °C and not zero.

Three identical thermometers, marked in kelvin, celcius and fahrenheit, placed in a fixed temperature bath, each thermometer showing

the same rise in the level of mercury for human body temperature, are depicted in Fig. 7.3.

The relation between the three scales of temperature is as given in Eq. (7.3) .

$$\frac{T_C}{100} = \frac{T_F - 32}{180} = \frac{T_K - 273.15}{100} \quad \text{--- (7.3)}$$

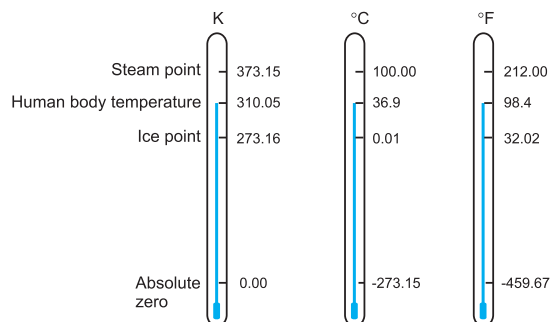


Fig. 7.3: Comparison of the kelvin, celsius and fahrenheit temperature scales (Thermometer reading are not to the scale).

Example 7.5: Express $T = 24.57$ K in celsius and fahrenheit.

Solution: We have

$$\frac{T_F - 32}{180} = \frac{T_C}{100} = \frac{T_K - 273.15}{100}$$

$$\begin{aligned} \therefore T_C &= T_K - 273.15 \\ &= 24.57 - 273.15 \\ &= -248.58^\circ\text{C} \end{aligned}$$

$$\frac{T_F - 32}{180} = \frac{T_K - 273.15}{100}$$

$$\begin{aligned} \therefore T_F &= \frac{180}{100} (T_K - 273.15) + 32 \\ &= \frac{9}{5} (24.57 - 273.15) + 32 \\ &= -447.44 + 32 \\ &= -415.44^\circ\text{F} \end{aligned}$$

Example 7.6: Calculate the temperature which has the same value on fahrenheit scale and kelvin scale.

Solution: Let the required temperature be y . i.e., $T_F = T_K = y$ then we have

$$\frac{y - 32}{180} = \frac{y - 273.15}{100}$$

$$\text{or, } 5y - 160 = 9y - 2458.35$$

$$\text{or, } 4y = -160 + 2458.35$$

$$\therefore y = 574.59$$

Thus 574.59°F and 574.59 K are equivalent temperatures.

7.4.2 Ideal Gas Equation:

The relation between three properties of a gas i.e., pressure, volume and temperature is called ideal gas equation. You will learn more about the properties of gases in chemistry.

Using absolute temperatures, the gas laws can be stated as given below.

- 1) **Charles' law-** In Fig. 7.2 (b), the volume-temperature graph passes through the origin if temperatures are measured on the kelvin scale, that is if we take 0 K as the origin. In that case the volume V is directly proportional to the absolute temperature T .

Thus $V \propto T$

$$\text{or, } \frac{V}{T} = \text{constant} \quad \text{--- (7.4)}$$

Thus Charles' law can be stated as, the volume of a fixed mass of gas is directly proportional to its absolute temperature if the pressure is kept constant.

- 2) **Pressure (Gay Lussac's) law-** From Fig.7.2, it can be seen that the pressure-temperature graph is similar to the volume-temperature graph.

Thus $P \propto T$

$$\text{or, } \frac{P}{T} = \text{constant} \quad \text{--- (7.5)}$$

Pressure law can be stated as the pressure of a fixed mass of gas is directly proportional to its absolute temperature if the volume is kept constant.

- 3) **Boyle's law-** For fixed mass of gas at constant temperature, pressure is inversely proportional to volume.

$$\text{Thus } P \propto \frac{1}{V}$$

$$PV = \text{constant} \quad \text{--- (7.6)}$$

Combining above three equations, we get

$$\frac{PV}{T} = \text{constant} \quad \text{--- (7.7)}$$

For one mole of a gas, the constant of proportionality is written as R

$$\therefore \frac{PV}{T} = R \quad \text{or} \quad PV = RT \quad \text{--- (7.8)}$$

If given mass of a gas consists of n moles, then Eq. (7.8) can be written as

$$PV = nRT \quad \text{--- (7.9)}$$

This relation is called ideal gas equation. The value of constant R is same for all gases. Therefore, it is known as universal gas constant. Its numerical value is $8.31 \text{ J K}^{-1} \text{ mol}^{-1}$.

Example 7.7: The pressure reading in a thermometer at steam point is $1.367 \times 10^3 \text{ Pa}$. What is pressure reading at triple point knowing the linear relationship between temperature and pressure?

Solution: We have $P_{\text{triple}} = 273.16 \times \left(\frac{P}{T}\right)$ where P_{triple} and P are the pressures at temperature of triple point (273.16 K) and T respectively. We are given that $P = 1.367 \times 10^3 \text{ Pa}$ at steam point i.e., at $273.15 + 100 = 373.15 \text{ K}$.

$$\therefore P_{\text{triple}} = 273.16 \times \left(\frac{1.367 \times 10^3}{373.15}\right) \\ = 1.000 \times 10^3 \text{ Pa}$$

7.5 Thermal Expansion:

When matter is heated, it normally expands and when cooled, it normally contracts. The atoms in a solid vibrate about their mean positions. When heated, they vibrate faster and force each other to move a little farther apart. This results into expansion. The molecules in a liquid or gas move with certain speed. When heated, they move faster and force each other to move a little farther apart. This results in expansion of liquids and gases on heating. The expansion is more in liquids than in solids; gases expand even more.

A change in the temperature of a body causes change in its dimensions. The increase in the dimensions of a body due to an increase in its temperature is called thermal expansion. There are three types of thermal expansion: 1) Linear expansion, 2) Areal expansion, 3) Volume expansion.

7.5.1 Linear Expansion:

The expansion in length due to thermal energy is called linear expansion.

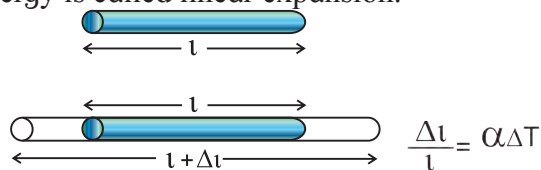


Fig. 7.4: Linear expansion Δl is exaggerated for explanation.

If the substance is in the form of a long rod of length l , then for small change ΔT , in temperature, the fractional change $\Delta l/l$, in length (shown in Fig.7.4), is directly proportional to ΔT .

$$\frac{\Delta l}{l} \propto \Delta T \\ \text{or } \frac{\Delta l}{l} = \alpha \Delta T \quad \text{--- (7.10)}$$

where α is called the coefficient of linear expansion of solid. Its value depends upon nature of the material. Rearranging Eq. (7.10), we get

$$\alpha = \frac{\Delta l}{l \Delta T} \\ = \frac{l_T - l_0}{l_0 (T - T_0)} \quad \text{--- (7.11)}$$

where l_0 = length of rod at 0°C

l_T = length of rod when heated to $T^\circ \text{C}$

$T_0 = 0^\circ \text{C}$ is initial temperature

T = final temperature

$\Delta l = l_T - l_0$ = change in length

$\Delta T = T - T_0$ = rise in temperature

Referring to Eq. (7.11), if $l_0 = 1$ and $T - T_0 = 1^\circ \text{C}$, then

$$\alpha = l_T - l_0 \text{ (numerically).}$$

Coefficient of linear expansion of a solid is thus defined as increase in the length per unit original length at 0°C for one degree centigrade rise in temperature.

The unit of coefficient of linear expansion is per degree celcius or per kelvin. The magnitude of α is very small and it varies only a little with temperature. For most practical purposes, α can be assumed to be constant for a particular material. Therefore, it is not necessary that initial temperature be taken as 0°C . Equation (7.11) can be rewritten as

$$\alpha = \frac{l_2 - l_1}{l_1 (T_2 - T_1)} \quad \text{--- (7.12)}$$

where l_1 = initial length at temperature $T_1^\circ \text{C}$

l_2 = final length at temperature $T_2^\circ \text{C}$.

Table 7.1 lists average values of coefficient of linear expansion for some materials in the

temperature range 0°C to 100 °C.

Table 7.1: Values of coefficient of linear expansion for some common materials.

Materials	α (K ⁻¹)
Carbon (diamond)	0.1×10^{-5}
Glass	0.85×10^{-5}
Iron	1.2×10^{-5}
Steel	1.3×10^{-5}
Gold	1.4×10^{-5}
Copper	1.7×10^{-5}
Silver	1.9×10^{-5}
Aluminium	2.5×10^{-5}
Sulphur	6.1×10^{-5}
Mercury	6.1×10^{-5}
Water	6.9×10^{-5}
Carbon (graphite)	8.8×10^{-5}

Example 7.8: The length of a metal rod at 27 °C is 4 cm. The length increases to 4.02 cm when the metal rod is heated upto 387 °C. Determine the coefficient of linear expansion of the metal rod.

Solution: Given

$$T_1 = 27 \text{ }^\circ\text{C}$$

$$T_2 = 387 \text{ }^\circ\text{C}$$

$$l_1 = 4 \text{ cm} = 4 \times 10^{-2} \text{ m}$$

$$l_2 = 4.02 \text{ cm} = 4.02 \times 10^{-2} \text{ m}$$

We have

$$\begin{aligned} \alpha &= \frac{l_2 - l_1}{l_1(T_2 - T_1)} \\ &= \frac{(4.02 - 4.0) \times 10^{-2}}{4 \times 10^{-2}(387 - 27)} \\ &= \frac{0.02 \times 10^{-2}}{4 \times 10^{-2} \times 360} \\ &= 1.39 \times 10^{-5} / ^\circ\text{C} \end{aligned}$$

Example 7.9: Length of an iron rod at temperature 27°C is 4.256 m. Find the temperature at which the length of the same rod increases to 4.268 m. (α for iron = $1.2 \times 10^{-5} \text{ K}^{-1}$)

Solution: Given

$$T_1 = 27^\circ\text{C}, l_1 = 4.256 \text{ m},$$

$$l_2 = 4.268 \text{ m}, \alpha = 1.2 \times 10^{-5} \text{ K}^{-1}$$

We have

$$\begin{aligned} \alpha &= \frac{l_2 - l_1}{l_1(T_2 - T_1)} = \frac{l_2 - l_1}{l_1 T_2 - l_1 T_1} \\ \therefore l_1 T_2 &= l_1 T_1 + \frac{l_2 - l_1}{\alpha} \\ T_2 &= \frac{1}{l_1} \left[l_1 T_1 + \frac{l_2 - l_1}{\alpha} \right] \\ &= \frac{1}{4.256} \left[(4.256 \times 27) + \frac{4.268 - 4.256}{1.2 \times 10^{-5}} \right] \\ &= \frac{1}{4.256} \left[114.912 + \frac{0.012}{1.2 \times 10^{-5}} \right] \\ &= \frac{1}{4.256} [114.912 + 1000] \\ &= 261.96 \text{ }^\circ\text{C} \end{aligned}$$

7.5.2 Areal Expansion:

The increase ΔA , in the surface area, on heating is called areal expansion or superficial expansion.

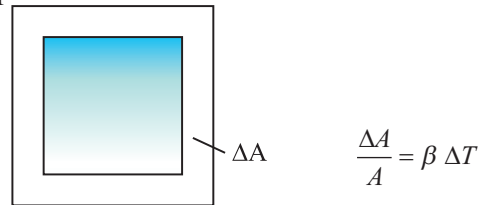


Fig. 7.5: Areal expansion ΔA is exaggerated for explanation.

If a substance is in the form of a plate of area A , then for small change ΔT in temperature, the fractional change in area, $\Delta A/A$ (as shown in Fig. 7.5), is directly proportional to ΔT .

$$\begin{aligned} \frac{\Delta A}{A} &\propto \Delta T \\ \text{or } \frac{\Delta A}{A} &= \beta \Delta T \end{aligned} \quad \text{--- (7.13)}$$

where β is called the coefficient of areal expansion of solid. It depends on the material of the solid. Rearranging Eq. (7.13), we get

$$\beta = \frac{\Delta A}{A \Delta T} = \frac{A_T - A_0}{A_0(T - T_0)} \quad \text{--- (7.14)}$$

where A_0 = area of plate at 0 °C

A_T = area of plate when heated to T °C

T_0 = 0 °C is initial temperature

T = final temperature

$\Delta A = A_T - A_0$ = change in area

$\Delta T = T - T_0$ = rise in temperature.

If $A_0 = 1 \text{ m}^2$ and $T - T_0 = 1 \text{ }^\circ\text{C}$, then

$\beta = A_T - A_0$ (numerically).

Therefore, **coefficient of areal expansion of a solid is defined as the increase in the area per unit original area at 0°C for one degree rise in temperature.**

The unit of β is per degree celcius or per kelvin.

As in the case of α , β also does not vary much with temperature. Hence, if A_1 is the area of a metal plate at T_1 °C and A_2 is the area at higher temperature T_2 °C, then

$$\beta = \frac{A_2 - A_1}{A_1(T_2 - T_1)} \quad \text{--- (7.15)}$$

Example 7.10: A thin aluminium plate has an area 286 cm² at 20 °C. Find its area when it is heated to 180 °C.

(β for aluminium = 4.9×10^{-5} /°C)

Solution: Given

$$T_1 = 20 \text{ °C}$$

$$T_2 = 180 \text{ °C}$$

$$A_1 = 286 \text{ cm}^2$$

$$\beta = 4.9 \times 10^{-5} \text{ /°C}$$

We have

$$\beta = \frac{A_2 - A_1}{A_1(T_2 - T_1)}$$

$$\begin{aligned} \therefore A_2 &= A_1 [1 + \beta (T_2 - T_1)] \\ &= 286 [1 + 4.9 \times 10^{-5} (180 - 20)] \\ &= 286 [1 + 4.9 \times 10^{-5} \times 160] \\ &= 286 [1 + 784.0 \times 10^{-5}] \\ &= 286 [1 + 0.00784] \\ &= 286 [1.00784] \\ \therefore A_2 &= 288.24 \text{ cm}^2 \end{aligned}$$

7.5.3 Volume expansion

The increase in volume due to heating is called volume expansion or cubical expansion.

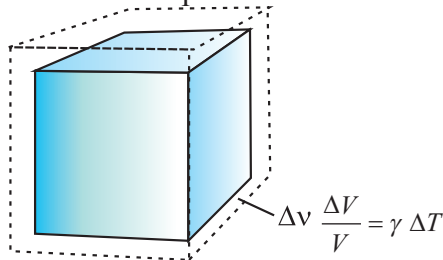


Fig. 7.6: Volume expansion ΔV is exaggerated for explanation.

If the substance is in the form of a cube of volume V , then for small change ΔT in temperature, the fractional change, $\Delta V/V$ (as shown in Fig.7.6), in volume is directly proportional to ΔT .

$$\frac{\Delta V}{V} \propto \Delta T$$

$$\text{or } \frac{\Delta V}{V} = \gamma \Delta T \quad \text{--- (7.16)}$$

where γ is called coefficient of cubical or volume expansion. It depends upon the nature of the material. Its unit is per degree celcius or per kelvin. From Eq.(7.16), we can write

$$\gamma = \frac{\Delta V}{V \Delta T} = \frac{V_T - V_0}{V_0(T - T_0)} \quad \text{--- (7.17)}$$

where V_0 = volume at 0 °C

V_T = volume when heated to T °C

$T_0 = 0$ °C is initial temperature

T = final temperature

$\Delta V = V_T - V_0$ = change in volume

$\Delta T = T - T_0$ = rise in temperature.

If $V_0 = 1 \text{ m}^3$, $T - T_0 = 1$ °C, then

$\gamma = V_T - V_0$ (numerically).

The coefficient of cubical expansion of a solid is therefore defined as increase in volume per unit original volume at 0°C for one degree rise in the temperature.

If V_1 is the volume of a body at T_1 °C and V_2 is the volume at higher temperature T_2 °C, then

$$\gamma_1 = \frac{V_2 - V_1}{V_1(T_2 - T_1)} \quad \text{--- (7.18)}$$

γ_1 is the coefficient of volume expansion at temperature T_1 °C.

Since fluids possess definite volume and take the shape of the container, only change in volume is significant. Equations (7.17) and (7.18) are valid for cubical or volume expansion of fluids. It is to be noted that since fluids are kept in containers, when one deals with the volume expansion of fluids, expansion of the container is also to be considered. If expansion of fluid results in a volume greater than the volume of the container, the fluid will overflow if the container is open. If the container is closed, volume expansion of fluid will cause additional

pressure on the walls of the container. Can you now tell why the balloon bursts sometimes on its own on a hot day?

Normally solids and liquids expand on heating. Hence their volume increases on heating. Since the mass is constant, it results in a decrease in the density on heating. You have learnt about the anomalous behaviour of water. Water expands on cooling from 4°C to 0°C. Hence its density decreases on cooling in this temperature range.

In Table 7.2 are given typical average values of the coefficient of volume expansion γ for some materials in the temperature range 0°C to 100°C.

Table 7.2: Values of coefficient of volume expansion for some common materials.

Materials	γ (K ⁻¹)
Invar	2×10^{-6}
Glass (ordinary)	2.5×10^{-5}
Steel	$(3.3-3.9) \times 10^{-5}$
Iron	3.55×10^{-5}
Gold	4.2×10^{-5}
Brass	5.7×10^{-5}
Aluminium	6.9×10^{-5}
Mercury	18.2×10^{-5}
Water	20.7×10^{-5}
Paraffin	58.8×10^{-5}
Gasoline	95.0×10^{-5}
Alcohol (ethyl)	110×10^{-5}

γ is also characteristic of the substance but is not strictly a constant. It depends in general on temperature as shown in Fig.7.7. It is seen that γ becomes constant only at very high temperatures.

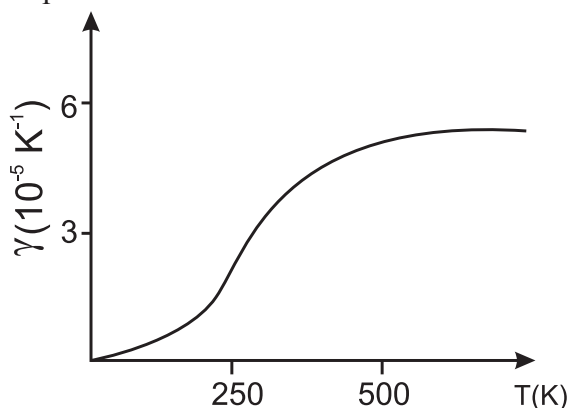


Fig. 7.7: Coefficient of volume expansion of copper as a function of temperature.

Example 7.11 : A liquid at 0 °C is poured in a glass beaker of volume 600 cm³ to fill it completely. The beaker is then heated to 90 °C. How much liquid will overflow?

$$(\gamma_{\text{liquid}} = 1.75 \times 10^{-4} / ^\circ\text{C}, \gamma_{\text{glass}} = 2.75 \times 10^{-5} / ^\circ\text{C})$$

Solution: Given

$$V_1 = 600 \text{ cm}^3$$

$$T_1 = 0 ^\circ\text{C}$$

$$T_2 = 90 ^\circ\text{C}$$

$$\text{We have } \gamma = \frac{V_2 - V_1}{V_1(T_2 - T_1)}$$

$$\therefore \text{increase in volume} = V_2 - V_1 = \gamma V_1(T_2 - T_1)$$

Increase in volume of beaker

$$\begin{aligned} &= \gamma_{\text{glass}} \times V_1(T_2 - T_1) \\ &= 2.75 \times 10^{-5} \times 600 \times (90 - 0) \\ &= 2.75 \times 10^{-5} \times 600 \times 90 \\ &= 148500 \times 10^{-5} \text{ cm}^3 \end{aligned}$$

$$\therefore \text{increase in volume of beaker} = 1.485 \text{ cm}^3$$

Increase in volume of liquid

$$\begin{aligned} &= \gamma_{\text{liquid}} \times V_1(T_2 - T_1) \\ &= 1.75 \times 10^{-4} \times 600 \times (90 - 0) \\ &= 1.75 \times 10^{-4} \times 600 \times 90 \\ &= 94500 \times 10^{-4} \text{ cm}^3 \end{aligned}$$

$$\therefore \text{increase in volume of liquid} = 9.45 \text{ cm}^3$$

$\therefore$ volume of liquid which overflows

$$\begin{aligned} &= (9.45 - 1.485) \text{ cm}^3 \\ &= 7.965 \text{ cm}^3 \end{aligned}$$

7.5.4 Relation between Coefficients of Expansion:

i) Relation between β and α :

Consider a square plate of side l_0 at 0 °C and l_T at T °C.

$$\therefore l_T = l_0(1 + \alpha T) \text{ from Eq. (7.11).}$$

If area of plate at 0 °C is A_0 , $A_0 = l_0^2$.

If area of plate at T °C is A_T ,

$$A_T = l_T^2 = l_0^2(1 + \alpha T)^2$$

$$\text{or } A_T = A_0(1 + \alpha T)^2 \quad \text{--- (7.19)}$$

Also from Eq. (7.14),

$$A_T = A_0(1 + \beta T) \quad \text{--- (7.20)}$$

Using Eqs. (7.19) and (7.20), we get
 $A_0 (1+\alpha T)^2 = A_0 (1+\beta T)$
 or $1+2\alpha T + \alpha^2 T^2 = 1+\beta T$
 Since the values of α are very small, the term $\alpha^2 T^2$ is very small and may be neglected.
 $\therefore \beta = 2\alpha$ --- (7.21)



Can you tell?

1. Why the metal wires for electrical transmission lines sag?
2. Why a railway track is not a continuous piece but is made up of segments separated by gaps?
3. How a steel wheel is mounted on an axle to fit exactly?
4. Why lakes freeze first at the surface?

The result is general because any solid can be regarded as a collection of small squares.

ii) Relation between γ and α :

Consider a cube of side l_0 at 0°C and l_T at $T^\circ\text{C}$.

$$\therefore l_T = l_0 (1+\alpha T) \text{ from Eq. (7.11).}$$

If volume of the cube at 0°C is V_0 , $V_0 = l_0^3$.

If volume of the cube at $T^\circ\text{C}$ is V_T ,

$$V_T = l_T^3 = l_0^3 (1+\alpha T)^3$$

$$\text{or } V_T = V_0 (1+\alpha T)^3 \quad \text{--- (7.22)}$$

Also from Eq. (7.17),

$$V_T = V_0 (1+\gamma T) \quad \text{--- (7.23)}$$

Using Eqs. (7.22) and (7.23), we get

$$V_0 (1+\alpha T)^3 = V_0 (1+\gamma T)$$

$$\text{or } 1+3\alpha T+3\alpha^2 T^2+\alpha^3 T^3=1+\gamma T$$

Since the values of α are very small, the terms with higher powers of α may be neglected.

$$\therefore \gamma = 3\alpha \quad \text{--- (7.24)}$$

Again the result is general because any solid can be regarded as a collection of small cubes.

Finally, the relation between α , β and γ is

$$\alpha = \frac{\beta}{2} = \frac{\gamma}{3} \quad \text{--- (7.25)}$$

Example 7.12: A sheet of brass is 50 cm long and 8 cm broad at 0°C . If the surface area at 100°C is 401.57cm^2 , find the coefficient of linear expansion of brass.

Solution: Given

$$T_1 = 0^\circ\text{C}$$

$$T_2 = 100^\circ\text{C}$$

$$A_1 = 50 \times 8 = 400\text{ cm}^2$$

$$A_2 = 401.57\text{ cm}^2$$

We have

$$\begin{aligned} \beta = 2\alpha &= \frac{A_2 - A_1}{A_1(T_2 - T_1)} \\ &= \frac{(401.57 - 400)\text{ cm}^2}{400\text{ cm}^2 \times (100 - 0)^\circ\text{C}} \\ &= \frac{1.57}{400 \times 100} = 0.3925 \times 10^{-4}^\circ\text{C}^{-1} \\ \therefore \alpha &= 0.1962 \times 10^{-4}^\circ\text{C}^{-1} \\ &= 1.962 \times 10^{-5}^\circ\text{C}^{-1} \end{aligned}$$

$\therefore$ Coefficient of linear expansion of brass is $1.962 \times 10^{-5} / ^\circ\text{C}$.



Do you know ?

- * When pressure is held constant, due to change in temperature, the volume of a liquid or solid changes very little in comparison to the volume of a gas.
- * The coefficient of volume expansion, γ , is generally an order of magnitude larger for liquids than for solids.
- * Metals have high values for the coefficient for linear expansion, α , than non-metals.
- * γ changes more with temperature than α and β .
- * We know that water expands on freezing from 4°C to 0°C . Other two substances, that expand on freezing are metals bismuth (Bi) and antimony (Sb). Thus the density of liquid is more than corresponding solid and hence solid Bi or Sb float on their liquids like ice floats on water.

7.6 Specific Heat Capacity:

7.6.1 Specific Heat Capacity of Solids and Liquids

If 1 kg of water and 1 kg of paraffin are heated in turn for the same time by the same heater, the temperature rise of paraffin is about twice that of water. Since the heater gives equal amounts of heat energy to each liquid, it seems that different substances require different

amounts of heat to cause the same temperature rise of 1°C in the same mass of 1 kg.

If ΔQ stands for the amount of heat absorbed or given out by a substance of mass m when it undergoes a temperature change ΔT , then the specific heat capacity of that substance is given by

$$s = \frac{\Delta Q}{m\Delta T} \quad \text{--- (7.26)}$$

If $m = 1 \text{ kg}$ and $\Delta T = 1^{\circ}\text{C}$ then $s = \Delta Q$.

Thus specific heat capacity is defined as the amount of heat per unit mass absorbed or given out by the substance to change its temperature by one unit (one degree) 1°C or 1K .

Table 7.3: Specific heat capacity of some substances at room temperature and atmospheric pressure.

Substance	Specific heat capacity ($\text{J kg}^{-1} \text{K}^{-1}$)
Steel	120
Lead	128
Gold	129
Tungsten	134.4
Silver	234
Copper	387
Iron	448
Carbon	506.5
Glass	837
Aluminium	903.0
Kerosene	2118
Paraffin oil	2130
Alcohol (ethyl)	2400
Ethanol	2500
Water	4186.0

The SI unit of specific heat capacity is $\text{J/kg } ^{\circ}\text{C}$ or J/kg K and C.G.S. unit is $\text{erg/g } ^{\circ}\text{C}$ or erg/g K . The specific heat capacity is a property of the substance and weakly depends on its temperature. Except for very low temperatures, the specific heat capacity is almost constant for all practical purposes.

If the amount of substance is specified in terms of moles μ instead of mass m in kg, then the specific heat is called molar specific heat (C) and is given by

$$C = \frac{1}{\mu} \frac{\Delta Q}{\Delta T} \quad \text{--- (7.27)}$$

The SI unit of molar specific heat capacity is $\text{J/mol } ^{\circ}\text{C}$ or J/mol K . Like specific heat, molar specific heat also depends on the nature of the substance and its temperature. Table 7.3 lists the values of specific heat capacity for some common materials.

From Table 7.3, it can be seen that water has the highest specific heat capacity compared to other substances. For this reason, water is used as a coolant in automobile radiators as well as for fomentation using hot water bags.

7.6.2 Specific Heat Capacity of Gas:

In case of a gas, slight change in temperature is accompanied with considerable changes in both, the volume and the pressure. If gas is heated at constant pressure, volume changes and therefore some work is done on the surroundings during expansion requiring additional heat. As a result, specific heat at constant pressure (S_p) is greater than specific heat at constant volume (S_v). It is thus necessary to define two principal specific heat capacities for a gas.

Principal specific heat capacities of gases:

- The principal specific heat capacity of a gas at constant volume (S_v) is defined as the quantity of heat absorbed or released for the rise or fall of temperature of unit mass of a gas through 1 K (or 1°C) when its volume is kept constant.
- The principal specific heat capacity of a gas at constant pressure (S_p) is defined as the quantity of heat absorbed or released for the rise or fall of temperature of unit mass of a gas through 1 K (1°C) when its pressure is kept constant.

Molar specific heat capacities of gases:

- Molar specific heat capacity of a gas at constant volume (C_v) is defined as the quantity of heat absorbed or released for the rise or fall of temperature of one mole of the gas through 1 K (or 1°C), when its volume is kept constant.

- b) Molar specific heat capacity of a gas at constant pressure (C_p) is defined as the quantity of heat absorbed or released for the rise or fall of temperature of one mole of the gas through 1K (or 1°C), when its pressure is kept constant.

Relation between Principal and Molar Specific Heat Capacities:

A relation between principal specific heat capacity and molar specific heat capacity is given by the following expression

Molar specific heat capacity = Molecular weight $\times$ principal specific heat capacity.

$$\text{i.e. } C_p = M \times S_p \text{ and } C_v = M \times S_v$$

where M is the molecular weight of the gas.

Table 7.4 lists values of molar specific heat capacity for some commonly known gases.

Table 7.4: Molar specific heat capacity of some gases.

Gas	C_p (J mol ⁻¹ K ⁻¹)	C_v (J mol ⁻¹ K ⁻¹)
He	20.8	12.5
H ₂	28.8	20.4
N ₂	29.1	20.8
O ₂	29.4	21.1
CO ₂	37.0	28.5

7.6.3 Heat Equation:

If a substance has a specific heat capacity of 1000 J/kg °C, it means that heat energy of 1000 J raises the temperature of 1 kg of that substance by 1°C or 6000 J will raise the temperature of 2 kg of the substance by 3 °C. If the temperature of 2 kg mass of the substance falls by 3 °C, the heat given out would also be 6000 J. In general we can write the heat equation as

Heat received or given out (Q) = mass (m) $\times$ temperature change (Δt) $\times$ specific heat capacity (s).

$$\text{or } Q = m \times \Delta T \times s \quad \text{--- (7.28)}$$

Example 7.13: If the temperature of 4 kg mass of a material of specific heat capacity 300 J/kg °C rises from 20 °C to 30 °C. Find the heat received.

Solution:

$$Q = 4 \text{ kg} \times (30-20) \text{ °C} \times 300 \text{ J/kg °C}$$

$$= 4 \times 10 \times 300 \text{ J}$$

$$\therefore Q = 12000 \text{ J}$$

7.6.4 Heat Capacity (Thermal Capacity):

Heat capacity or thermal capacity of a body is the quantity of heat needed to raise or lower the temperature of the whole body by 1°C (or 1K).

$\therefore$ Thermal heat capacity can be written as
Heat received or given out

$$= \text{mass} \times 1 \times \text{specific heat capacity}$$

$$\text{Heat capacity} = Q = m \times s \quad \text{--- (7.29)}$$

Heat capacity (thermal capacity) is measured in J/°C.

Example 7.14: Find thermal capacity for a copper block of mass 0.2 kg, if specific heat capacity of copper is 290 J/kg °C.

Solution: Given

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

$$s = 290 \text{ J/kg °C}$$

$$\text{Thermal capacity} = m \times s = 0.2 \text{ kg} \times 290 \text{ J/kg °C} = 58 \text{ J/°C}$$

7.7 Calorimetry:

Calorimetry is an experimental technique for the quantitative measurement of heat exchange. To make such measurement a calorimeter is used. Figure 7.8 shows a simple water calorimeter.

It consists of cylindrical vessel made of copper or aluminium and provided with a stirrer and a lid. The calorimeter is well-insulated to prohibit any transfer of heat into or out of the calorimeter.

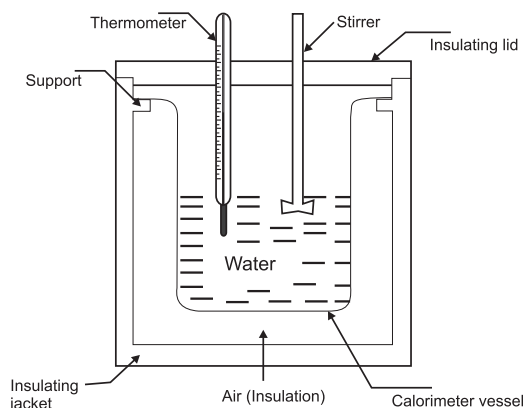


Fig. 7.8: Calorimeter.

One important use of calorimeter is to determine the specific heat of a substance using

the principle of conservation of energy. Here we are dealing with heat energy and the system is isolated from surroundings. Therefore, heat gained is equal to the heat lost.

In the technique known as the “method of mixtures”, a sample 'A' of the substance is heated to a high temperature which is accurately measured. The sample 'A' is then placed quickly in the calorimeter containing water. The contents are stirred constantly until the mixture attains a final common temperature. The heat lost by the sample 'A' will be gained by the water and the calorimeter. The specific heat of the sample 'A' of the substance can be calculated as under:

Let

m_1 = mass of the sample 'A'

m_2 = mass of the calorimeter and the stirrer

m_3 = mass of the water in calorimeter

s_1 = specific heat capacity of the substance of sample 'A'

s_2 = specific heat capacity of the material of calorimeter (and stirrer)

s_3 = specific heat capacity of water

T_1 = initial temperature of the sample 'A'

T_2 = initial temperature of the calorimeter stirrer and water

T = final temperature of the combined system

We have the data as follows:

Heat lost by the sample 'A' = $m_1 s_1 (T_1 - T)$

Heat gained by the calorimeter and the stirrer
= $m_2 s_2 (T - T_2)$

Heat gained by the water = $m_3 s_3 (T - T_2)$

Assuming no loss of heat to the surroundings, the heat lost by the sample goes into the calorimeter, stirrer and water. Thus writing heat equation as,

$$m_1 s_1 (T_1 - T) = m_2 s_2 (T - T_2) + m_3 s_3 (T - T_2) \quad \text{---(7.30)}$$

Knowing the specific heat capacity of water ($s_3 = 4186 \text{ J kg}^{-1} \text{ K}^{-1}$) and copper ($s_2 = 387 \text{ J kg}^{-1} \text{ K}^{-1}$) being the material of the calorimeter and the stirrer, one can calculate specific heat capacity (s_1) of material of sample 'A', from Eq. (7.30) as

$$s_1 = \frac{(m_2 s_2 + m_3 s_3)(T - T_2)}{m_1 (T_1 - T)} \quad \text{--- (7.31)}$$

Also, one can find specific heat capacity of water or any liquid using the following expression, if the specific heat capacity of the material of calorimeter and sample is known

$$s_3 = \frac{m_1 s_1 (T_1 - T)}{m_3 (T - T_2)} - \frac{m_2 s_2}{m_3} \quad \text{--- (7.32)}$$

Note - In the experiment, the heat from the solid sample 'A' is given to the liquid and therefore the sample should be denser than the liquid, so that sample does not float on the liquid.

Example 7.15: A sphere of aluminium of 0.06 kg is placed for sufficient time in a vessel containing boiling water so that the sphere is at 100°C . It is then immediately transferred to 0.12 kg copper calorimeter containing 0.30 kg of water at 25°C . The temperature of water rises and attains a steady state at 28°C . Calculate the specific heat capacity of aluminium. (Specific heat capacity of water, $s_w = 4.18 \times 10^3 \text{ J kg}^{-1} \text{ K}^{-1}$, specific heat capacity of copper $s_{Cu} = 0.387 \times 10^3 \text{ J kg}^{-1} \text{ K}^{-1}$)

Solution : Given

Mass of aluminium sphere = $m_1 = 0.06 \text{ kg}$

Mass of copper calorimeter = $m_2 = 0.12 \text{ kg}$

Mass of water in calorimeter

= $m_3 = 0.30 \text{ kg}$

Specific heat capacity of copper

= $s_{Cu} = s_2 = 0.387 \times 10^3 \text{ J kg}^{-1} \text{ K}^{-1}$

Specific heat capacity of water

= $s_w = s_3 = 4.18 \times 10^3 \text{ J kg}^{-1} \text{ K}^{-1}$

Initial temperature of aluminium sphere

= $T_1 = 100^\circ\text{C}$

Initial temperature of calorimeter and water = $T_2 = 25^\circ\text{C}$

Final temperature of the mixture

= $T = 28^\circ\text{C}$

We have

$$\begin{aligned}
 s_1 &= \frac{(m_2 s_2 + m_3 s_3)(T - T_2)}{m_1 (T_1 - T)} \\
 &= \frac{[(0.12 \times 387) + (0.30 \times 4180)](28 - 25)}{(0.06)(100 - 28)} \\
 &= \frac{(46.44 + 1254) \times 3}{(0.06) \times 72} = \frac{3901.32}{4.32} \\
 &= 903.08 \text{ J kg}^{-1} \text{ K}^{-1}
 \end{aligned}$$

∴ Specific heat capacity of aluminium is $903.08 \text{ J kg}^{-1} \text{ K}^{-1}$.

7.8 Change of State:

Matter normally exists in three states: solid, liquid and gas. A transition from one of these states to another is called a change of state. Two common changes of states are solid to liquid and liquid to gas (and vice versa). These changes can occur when exchange of heat takes place between the substance and its surroundings.



Activity

To understand the process of change of state

Take some cubes of ice in a beaker. Note the temperature of ice (0°C). Start heating it slowly on a constant heat source. Note the temperature after every minute. Continuously stir the mixture of water and ice. Observe the change in temperature. Continue heating even after the whole of ice gets converted into water. Observe the change in temperature as before till vapours start coming out. Plot the graph of temperature (along y-axis) versus time (along x-axis). You will obtain a graph of temperature versus time as shown in Fig. 7.9.

Analysis of observations :

1) From point A to B:

There is no change in temperature from point A to point B, this means the temperature of the ice bath does not change even though heat is being continuously supplied. That is the temperature remains constant until the entire amount of the ice melts. The heat supplied is being utilised in changing the state from solid

(ice) to liquid (water).

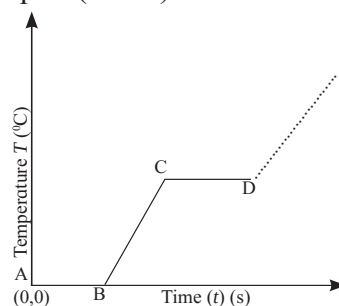


Fig. 7.9 : Variation of temperature with time.

- The change of state from solid to liquid is called melting and from liquid to solid is called solidification.
- Both the solid and liquid states of the substance co-exist in thermal equilibrium during the change of states from solid to liquid or vice versa.
- The temperature at which the solid and the liquid states of the substance are in thermal equilibrium with each other is called the melting point of solid (here ice) or freezing point of liquid (here water). It is characteristic of the substance and also depends on pressure.
- The melting point of a substance at one standard atmospheric pressure is called its normal melting point.
- At one standard atmospheric pressure, the freezing point of water and melting point of ice is 0°C or 32°F . The freezing point describes the liquid to solid transition while melting point describes solid-to-liquid transition.

2) From point B to D:

The temperature begins to rise from point B to point C, i.e., after the whole of ice gets converted into water and we continue further heating. We see that temperature begins to rise. The temperature keeps on rising till it reaches point C i.e., nearly 100°C . Then it again becomes steady. It is observed that the temperature remains constant until the entire amount of the liquid is converted into vapour. The heat supplied is now being utilized to change water from liquid state to vapour or gaseous state.

- The change of state from liquid to vapour

is called vapourisation while that from vapour to liquid is called condensation.

- b) Both the liquid and vapour states of the substance coexist in thermal equilibrium during the change of state from liquid to vapour.
- c) The temperature at which the liquid and the vapour states of the substance coexist is called the boiling point of liquid, here water or steam point. This is also the temperature at which water vapour condenses to form water.
- d) The boiling point of a substance at one standard atmospheric pressure is called its normal boiling point.



Can you tell?

1. What after point D in graph ? Can steam be hotter than 100°C ?
2. Why steam at 100°C causes more harm to our skin than water at 100°C ?



Do you know ?

You must have seen that water spilled on floor dries up after some time. Where does the water disappear? It is converted into water vapour and mixes with air. We say that water has evaporated. You also know that water can be converted into water vapour if you heat the water till its boiling point. **What is then the difference between boiling and evaporation?**

Both evaporation and boiling involve change of state, evaporation can occur at any temperature but boiling takes place at a fixed temperature for a given pressure, unique for each liquid. Evaporation takes place from the surface of liquid while boiling occurs in the whole liquid.

As you know, molecules in a liquid are moving about randomly. The average kinetic energy of the molecules decides the temperature of the liquid. However, all molecules do not move with the same speed. One with higher kinetic energy may escape from the surface region by overcoming the interatomic forces. This process can take place at any temperature. This is evaporation. If the temperature of the liquid is higher, more is the average kinetic energy. Since the number of molecules is fixed, it implies that the number of fast moving molecules

is more. Hence the rate of losing such molecules to atmosphere will be larger. Thus, higher is the temperature of the liquid, greater is the rate of evaporation. Since faster molecules are lost, the average kinetic energy of the liquid is reduced and hence the temperature of the liquid is lowered. Hence the phenomenon of evaporation gives a cooling effect to the remaining liquid. Since evaporation takes place from the surface of a liquid, the rate of evaporation is more if the area exposed is more and if the temperature of the liquid is higher.

You might have seen that if your mother wants her sari/clothes to dry faster, she does not fold them. More is the area exposed, faster is the drying because the water gets evaporated faster. The presence of wind or strong breeze and content of water vapour in the atmosphere are two other important factors determining the drying of clothes but we do not refer to them here.

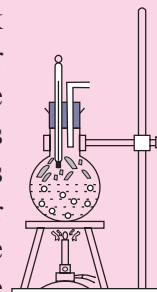
Before giving an injection to a patient, normally a spirit swab is used to disinfect the region. We feel a cooling effect on our skin due to evaporation of the spirit as explained before.



Activity

Activity to understand the dependence of boiling point on pressure

Take a round bottom flask, more than half filled with water. Keep it over a burner and fix a thermometer and steam outlet through the cork of the flask as shown in figure. As water in the flask gets heated, note that first the air, which was dissolved in the water comes out as small bubbles. Later bubbles of steam form at the bottom but as they rise to the cooler water near the top, they condense and disappear. Finally, as the temperature of the entire mass of the water reaches 100°C , bubbles of steam reach the surface and boiling is said to occur. The steam in the flask may not be visible but as it comes out of the flask, it condenses as tiny droplets of water giving a foggy appearance.



If now the steam outlet is closed for a few seconds to increase the pressure in the

flask, you will notice that boiling stops. More heat would be required to raise the temperature (depending on the increase in pressure) before boiling starts again. Thus boiling point increases with increase in pressure.

Let us now remove the burner. Allow water to cool to about 80°C. Remove the thermometers and steam outlet. Close the flask with a air tight cork. Keep the flask turned upside down on a stand. Pour ice-cold water on the flask. Water vapours in the flask condense reducing the pressure on the water surface inside the flask. Water begins to boil again, now at a lower temperature. Thus boiling point decreases with decrease in pressure and increases with increase in pressure.



Can you tell?

1. Why cooking is difficult at high altitude?
2. Why cooking is faster in pressure cooker?

7.8.1 Sublimation:

Have you seen what happens when camphor is burnt? All substances do not pass through the three states: solid-liquid-gas. There are certain substances which normally pass from the solid to the vapour state directly and vice versa. The change from solid state to vapour state without passing through the liquid state is called sublimation and the substance is said to sublime. Dry ice (solid CO_2) and iodine sublime. During the sublimation process, both the solid and vapour states of a substance coexist in thermal equilibrium. Most substances sublime at very low pressures.

7.8.2 Phase Diagram:

A pressure - temperature (PT) diagram often called a phase diagram, is particularly convenient for comparing different phases of a substance.

A phase is a homogeneous composition of a material. A substance can exist in different phases in solid state, e.g., you are familiar with two phases of carbon- graphite and diamond. Both are solids but the regular geometric

arrangement of carbon atoms is different in the two cases. Figure 7.10 shows the phase diagram of water and CO_2 . Let us try to understand the diagram.

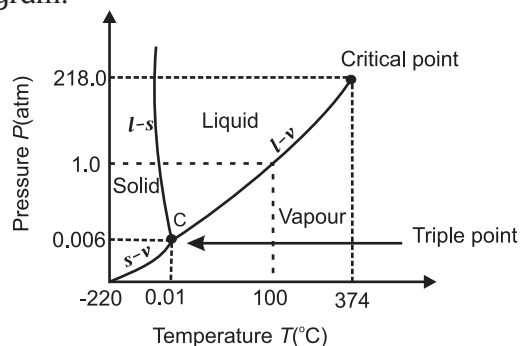


Fig. 7.10 (a): Phase diagram of water (not to scale).

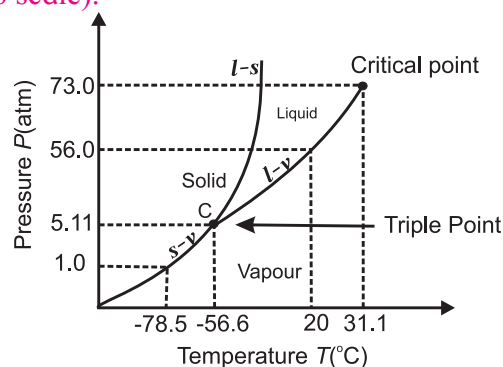


Fig. 7.10 (b): Phase diagram of CO_2 (not to scale).

i) Vapourisation curve $l-v$: The curve labelled $l-v$ represents those points where the liquid and vapour phases are in equilibrium. Thus it is a graph of boiling point versus pressure. Note that the curves correctly show that at a pressure of 1 atmosphere, the boiling points of water is 100°C and that the boiling point is lowered for a decreased pressure.

ii) Fusion curve $l-s$: The curve $l-s$ represents the points where the solid and liquid phases coexist in equilibrium. Thus it is a graph of the freezing point versus pressure. At one standard atmosphere pressure, the freezing point of water is 0 °C as shown in Fig. 7.10 (a). Also notice that at a pressure of one standard atmosphere water is in the liquid phase if the temperature is between 0 °C and 100 °C but is in the solid or vapour phase if the temperature is below 0 °C or above 100 °C. Note that $l-s$ curve for water slopes upward to the left i.e., fusion curve of water has a slightly negative slope. This is true

only of substances that expand upon freezing. However, for most materials like CO_2 , the l - s curve slopes upwards to the right i.e., fusion curve has a positive slope. The melting point of CO_2 is -56°C at higher pressure of 5.11 atm.

iii) Sublimation curve s - v : The curve labelled s - v is the sublimation point versus pressure curve. Water sublimates at pressure less than 0.0060 atmosphere, while carbon dioxide, which in the solid state is called dry ice, sublimates even at atmospheric pressure at temperature as low as -78°C .

iv) Triple point: The temperature and pressure at which the fusion curve, the vapourisation curve and the sublimation curve meet and all the three phases of a substance coexist is called the triple point of the substance. That is, the triple point of water is that point where water in solid, liquid and gaseous states coexist in equilibrium and this occurs only at a unique temperature and pressure. The triple point of water is 273.16 K and $6.11 \times 10^{-3}\text{ Pa}$ and that of CO_2 is -56.6°C and $5.1 \times 10^{-5}\text{ Pa}$.

7.8.3 Gas and Vapour:

The terms gas and vapour are sometimes used quite randomly. Therefore, it is important to understand the difference between the two. A gas cannot be liquefied by pressure alone, no matter how high the pressure is. In order to liquefy a gas, it must be cooled to a certain temperature. This temperature is called critical temperature.

Critical temperatures for some common gases and water vapour are given in Table 7.5. Thus, nitrogen must be cooled below -147°C to liquefy it by pressure.

Table 7.5: Critical Temperatures of some common gases and water vapour.

Gas	Critical Temperature	
	($^\circ\text{C}$)	(K)
Air	-190	83
N_2	-147	126
O_2	-118	155
CO_2	31.1	241.9
Water vapour	374	647

Gas and vapour can thus be defined as-

- 1) A substance which is in the gaseous phase and is above its critical temperature is called a gas.
- 2) A substance which is in the gaseous phase and is below its critical temperature is called a vapour.

Vapour can be liquefied simply by increasing the pressure, while gas cannot. Vapour also exerts pressure like a gas.

7.8.4 Latent Heat:

Whenever there is a change in the state of a substance, heat is either absorbed or given out but there is no change in the temperature of the substance.

Latent heat of a substance is the quantity of heat required to change the state of unit mass of the substance without changing its temperature.

Thus if mass m of a substance undergoes a change from one state to the other then the quantity of heat absorbed or released is given by $Q = mL$ --- (7.33)

where L is known as latent heat and is characteristic of the substance. Its SI unit is J kg^{-1} . The value of L depends on the pressure and is usually quoted at one standard atmospheric pressure.

The quantity of heat required to convert unit mass of a substance from its solid state to the liquid state, at its melting point, without any change in its temperature is called its latent heat of fusion (L_f).

The quantity of heat required to convert unit mass of a substance from its liquid state to vapour state, at its boiling point without any change in its temperature is called its latent heat of vapourization (L_v).

A plot of temperature versus heat energy for a given quantity of water is shown in Fig. 7.11.

From Fig. 7.11, we see that when heat is added (or removed) during a change of state, the temperature remains constant. Also the slopes of the phase lines are not all the same, which indicates that specific heats of the various states are not equal. For water the latent heat

of fusion and vaporisation are $L_f = 3.33 \times 10^5 \text{ J kg}^{-1}$ and $L_v = 22.6 \times 10^5 \text{ J kg}^{-1}$ respectively. That is $3.33 \times 10^5 \text{ J}$ of heat is needed to melt 1 kg of ice at 0°C and $22.6 \times 10^5 \text{ J}$ of heat is needed to convert 1 kg of water to steam at 100°C . Hence, steam at 100°C carries $22.6 \times 10^5 \text{ J kg}^{-1}$ more heat than water at 100°C . This is why burns from steam are usually more serious than those from boiling water. Melting points, boiling point and latent heats for various substances are given in Table 7.6.

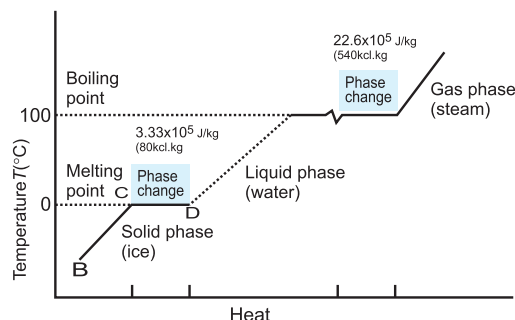


Fig. 7.11: Temperature versus heat for water at one standard atmospheric pressure (not to scale).

Table 7.6 : Temperature of change of state and latent heats for various substances at one standard atmosphere pressure.

Substance	Melting point ($^\circ \text{C}$)	L_f ($\times 10^5 \text{ J kg}^{-1}$)	Boiling point ($^\circ \text{C}$)	L_v ($\times 10^5 \text{ J kg}^{-1}$)
Gold	1063	0.645	2660	15.8
Lead	328	0.25	1744	8.67
Water	0	3.33	100	22.6
Ethyl alcohol	-114	1.0	78	8.5
Mercury	-39	0.12	357	2.7
Nitrogen	-210	0.26	-196	2.0
Oxygen	-219	0.14	-183	2.1

Example 7.16: When 0.1 kg of ice at 0°C is mixed with 0.32 kg of water at 35°C in a container. The resulting temperature of the mixture is 7.8°C . Calculate the heat of fusion of ice ($s_{\text{water}} = 4186 \text{ J kg}^{-1} \text{ K}^{-1}$).

Solution: Given

$$m_{\text{ice}} = 0.1 \text{ kg}$$

$$m_{\text{water}} = 0.32 \text{ kg}$$

$$T_{\text{ice}} = 0^\circ \text{C}$$

$$T_{\text{water}} = 35^\circ \text{C}$$

$$T_F = 7.8^\circ \text{C}$$

$$s_{\text{water}} = 4186 \text{ J kg}^{-1} \text{ K}^{-1}$$

Heat lost by water

$$= m_{\text{water}} s_{\text{water}} (T_F - T_{\text{water}})$$

$$= 0.32 \text{ kg} \times 4186 \text{ J} \times (7.8 - 35)^\circ \text{C}$$

= - 36434.944 J (here negative sign indicates loss of heat energy)

$$\text{Heat required to melt ice} = m_{\text{ice}} L_f = 0.1 \times L_f$$

Heat required to raise temperature of water (from ice) to final temperature

$$= m_{\text{ice}} s (T - T_{\text{ice}})$$

$$= 0.1 \text{ kg} \times 4186 \text{ J} \times (7.8 - 0)^\circ \text{C}$$

$$= 3265.08 \text{ J}$$

Head lost = Heat gained

$$36434.944 = 0.1 L_f + 3265.08$$

$$L_f = \frac{36434.944 - 3265.08}{0.1} = 3316.9864$$

$$= 3.31698 \times 10^5 \text{ J kg}^{-1}$$



Do you know ?

The latent heat of vapourization is much larger than the latent heat of fusion. The energy required to completely separate the molecules or atoms is greater than the energy needed to break the rigidity (rigid bonds between the molecules or atoms) in solids. Also when the liquid is converted into vapour, it expands. Work has to be done against the surrounding atmosphere to allow this expansion.

7.9 Heat Transfer:

Heat may be transferred from one point of body to another in three different ways- by conduction, convection and radiation. Heat transfers through solids by conduction. In this process, heat is passed on from one molecule to other molecule but the molecules do not leave their mean positions. Liquids and gases are heated by convection. In this process, there is a bodily movement of the heated molecules. In order to transfer heat by conduction and convection a material medium is required. However transfer of heat by radiation does not need any medium. Radiation of heat energy takes place by electromagnetic (EM) waves that travel with a speed of $3 \times 10^8 \text{ ms}^{-1}$ in the space/vacuum. The energy from the Sun comes to us by radiation. It may be noted that conduction is a slow process of heat transfer while convection is a rapid process. However radiation is the fastest process because the transfer of heat takes place at the speed of light.

7.9.1 Conduction:

Conduction is the process by which heat flows from the hot end to the cold end of a solid body without any net bodily movement of the particles of the body.

Heat passes through solids by conduction only. When one end of a metal rod is placed in a flame while the other end is held in hand, the end held in hand slowly gets hotter, although it itself is not in direct contact with flame. We say that heat has been conducted from the hot end to the cold end. When one end of the rod is heated, the molecules there vibrate faster. As they collide with their slow moving neighbours, they transfer some of their energy by collision to these molecules which in turn transfer energy to their neighbouring molecules still farther down the length of the rod. Thus the energy of thermal motion is transferred by molecular collisions down the rod. The transfer of heat continues till the two ends of the rod are at the same temperature in principle but this will take infinite time. Normally various sections of the rod will attain a temperature which remains constant but not same through out the length of

the rod. This method of heat transfer is called conduction.

Those solid substances which conduct heat easily are called good conductors of heat e.g. silver, copper, aluminium, brass etc. All metals are good conductors of heat. Those substances which do not conduct heat easily are called bad conductors of heat e.g. wood, cloth, air, paper, etc. In general, good conductors of heat are also good conductors of electricity. Similarly bad conductors of heat are bad conductors of electricity also.

7.9.1.1 Thermal Conductivity:

Thermal conductivity of a solid is a measure of the ability of the solid to conduct heat through it. Thus good conductors of heat have higher thermal conductivity than bad conductors.

Suppose that one end of a metal rod is heated (see Fig 7.12 (a)). The heat flows by conduction from hot end to the cold end. As a result the temperature of every section of the rod starts increasing. Under this condition, the rod is said to be in a variable temperature state. After some time the temperature at each section of the rod becomes steady i.e. does not change. Note that temperature of each cross-section of the rod is constant but not the same. This is called steady state condition. Under steady state condition, the temperature at points within the rod decreases uniformly with distance from the hot end to the cold end. The fall of temperature with distance between the ends of the rod in the direction of flow of heat, is called temperature gradient.

$$\therefore \text{Temperature gradient} = \frac{T_1 - T_2}{x}$$

where T_1 = temperature of hot end

T_2 = temperature of cold end

x = length of the rod

7.9.1.2 Coefficient of Thermal Conductivity:

Consider a cube of each side x and each face of cross-sectional area A . Suppose its opposite faces are maintained at temperatures T_1 and T_2 ($T_1 > T_2$) as shown in Fig. 7.12 (b). Experiments show that under steady state condition, the

quantity of heat 'Q' that flows from the hot face to the cold face is

- directly proportional to the cross-sectional area A of the face. i.e., $Q \propto A$
- directly proportional to the temperature difference between the two faces i.e., $Q \propto (T_1 - T_2)$
- directly proportional to time t (in seconds) for which heat flows i.e. $Q \propto t$
- inversely proportional to the perpendicular distance x between hot and cold faces i.e., $Q \propto 1/x$

Combining the above four factors, we have the quantity of heat

$$Q \propto \frac{A(T_1 - T_2)t}{x}$$

$$\therefore Q = \frac{kA(T_1 - T_2)t}{x} \quad \text{--- (7.34)}$$

where k is a constant of proportionality and is called coefficient of thermal conductivity. Its value depends upon the nature of the material.

If $A = 1 \text{ m}^2$, $T_1 - T_2 = 1^\circ\text{C}$ (or 1 K), $t = 1 \text{ s}$ and $x = 1 \text{ m}$, then from Eq. (7.34), $Q = k$.

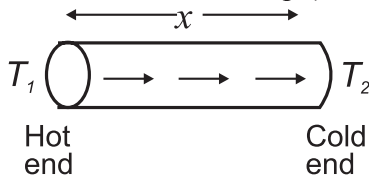


Fig 7.12 (a): Section of a metal bar in the steady state.

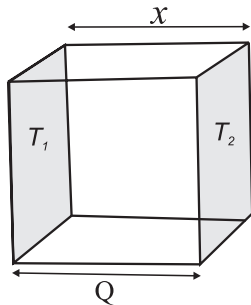


Fig 7.12 (b): Section of a cube in the steady state.

Thus the **coefficient of thermal conductivity of a material is defined as the quantity of heat that flows in one second between the opposite faces of a cube of side 1 m, the faces being kept at a temperature difference of 1°C (or 1 K).**

From Eq. (7.34), we have

$$k = \frac{Qx}{A(T_1 - T_2)t} \quad \text{--- (7.35)}$$

SI unit of coefficient of thermal conductivity k is $\text{J s}^{-1} \text{ m}^{-1} ^\circ\text{C}^{-1}$ or $\text{J s}^{-1} \text{ m}^{-1} \text{ K}^{-1}$ and its dimensions are $[\text{L}^1 \text{ M}^1 \text{ T}^3 \text{ K}^{-1}]$.

From Eq. (7.34), we also have

$$\frac{Q}{t} = \frac{kA(T_1 - T_2)}{x} \quad \text{--- (7.36)}$$

The quantity Q/t , denoted by P_{cond} , is the time rate of heat flow (i.e. heat flow per second) from the hotter face to the colder face, at right angles to the faces. Its SI unit is watt (W). SI unit of k can therefore be written as $\text{W m}^{-1} ^\circ\text{C}^{-1}$ or $\text{W m}^{-1} \text{ K}^{-1}$.

Using calculus, Eq. (7.36) may be written as

$$\frac{dQ}{dt} = -kA \frac{dT}{dx}$$

where $\frac{dT}{dx}$ is the temperature gradient.

The negative sign indicates that heat flow is in the direction of decreasing temperature.

If $A = 1 \text{ m}^2$ and $\frac{dT}{dx} = 1$, then $\frac{dQ}{dt} = k$ (numerically).

Hence the **coefficient of thermal conductivity of a material may also be defined as the rate of flow of heat per unit area per unit temperature gradient when the heat flow is at right angles to the faces of a thin parallel-sided slab of material.**

The coefficients of thermal conductivity of some materials are given in Table 7.7.

7.9.1.3 Thermal Resistance (R_T):

Conduction rate P_{cond} is the amount of energy transferred per unit time through a slab of area A and thickness x , the two sides of the slab being at temperatures T_1 and T_2 ($T_1 > T_2$), and is given by Eq. (7.36)

$$P_{\text{cond}} = \frac{Q}{t} = kA \frac{T_1 - T_2}{x} \quad \text{--- (7.37)}$$

As discussed earlier, k depends on the material of the slab. A material that readily transfers heat energy by conduction is a good thermal conductor and has high value of k .

Table 7.7: Coefficient of thermal conductivity (k).

Substance	Coefficient of thermal conductivity ($\text{J s}^{-1} \text{m}^{-1} \text{K}^{-1}$)
Silver	406
Copper	385
Aluminium	205
Steel	50.2
Insulating brick	0.15
Glass	0.8
Brick and concrete	0.8
Water	0.8
Wood	0.04-0.12
Air at 0°C	0.024

In western countries, where the temperature falls below 0°C in winter season, insulating the house from the surroundings is very important. In our country, if we wish to carry cold drinks with us for picnic or wish to bring ice-cream from the shop to our house, we need to keep them in containers (made up of say thermocol) that are poor thermal conductors. Hence the concept of thermal resistance R_T , similar to electrical resistance, is introduced. The opposition of a body, to the flow of heat through it, is called thermal resistance. The greater the thermal conductivity of a material, the smaller is its thermal resistance and vice versa. Thus bad thermal conductors are those which have high thermal resistance.

From Eq. (7.37)

$$\frac{(T_1 - T_2)}{P_{\text{cond}}} = \frac{x}{kA} \quad \text{--- (7.38)}$$

We know that when a current flows through a conductor, the ratio V/I is called the electrical resistance of the conductor where V is the electrical potential difference between the ends of the conductor and I is the current or rate of flow of charge. In Eq. (7.38), $(T_1 - T_2)$ is the temperature difference between the ends of the conductor and P_{cond} is the rate of flow of heat. Therefore in analogy with electrical resistance, $(T_1 - T_2)/P_{\text{cond}}$ is called thermal resistance R_T of the material i.e.,

$$\text{Thermal resistance } R_T = \frac{x}{kA}$$

The SI unit of thermal resistance is $^\circ\text{C s/kcal}$

or $^\circ\text{C s/J}$ and its dimensional formula is $[\text{M}^{-1} \text{L}^{-2} \text{T}^3 \text{K}^1]$.

The lower the thermal conductivity k , the higher is the thermal resistance. R_T A material with high R_T value is a poor thermal conductor and is a good thermal insulator. Thermal resistivity ρ_T is the reciprocal of thermal conductivity k and is characteristic of a material while thermal resistance is that of slab (or of rod) and depends on the material and on the thickness of slab (or length of rod).

Example 7.17: What is the rate of energy loss in watt per square metre through a glass window 5 mm thick if outside temperature is -20°C and inside temperature is 25°C ? ($k_{\text{glass}} = 1 \text{ W/m K}$)

Solution : Given

$$k_{\text{glass}} = 1 \text{ W/m K}$$

$$T_1 = 25^\circ\text{C}$$

$$T_2 = -20^\circ\text{C}$$

$$x = 5 \text{ mm} = 5 \times 10^{-3} \text{ m}$$

$$\therefore T_1 - T_2 = 25 - (-20)^\circ\text{C} = 45 \text{ K}$$

$$\text{We have } P_{\text{cond}} = \frac{Q}{t} = kA \frac{T_1 - T_2}{x}.$$

$\therefore$ The rate of energy loss per square metre is

$$\frac{P_{\text{cond}}}{A} = k \frac{T_1 - T_2}{x}$$

$$= 1 \text{ W m}^{-1} \text{K}^{-1} \times 45 \text{ K} / (5 \times 10^{-3} \text{ m})$$

$$= 9 \times 10^3 \text{ W/m}^2$$

7.9.1.4 Applications of Thermal Conductivity:

- Cooking utensils are made of metals but are provided with handles of bad conductors.

Since metals are good conductors of heat, heat can be easily conducted through the base of the utensils. The handles of utensils are made of bad conductors of heat (e.g., wood, ebonite etc.) so that they can not conduct heat from the utensils to our hands.

- Thick walls are used in the construction of cold storage rooms. Brick is a bad conductor of heat so that it reduces the flow of heat from the surroundings to the rooms. Still better heat insulation is obtained by using hollow bricks. Air

being a poorer conductor than a brick, it further avoids the conduction of heat from outside.

- iii) To prevent ice from melting it is wrapped in a gunny bag. A gunny bag is a poor conductor of heat and reduces the heat flow from outside to ice. Moreover, the air filled in the interspaces of a gunny bag, being very bad conductor of heat, further avoids the conduction of heat from outside.

Low thermal conductivity can also be a disadvantage. When hot water is poured in a glass beaker the inner surface of the glass expands on heating. Since glass is a bad conductor of heat, the heat from inside does not reach the outside surface so quickly. Hence the outer surface does not expand thereby causing a crack in the glass.

Example 7.18: The temperature difference between two sides of an iron plate, 2 cm thick, is 10 °C. Heat is transmitted through the plate at the rate of 600 kcal per minute per square metre at steady state. Find the thermal conductivity of iron.

Solution: Given

$$\frac{Q}{At} = 600 \text{ kcal / min m}^2 = \frac{600}{60} = 10 \text{ kcal / s m}^2$$

$$x = 2 \text{ cm} = 2 \times 10^{-2} \text{ m}$$

$$T_1 - T_2 = 10^\circ\text{C}$$

From Eq.(7.34), we have

$$Q = \frac{kA(T_1 - T_2)t}{x}$$

$$\therefore k = \frac{Q}{At} \frac{x}{T_1 - T_2}$$

$$= \frac{10 \text{ kcal / s m}^2 \times 2 \times 10^{-2} \text{ m}}{10^\circ\text{C}}$$

$$= 0.02 \text{ kcal / m s }^\circ\text{C}$$

Example 7.19: Calculate the rate of loss of heat through a glass window of area 1000 cm² and thickness of 4 mm, when temperature inside is 27 °C and outside is -5 °C. Coefficient of thermal conductivity of glass is 0.022 cal/ s cm °C.

Solution : Given

$$A = 1000 \text{ cm}^2 = 1000 \times 10^{-4} \text{ m}^2$$

$$k = 0.022 \text{ cal/ s cm }^\circ\text{C} = 0.022 \times 10^2 \text{ cal/m }^\circ\text{C}$$

$$x = 4 \text{ mm} = 0.4 \times 10^{-2} \text{ m}$$

$$T_1 = 27^\circ\text{C}, T_2 = -5^\circ\text{C}$$

From Eq. (7.34), we have

$$Q = \frac{kA(T_1 - T_2)t}{x}$$

$$\therefore \frac{Q}{t} = \frac{kA(T_1 - T_2)}{x}$$

$$= \frac{0.022 \times 10^2 \times 1000 \times 10^{-4} \times (27 - (-5))}{0.4 \times 10^{-2}}$$

$$= 1.76 \times 10^3 \text{ cal / s} = 1.76 \text{ kcal / s}$$

7.9.2 Convection:

We have seen that heat is transmitted through solids by conduction wherein energy is transferred from one molecule to another but the molecules themselves vibrating with larger amplitude do not leave their mean positions. But in convection, heat is transmitted from one point to another by the actual bodily movement of the heated (energised) molecules within the fluid.

In liquids and gases heat is transmitted by convection because their molecules are quite free to move about. The mechanism of heat transfer by convection in liquids and gases is described below.

Consider water being heated in a vessel from below. The water at the bottom of the vessel is heated first and consequently its density decreases i.e., water molecules at the bottom are separated farther apart. These hot molecules have high kinetic energy and rise upward to cold region while the molecules from cold region come down to take their place. Thus each molecule at the bottom gets heated and rises then cools and descends. This action sets up the flow of water molecules called convection currents. The convection currents transfer heat to the entire mass of water. Note that transfer of heat is by the bodily/ physical movement of the water molecules.

Always remember:

The process by which heat is transmitted through a substance from one point to another due to the actual bodily movement of the heated particles of the substance is called convection.

7.9.2.1 Applications of Convection:**i) Heating and cooling of rooms**

The mechanism of heating a room by a heat convector or heater is entirely based on convection. The air molecules in immediate contact with the heater are heated up. These air molecules acquire sufficient energy and rise upward. The cool air at the top being denser moves down to take their place. This cool air in turn gets heated and moves upward. In this way, convection currents are set up in the room which transfer heat to different parts of the room. The same principle but in opposite direction is used to cool a room by an air-conditioner.

ii) Cooling of transformers

Due to current flowing in the windings of the transformer, enormous heat is produced. Therefore, transformer is always kept in a tank containing oil. The oil in contact with transformer body heats up, creating convection currents. The warm oil comes in contact with the cooler tank, gives heat to it and descends to the bottom. It again warms up to rise upward. This process is repeated again and again. The heat of the transformer body is thus carried away by convection to the cooler tank. The cooler tank, in turn loses its heat by convection to the surrounding air.

7.9.2.2 Free and Forced Convection:

- i) When a hot body is in contact with air under ordinary conditions, like air around a firewood, the air removes heat from the body by a process called free or natural convection. Land and sea breezes are also formed as a result of free convection currents in air.
- ii) The convection process can be accelerated by employing a fan to create a rapid

circulation of fresh air. This is called forced convection. Example in section 7.9.2.1 are of forced convection, namely, heat convector, air conditioner, heat radiators in IC engine etc.

7.9.3 Radiation:

The transfer of heat energy from one place to another via emission of EM energy (in a straight line with the speed of light) without heating the intervening medium is called radiation.

For transfer of heat by radiation, molecules are not needed i.e. medium is not required. The fact that Earth receives large quantities of heat from the Sun shows that heat can pass through empty space (i.e., vacuum) between the Sun and the atmosphere that surrounds the Earth. In fact, transfer of heat by radiation has the same properties as light (or EM wave).

A natural question arises as to how heat transfer occurs in the absence of a medium (i.e., molecules). All objects possess thermal energy due to their temperature T ($T > 0 \text{ K}$). The rapidly moving molecules of a hot body emit EM waves travelling with the velocity of light. These are called **thermal radiations**. These carry energy with them and transfer it to the low-speed molecules of a cold body on which they fall. This results in an increase in the molecular motion of the cold body and its temperature rises. Thus transfer of heat by radiation is a two-fold process- the conversion of thermal energy into waves and reconversion of waves into thermal energy by the body on which they fall. We will learn about EM waves in Chapter 13.

7.10 Newton's Laws of Cooling:

If hot water in a vessel is kept on table, it begins to cool gradually. To study how a given body can cool on exchanging heat with its surroundings, following experiment is performed.

A calorimeter is filled up to two third of its capacity with boiling water and is covered. A thermometer is fixed through a hole in the lid and its position is adjusted so that the bulb of the thermometer is fully immersed in water.

The calorimeter vessel is kept in a constant temperature enclosure or just in open air since room temperature will not change much during experiment. The temperature on the thermometer is noted at one minute interval until the temperature of water decreases by about 25 °C. A graph of temperature T (along y-axis) is plotted against time t (along x-axis). This graph is called cooling curve (Fig 7.13 (a)). From this graph you can infer how the cooling of hot water depends on the difference of its temperature from that of its surroundings. You will also notice that initially the rate of cooling is higher and it decreases as the temperature of the water falls. A tangent is drawn to the curve at suitable points on the curve. The slope of each tangent (dT/dt) gives the rate of fall of temperature at that temperature. Taking (0,0) as the origin, if a graph of dT/dt is plotted against corresponding temperature difference ($T-T_0$), the curve is a straight line as shown in Fig 7.13 (b).

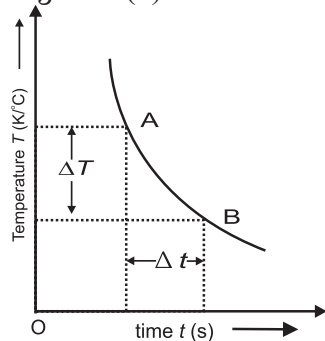


Fig 7.13 (a):
Temperature versus time graph. $\lim_{\Delta t \rightarrow 0} \frac{\Delta T}{\Delta t}$ gives the slope of the tangent drawn to the curve at point A and indicates the rate of fall of temperature.

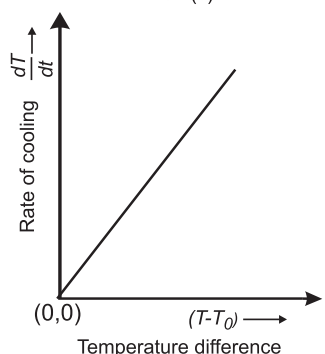


Fig 7.13 (b):
Rate of change of temperature versus time graph.

The above activity shows that a hot body loses heat to its surroundings in the form of heat radiation. The rate of loss of heat depends on the difference in the temperature of the body and its surroundings. Newton was the first to study the relation between the heat lost by a body in a given enclosure and its temperature in a systematic manner.

According to Newton's law of cooling the rate of loss of heat dT/dt of the body is directly proportional to the difference of temperature ($T - T_0$) of the body and the surroundings provided the difference in temperatures is small. Mathematically this may be expressed as

$$\frac{dT}{dt} \propto (T - T_0)$$

$$\therefore \frac{dT}{dt} = C (T - T_0) \quad \text{--- (7.39)}$$

where C is constant of proportionality.

Example 7.20: A metal sphere cools at the rate of 1.6 °C/min when its temperature is 70°C. At what rate will it cool when its temperature is 40°C. The temperature of surroundings is 30°C.

Solution: Given $T_1 = 70^\circ \text{C}$
 $T_2 = 40^\circ \text{C}$
 $T_0 = 30^\circ \text{C}$

$$\left(\frac{dT}{dt} \right)_1 = 1.6^\circ \text{C/min}$$

According to Newton's law of cooling, if C is the constant of proportionality

$$\left(\frac{dT}{dt} \right)_1 = C (T_1 - T_0)$$

$$\text{or, } 1.6 = C (70 - 30)$$

$$\therefore C = \frac{1.6}{40} = 0.04 / \text{min}$$

$$\text{Also } \left(\frac{dT}{dt} \right)_2 = C (T_2 - T_0)$$

$$= 0.04(40 - 30) = 0.4^\circ \text{C/min}$$

Thus the rate of cooling drops by a factor of four when the difference in temperature of the metal sphere and its surroundings drops by a factor of four.



Internet my friend

1. <https://hyperphysics.phy-astr.gsu.edu/hbase/hframe.html>
2. <https://youtu.be/7ZKHc5J6R5Q>
3. <https://physics.info/expansion>



Exercises

1. Choose the correct option.

- i) Range of temperature in a clinical thermometer, which measures the temperature of human body, is
(A) 70°C to 100°C
(B) 34°C to 42°C
(C) 0°F to 100°F
(D) 34°F to 80°F
- ii) A glass bottle completely filled with water is kept in the freezer. Why does it crack?
(A) Bottle gets contracted
(B) Bottle is expanded
(C) Water expands on freezing
(D) Water contracts on freezing
- iii) If two temperatures differ by 25°C on Celsius scale, the difference in temperature on Fahrenheit scale is
(A) 65° (B) 45°
(C) 38° (D) 25°
- iv) If α , β and γ are coefficients of linear, area l and volume expansion of a solid then
(A) $\alpha:\beta:\gamma$ 1:3:2 (B) $\alpha:\beta:\gamma$ 1:2:3
(C) $\alpha:\beta:\gamma$ 2:3:1 (D) $\alpha:\beta:\gamma$ 3:1:2
- v) Consider the following statements-
(I) The coefficient of linear expansion has dimension K^{-1}
(II) The coefficient of volume expansion has dimension K^{-1}
(A) I and II are both correct
(B) I is correct but II is wrong
(C) II is correct but I is wrong
(D) I and II are both wrong
- vi) Water falls from a height of 200 m. What is the difference in temperature between the water at the top and bottom of a water fall given that specific heat of water is $4200 \text{ J kg}^{-1} \text{ }^{\circ}\text{C}^{-1}$?
(A) 0.96°C (B) 1.02°C
(C) 0.46°C (D) 1.16°C

2. Answer the following questions.

- i) Clearly state the difference between heat and temperature?
- ii) How a thermometer is calibrated?
- iii) What are different scales of temperature? What is the relation between them?

- iv) What is absolute zero?
- v) Derive the relation between three coefficients of thermal expansion.
- vi) State applications of thermal expansion.
- vii) Why do we generally consider two specific heats for a gas?
- viii) Are freezing point and melting point same with respect to change of state? Comment.
- ix) Define (i) Sublimation (ii) Triple point.
- x) Explain the term 'steady state'.
- xi) Define coefficient of thermal conductivity. Derive its expression.
- xii) Give any four applications of thermal conductivity in every day life.
- xiii) Explain the term thermal resistance. State its SI unit and dimensions.
- xiv) How heat transfer occurs through radiation in absence of a medium?
- xv) State Newton's law of cooling and explain how it can be experimentally verified.
- xvi) What is thermal stress? Give an example of disadvantages of thermal stress in practical use?
- xvii) Which materials can be used as thermal insulators and why?

3. Solve the following problems.

- i) A glass flask has volume $1 \times 10^{-4} \text{ m}^3$. It is filled with a liquid at 30°C . If the temperature of the system is raised to 100°C , how much of the liquid will overflow. (Coefficient of volume expansion of glass is $1.2 \times 10^{-5} ({}^{\circ}\text{C})^{-1}$ while that of the liquid is $75 \times 10^{-5} ({}^{\circ}\text{C})^{-1}$).

[Ans : $516.6 \times 10^{-8} \text{ m}^3$]

- ii) Which will require more energy, heating a 2.0 kg block of lead by 30 K or heating a 4.0 kg block of copper by 5 K? ($s_{\text{lead}} = 128 \text{ J kg}^{-1} \text{ K}^{-1}$, $s_{\text{copper}} = 387 \text{ J kg}^{-1} \text{ K}^{-1}$)

[Ans : copper]

- iii) Specific latent heat of vaporization of water is $2.26 \times 10^6 \text{ J/kg}$. Calculate the energy needed to change 5.0 g of water

- into steam at 100 °C.
[Ans : 11.3×10^3 J]
- iv) A metal sphere cools at the rate of 0.05 °C/s when its temperature is 70 °C and at the rate of 0.025 °C/s when its temperature is 50 °C. Determine the temperature of the surroundings and find the rate of cooling when the temperature of the metal sphere is 40 °C.
[Ans : 30 °C, 0.0125 °C/s]
- v) The volume of a gas varied linearly with absolute temperature if its pressure is held constant. Suppose the gas does not liquefy even at very low temperatures, at what temperature the volume of the gas will be ideally zero?
[Ans : -273.15 °C]
- vi) In olden days, while laying the rails for trains, small gaps used to be left between the rail sections to allow for thermal expansion. Suppose the rails are laid at room temperature 27 °C. If maximum temperature in the region is 45 °C and the length of each rail section is 10 m, what should be the gap left given that $\alpha = 1.2 \times 10^{-5} \text{ K}^{-1}$ for the material of the rail section?
[Ans : 2.16 mm]
- vii) A blacksmith fixes iron ring on the rim of the wooden wheel of a bullock cart. The diameter of the wooden rim and the iron ring are 1.5 m and 1.47 m respectively at room temperature of 27 °C. To what temperature the iron ring should be heated so that it can fit the rim of the wheel ($\alpha_{\text{iron}} = 1.2 \times 10^{-5} \text{ K}^{-1}$).
[Ans: 1727.7 °C]
- viii) In a random temperature scale X, water boils at 200 °X and freezes at 20 °X. Find the boiling point of a liquid in this scale if it boils at 62 °C.
[Ans: 131.6°X]
- ix) A gas at 900°C is cooled until both its pressure and volume are halved. Calculate its final temperature.
[Ans: 293.29K]
- x) An aluminium rod and iron rod show 1.5 m difference in their lengths when heated at all temperature. What are their lengths at 0 °C if coefficient of linear expansion for aluminium is $24.5 \times 10^{-6} / ^\circ\text{C}$ and for iron is $11.9 \times 10^{-6} / ^\circ\text{C}$
[Ans: 1.417m, 2.917m]
- xi) What is the specific heat of a metal if 50 cal of heat is needed to raise 6 kg of the metal from 20°C to 62 °C ?
[Ans: $s = 0.198 \text{ cal/kg } ^\circ\text{C}$]
- xii) The rate of flow of heat through a copper rod with temperature difference 30 °C is 1500 cal/s. Find the thermal resistance of copper rod.
[Ans: 0.02 °C s cal]
- xiii) An electric kettle takes 20 minutes to heat a certain quantity of water from 0°C to its boiling point. It requires 90 minutes to turn all the water at 100°C into steam. Find the latent heat of vaporisation. (Specific heat of water = 1 cal/g°C)
[Ans: 450 cal/g]
- xiv) Find the temperature difference between two sides of a steel plate 4 cm thick, when heat is transmitted through the plate at the rate of 400 k cal per minute per square metre at steady state. Thermal conductivity of steel is 0.026 kcal/m s K.
[Ans: 10.26°C or 10.26 K]
- xv) A metal sphere cools from 80 °C to 60 °C in 6 min. How much time will it take to cool from 60 °C to 40 °C if the room temperature is 30°C?
[Ans: 10 min]
- ***



Can you recall?

- | | |
|---------------------------------------|-------------------------------------|
| 1. What type of wave is a sound wave? | 2. Can sound travel in vacuum? |
| 3. What are reverberation and echo? | 4. What is meant by pitch of sound? |

8.1 Introduction:

We are all aware of the ripples created on the surface of water when a stone is dropped in it. The water molecules oscillate up and down around their equilibrium positions but they do not move from one point to another along the surface of water. The disturbance created by dropping the stone, however travels outwards. This type of wave is a periodic and regular disturbance in a medium which does not cause any flow of material but causes the flow of energy and momentum from one point to another. There are different types of waves and not all types require material medium to travel through. We know that light is a type of wave and it can travel through vacuum. Here we will first study different types of waves, learn about their common properties and then study sound waves in particular.

Types of waves:

(i) **Mechanical waves:** A wave is said to be mechanical if a material medium is essential for its propagation. Examples of these types of waves are water waves, waves along a stretched string, seismic waves, sound waves, etc.

(ii) **EM waves:** These are generated due to periodic vibrations in electric and magnetic fields. These waves can propagate through material media, however, material medium is not essential for their propagation. These will be studied in Chapter 13.

(iii) **Matter waves:** There is always a wave associated with any object if it is in motion. Such waves are matter waves. These are studied in quantum mechanics.

Travelling or progressive waves are waves in which a disturbance created at one place travels to distant points and keeps travelling unless stopped by some external

agencies. In such types of waves energy gets transferred from one point to another. Water waves mentioned above are travelling waves. They keep travelling outward from the point where stone was dropped until they are stopped by walls of the container or the boundary of the water body. Other type of waves are stationary waves about which we will learn in XIIth standard.

8.2 Common Properties of all Waves:

The properties described below are valid for all types of waves, however, here they are described for mechanical waves.

1) Amplitude (A): Amplitude of a wave motion is the largest displacement of a particle of the medium through which the wave is propagating, from its rest position. It is measured in metre in SI units.

2) Wavelength (λ): Wavelength is the distance between two successive particles which are in the same state of vibration. It is further explained below. It is measured in metre.

3) Period (T): Time required to complete one vibration by a particle of the medium is the period T of the wave. It is measured in seconds.

4) Double periodicity: Waves possess double periodicity. At every location the wave motion repeats itself at equal intervals of time, hence it is periodic in time. Similarly, at any given instant, the form of wave repeats at equal distances hence, it is periodic in space. In this way wave motion is a doubly periodic phenomenon i.e periodic in time and periodic in space.

5) Frequency (n): Frequency of a wave is the number of vibrations performed by a particle during each second. SI unit of frequency is hertz. (Hz) Frequency is a reciprocal of time period, i.e., $n = \frac{1}{T}$

6) Velocity (v): The distance covered by a wave per unit time is called the velocity of the wave. During the period (T), the wave covers a distance equal to the wavelength (λ). Therefore the magnitude of velocity of wave is given by,

$$\text{Magnitude of velocity} = \frac{\text{distance}}{\text{time}}$$

$$v = \frac{\text{wavelength}}{\text{period}}$$

$$v = \frac{\lambda}{T} \quad \text{--- (8.1)}$$

$$\text{but } \frac{1}{T} = n \text{ (frequency)} \quad \text{--- (8.2)}$$

$$\therefore v = n\lambda \quad \text{--- (8.3)}$$

This equation indicates that, the magnitude of velocity of a wave in a medium is constant. Increase in frequency of a wave causes decrease in its wavelength. When a wave goes from one medium to another medium, the frequency of the wave does not change. In such a case speed and wavelength of the wave change.

For mechanical waves to propagate through a medium, the medium should possess certain properties as given below:

- The medium should be continuous and elastic so that the medium regains original state after removal of deforming forces.
- The medium should possess inertia. The medium must be capable of storing energy and of transferring it in the form of waves.
- The frictional resistance of the medium must be negligible so that the oscillations will not be damped.

7) Phase and phase difference:

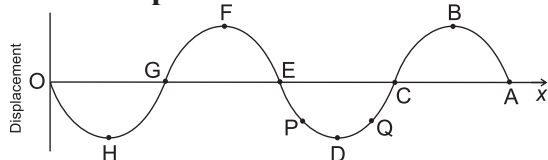


Fig. 8.1 (a): Displacement as a function of distance along the wave.

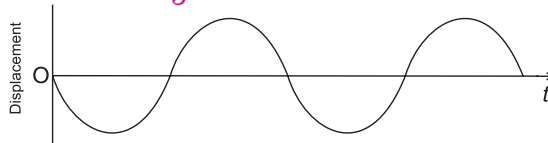


Fig. 8.1 (b): Displacement as a function of time.

In the Fig. 8.1 (a), displacements of various particles along a sinusoidal wave travelling along +ve x -axis are plotted against their respective distances from the source (at O) at a given instant. This plot is valid for transverse as well as longitudinal wave.

The state of oscillation of a particle is called its **phase**. In order to describe the phase at a place, we need to know (a) the displacement (b) the direction of velocity and (c) the oscillation number (during which oscillation) of the particle there.

In Fig. 8.1 (a), particles P and Q (or E and C or B and D) have same displacements but the directions of their velocities are opposite. Particles B and F have same magnitude of displacements and same direction of velocity. Such particles are said to be in phase during their respective oscillations. Also, these are successive particles with this property of having same phase. Separation between these two particles is wavelength λ . These two successive particles differ by '1' in their oscillation number, i.e., if particle B is at its n^{th} oscillation, particle F will be at its $(n+1)^{\text{th}}$ oscillation as the wave is travelling along + x direction. Most convenient way to understand phase is in terms of angle. For a sinusoidal wave, the variation in the displacement is a 'sine' function of distance from the source and of time as discussed below. For such waves it is possible for us to assign angles corresponding to the displacement (or time).

At the instant the above graph is drawn, the disturbance (energy) has just reached the particle A. The phase angle corresponding to this particle A can be taken as 0° . At this instant, particle B has completed quarter oscillation and reached its positive maximum ($\sin \theta = +1$). The phase angle θ of this particle B is $\pi/2 = 90^\circ$ at this instant. Similarly, phase angles of particles C and E are π° (180°) and $2\pi^\circ$ (360°) respectively. Particle F has completed one oscillation and is at its positive maximum during its second oscillation. Hence its phase angle is $2\pi^\circ + \frac{\pi^\circ}{2} = \frac{5\pi^\circ}{2}$.

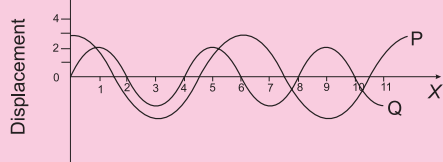
B and F are the successive particles in the same state (same displacement and same direction of velocity) during their respective oscillations. Separation between these two is wavelength (λ). Phase angle between these two differs by $2\pi^\circ$. Hence wavelength is better understood as the separation between two particles with phase difference of $2\pi^\circ$.

As noted above, waves possess double periodicity. This means the displacements of particles are periodic in space (as shown in Fig. 8.1 (a)) as well as periodic time. Figure 8.1 (b) shows the displacement of one particular particle as a function of time.



Activity :

- (1) Using axes of displacement and distance, sketch two waves A and B such that A has twice the wavelength and half the amplitude of B.
- (2) Determine the wavelength and amplitude of each of the two waves P and Q shown in figure below.



Characteristics of progressive wave

- 1) All vibrating particles of the medium have same amplitude, period and frequency.
- 2) State of oscillation i.e., phase changes from particle to particle.

Example 8.1: The speed of sound in air is 330 m/s and that in glass is 4500 m/s. What is the ratio of the wavelength of sound of a given frequency in the two media?

Solution:

$$\begin{aligned}
 v_{\text{air}} &= n \lambda_{\text{air}} \\
 v_{\text{glass}} &= n \lambda_{\text{glass}} \\
 \therefore \frac{\lambda_{\text{air}}}{\lambda_{\text{glass}}} &= \frac{v_{\text{air}}}{v_{\text{glass}}} = \frac{330}{4500} = 7.33 \times 10^{-2} \\
 &= 0.0733 \approx 7.33 \times 10^{-2}
 \end{aligned}$$

8.3 Transverse Waves and Longitudinal Waves:

Progressive waves can be of two types, transverse and longitudinal waves.

Transverse waves : A wave in which particles of the medium vibrate in a direction perpendicular to the direction of propagation of wave is called transverse wave. Water waves are transverse waves, as water molecules vibrate perpendicular to the surface of water while the wave propagates along the surface.

Characteristics of transverse waves.

- 1) All particles of the medium in the path of the wave vibrate in a direction perpendicular to the direction of propagation of the wave with same period and amplitude.
- 2) When transverse wave passes through a medium, the medium is divided into alternate the crests i.e., regions of positive displacements and troughs i.e., regions of negative displacements.
- 3) A crest and an adjacent trough form one cycle of a transverse wave. The distance measured along the wave between any two consecutive points in the same phase (crest or trough) is called the wavelength of the wave.
- 4) Crests and troughs advance in the medium and are responsible for transfer of energy.
- 5) Transverse waves can travel through solids and on surfaces of liquids only. They can not travel through liquids and gases. EM waves are transverse waves but they do not require material medium for propagation.
- 6) When transverse waves advance through a medium there is no change in the pressure and density at any point of medium, however shape changes periodically.
- 7) If vibrations of all the particles along the path of a wave are constrained to be in a single plane, then the wave is called polarised wave. Transverse wave can be polarised.
- 8) Medium conveying a transverse wave must possess elasticity of shape.

Longitudinal waves : A wave in which particles of the medium vibrate in a direction parallel to the direction of propagation of wave is called longitudinal wave. Sound waves are longitudinal waves.

Characteristics of longitudinal waves:

- 1) All the particles of medium along the path of the wave vibrate in a direction parallel to the direction of propagation of wave with same period and amplitude.
- 2) When longitudinal wave passes through a medium, the medium is divided into regions of alternate compressions and rarefactions. Compression is the region where the particles of medium are crowded (high pressure zone), while rarefaction is the region where the particles of medium are more widely separated, i.e. the medium gets rarefied (low pressure zone).
- 3) A compression and adjacent rarefaction form one cycle of longitudinal wave. The distance measured along the wave between any two consecutive points having the same phase is the wavelength of wave.
- 4) For propagation of longitudinal waves, the medium should possess the property of elasticity of volume. Thus longitudinal waves can travel through solids, liquids and gases. Longitudinal wave can not travel through vacuum or free space.
- 5) The compression and rarefaction advance in the medium and are responsible for transfer of energy.
- 6) When longitudinal wave advances through a medium there are periodic variations in pressure and density along the path of wave and also with time.
- 7) Longitudinal waves can not be polarised, as the direction of vibration of particles and direction of propagation of wave are same or parallel.

8.4 Mathematical Expression of a Wave:

Let us describe a progressive wave mathematically. Since it is a progressive wave, we require a function of both the position x and time t . This function will describe the shape of the wave at any instant of time. Another

requirement of the function is that it should describe the motion of the particle of the medium at that point. A sinusoidal progressive wave can be described by a sinusoidal function. Let us assume that the progressive wave is transverse and, therefore, the position of the particle of the medium is described by a fixed value of x . The displacement from the equilibrium position can be described by y . Such a sinusoidal wave can be written as follows:

$$y(x, t) = a \sin(kx - \omega t + \phi) \quad \text{--- (8.4)}$$

Hence a , k , ω and ϕ are constants.

Let us see the justification for writing this equation. At a particular instant say $t = t_0$,

$$\begin{aligned} y(x, t_0) &= a \sin(kx - \omega t_0 + \phi) \\ &= a \sin(kx + \text{constant}) \end{aligned}$$

Thus the shape of the wave at $t = t_0$, as a function of x is a sine wave.

Also, at a fixed location $x = x_0$,

$$\begin{aligned} y(x_0, t) &= a \sin(kx_0 - \omega t + \phi) \\ &= a \sin(\text{constant} - \omega t) \end{aligned}$$

Hence the displacement y , at $x = x_0$ varies as a sine function.

This means that the particles of the medium, through which the wave travels, execute simple harmonic motion around their equilibrium position. In addition x must increase in the positive direction as time t increases, so as to keep $(kx - \omega t + \phi)$ a constant. Thus the Eq. (8.4) represents a wave travelling along the positive x axis. A wave represented by

$$y(x, t) = a \sin(kx + \omega t + \phi) \quad \text{--- (8.5)}$$

is a wave travelling in the direction of the negative x axis.

Symbols in Eq. (8.4):

$y(x, t)$ is the displacement as a function of position (x) and time (t)

a is the amplitude of the wave.

ω is the angular frequency of the wave

k is the angular wave number

$(kx_0 - \omega t + \phi)$ is the argument of the sinusoidal wave and is the phase of the particle at x at time t .

8.5 The Speed of Travelling Waves

Speed of a mechanical wave depends upon the elastic properties and density of the medium. The same medium can support both transverse and longitudinal waves which have different speeds.

8.5.1 The speed of transverse waves

The speed of a wave is determined by the restoring force produced in the medium when it is disturbed. The speed also depends on inertial properties like mass density of the medium. The waves produced on a string are transverse waves. In this case the restoring force is provided by the tension T in the string. The inertial property i.e. the linear mass density m , can be determined from the mass of string M and its length L as $m = M/L$. The formula for speed of transverse wave on stretched string is given by

$$v = \sqrt{\frac{T}{m}} \quad \text{--- (8.6)}$$

The derivation of the formula is beyond the scope of this book.

The important point here is that the speed of a transverse wave depends only on the properties of the string, T and m . It does not depend on wavelength or frequency of the wave.

8.5.2 The speed of longitudinal waves

In case of longitudinal waves, the particles of the medium oscillate forward and backward along the direction of wave propagation. This causes compression and rarefaction which travel in the medium as the medium possess elastic property.

Speed of sound in liquids and solids is higher than that in gases. The speed of sound as a longitudinal wave in an ideal gas is given by Newton's formula as discussed below. Speed of sound in different media is given in table below.

Always remember:

When a sound wave goes from one medium to another its velocity changes along with its wavelength. Its frequency, which is decided by the source remains constant.

Table 8.1: Speed of Sound in Gas, Liquids, and Solids

Medium	Speed (m/s)
<u>Gases</u>	
Air [0°C]	331
Air [20°C]	343
Helium	965
Hydrogen	1284
<u>Liquids</u>	
Water (0°C)	1402
Water (20°C)	1482
Seawater	1522
<u>Solids</u>	
Vulcanised Rubber	54
Copper	3560
Steel	5941
Granite	6000
Aluminium	6420

8.5.3 Newton's formula for velocity of sound:

Propagation of longitudinal waves was studied by Newton. Sound waves travel through a medium in the form of compressions and rarefactions. The density of medium is greater at the compression while being smaller in the rarefaction. Hence the velocity of sound depends on elasticity and density of the medium. Newton formulated the relation as

$$v = \sqrt{\frac{E}{\rho}} \quad \text{--- (8.7)}$$

where E is the proper modulus of elasticity of medium and ρ is the density of medium.

Newton assumed that, during propagation of sound, there is no change in the average temperature of the medium. Hence sound wave propagation in air is an isothermal process (temperature remaining constant) and isothermal elasticity should be considered. The volume elasticity of air determined under isothermal change is called isothermal bulk modulus and is equal to the atmospheric pressure ' P '. Hence Newton's formula for speed of sound in air is given by

$$v = \sqrt{\frac{P}{\rho}} \quad \text{--- (8.8)}$$

As atmospheric pressure is given by $P = h \rho g$ and at NTP,

$$h = 0.76 \text{ m of Hg}$$

$$d = 13600 \text{ kg/m}^3 \text{ - density of mercury}$$

$$\rho = 1.293 \text{ kg/m}^3 \text{ - density of air}$$

and $g = 9.8 \text{ m/s}^2$

$$v = \sqrt{\frac{0.76 \times 13600 \times 9.8}{1.293}}$$

$$v = 279.9 \text{ m/s at NTP.}$$

This is the value of velocity of sound according to Newton's formula. But the experimental value of velocity of sound at 0°C as determined earlier by a number of scientists is 332 m/s . The difference between predicted value by Newton's formula and experimental value is large and it is not due to experimental error. The Experimental value is 16% greater than the value given by the formula. Newton could not give satisfactory explanation of this discrepancy. It was resolved by French physicist Pierre Simon Laplace (1749-1827).

Example 8.2: Suppose you are listening to an out-door live concert sitting at a distance of 150 m from the speakers. Your friend is listening to the live broadcast of the concert in another country and the radio signal has to travel 3000 km to reach him. Who will hear the music first and what will be the time difference between the two? Velocity of light $= 3 \times 10^8 \text{ m/s}$ and that of sound is 330 m/s .

Solution: Time taken by sound to reach you

$$= \frac{150}{330} \text{ s} = 0.4546$$

Time taken by the broadcasted sound (which is done by EM waves having velocity $= 3 \times 10^8 \text{ m/s}$)

$$= \frac{3000 \text{ km}}{3 \times 10^5 \text{ km/s}} = \frac{3 \times 10^3}{3 \times 10^5} = 10^{-2} \text{ s}$$

$\therefore$ your friend will hear the sound first. The time difference will be

$$= 0.4546 - 0.01$$

$$= 0.4446 \text{ s.}$$

8.5.4 Laplace's correction

According to Laplace, the generation of compression and rarefaction is not a slow

process but is a rapid process. If frequency is 256 Hz , the air is compressed and rarefied 256 times in a second. Such process must be a rapid process. Heat is produced during compression and is lost during rarefaction. This heat does not get sufficient time for dissipation. Due to this the total heat content remains the same. Such a process is called an adiabatic process and hence, adiabatic elasticity must be adiabatic and not isothermal elasticity, as was assumed by Newton.

Always remember:

In isothermal process temperature remains constant while in adiabatic process there is neither transfer of heat nor of mass.

The adiabatic modulus of elasticity of air is given by,

$$E = \gamma P \quad \text{--- (8.9)}$$

where P is the pressure of the medium (air) and γ is ratio of specific heat of air at constant pressure (C_p) to the specific heat of air at constant volume (C_v) called as the adiabatic ratio

$$\text{i.e., } \gamma = \frac{C_p}{C_v} \quad \text{--- (8.10)}$$

For air the ratio of C_p / C_v is 1.41

$$\text{i.e. } \gamma = 1.41$$

Newton's formula for speed of sound in air as modified by Laplace to give

$$v = \sqrt{\frac{\gamma P}{\rho}} \quad \text{--- (8.11)}$$

According to this formula velocity of sound at NTP is

$$v = \sqrt{\frac{1.41 \times 0.76 \times 13600 \times 9.8}{1.293}}$$

$$= 332.3 \text{ m/s}$$

This value is in close agreement with the experimental value. As seen above, the velocity of sound depends on the properties of the medium.

8.5.5 Factors affecting speed of sound:

As sound waves travel through atmosphere (open air), some factors related to air affect the speed of sound.



Do you know ?

C_p , the specific heat of gas at constant pressure, is defined as the quantity of heat required to raise the temperature of unit mass of gas through 1°K when pressure remains constant.

C_v , the specific heat of gas at constant volume, is defined as the quantity of heat required to raise the temperature of unit mass of gas through 1°K when volume remains constant.

When pressure is kept constant the volume of the gas increases with increase in temperature. Thus additional heat is required to increase the volume of gas against the external pressure. Therefore heat required to raise the temperature of unit mass of gas through 1°K when pressure is kept constant is greater than the heat required when volume is kept constant. i.e. $C_p > C_v$.

a) Effect of pressure on velocity of sound

According to Laplace's formula velocity of sound in air is

$$v = \sqrt{\frac{\gamma P}{\rho}}$$

If M is the mass and V is volume of air then

$$\rho = \frac{M}{V}$$

$$\therefore v = \sqrt{\frac{\gamma PV}{M}} \quad \text{--- (8.12)}$$

At constant temperature $PV = \text{constant}$ according to Boyle's law. Also M and γ are constant, hence $v = \text{constant}$.

Therefore at constant temperature, a change in pressure has no effect on velocity of sound in air. This can be seen in another way. For gaseous medium, $PV = nRT$, n being the number of moles.

$$\therefore v = \sqrt{\frac{\gamma nRT}{M}} \quad \text{--- (8.13)}$$

Hence for gaseous medium obeying ideal gas equation change in pressure has no effect on velocity of sound unless there is change in temperature.

Example 8.3: Consider a closed box of rigid walls so that the density of the air inside it is constant. On heating, the pressure of this enclosed air is increased from P_0 to P . It is now observed that sound travels 1.5 times faster than at pressure P_0 calculate P/P_0 .

Solution:

$$v_P = \sqrt{\frac{\gamma P_0}{\rho}}$$

$$v_{P_0} = \sqrt{\frac{\gamma P_0}{\rho}}$$

$$v_P = 1.5 v_{P_0}$$

$$\sqrt{\frac{\gamma P}{\rho}} = 1.5 \sqrt{\frac{\gamma P_0}{\rho}}$$

$$\frac{P}{\rho} = 2.25 \frac{P_0}{\rho}$$

$$P = 2.25 P_0$$

(b) Effect of temperature on speed of sound

Suppose v_0 and v are the speeds of sound at T_0 and T in kelvin respectively. Let ρ_0 and ρ be the densities of gas at these two temperatures. The velocity of sound at temperature T_0 and T can be written by using Eq. (8.13),

$$v_0 = \sqrt{\frac{\gamma RT_0}{M}} \quad \text{--- } M \text{ is molar mass, } n = 1$$

$$v = \sqrt{\frac{\gamma RT}{M}}$$

$$\therefore \frac{v}{v_0} = \sqrt{\frac{RT}{RT_0}}$$

$$\therefore \frac{v}{v_0} = \sqrt{\frac{T}{T_0}} \quad \text{--- (8.14)}$$

This equation shows that speed of sound in air is directly proportional to the square root of absolute temperature. Thus, speed of sound in air increases with increase in temperature. Taking $T_0 = 273 \text{ K}$ and writing $T = (273 + t) \text{ K}$ where t is the temperature in degree celsius. The ratio of velocity of sound in air at $t^\circ \text{C}$ to that at 0°C is given by,

$$\therefore \frac{v}{v_0} = \sqrt{\frac{273+t}{273}}$$

$$\therefore \frac{v}{v_0} = \sqrt{1 + \frac{t}{273}}$$

$$\therefore \frac{v}{v_0} = \sqrt{1 + \alpha t} \quad \text{where } \alpha = \frac{1}{273}$$

$$\text{or, } v = v_0 (1 + \alpha t)^{\frac{1}{2}}$$

As α is very small, we can write

$$v \approx v_0 \left(1 + \frac{1}{2} \alpha t \right)$$

$$v = v_0 \left(1 + \frac{1}{2} \times \frac{1}{273} t \right)$$

$$v = v_0 \left(1 + \frac{t}{546} \right)$$

$$v = v_0 + \frac{v_0}{546} t$$

But $v_0 = 332 \text{ m/s}$ at 0°C

$$\therefore v = v_0 + \frac{332}{546} t$$

$$\therefore v \approx v_0 + (0.61)t, \quad \text{--- (8.15)}$$

i.e., for 1°C rise in temperature velocity increases by 0.61 m/s . Hence for small variations in temperature ($< 50^\circ\text{C}$), the speed of sound changes linearly with temperature.

(c) Effect of humidity on speed of sound

Humidity (moisture) in air depends upon the presence of water vapour in it. Let ρ_m and ρ_d be the densities of moist and dry air respectively. If v_m and v_d are the speeds of sound in moist air and dry air then using Eq. (8.11).

$$v_m = \sqrt{\frac{\gamma P}{\rho_m}}$$

$$\text{and } v_d = \sqrt{\frac{\gamma P}{\rho_d}}$$

$$\therefore \frac{v_m}{v_d} = \sqrt{\frac{\rho_d}{\rho_m}} \quad \text{--- (8.16)}$$

Moist air is always less dense than dry air, i.e.,

$$\rho_m < \rho_d$$

$$(\rho_m = 0.81 \text{ kg/m}^3 \text{ (at } 0^\circ\text{C)}) \text{ and}$$

$$\rho_d = 1.29 \text{ kg/m}^3 \text{ (at } 0^\circ\text{C)})$$

$$\therefore v_m > v_d.$$

Thus, the speed of sound in moist air is greater than speed of sound in dry air. i.e speed increases with increase in the moistness of air.

8.6 Principle of Superposition of Waves:

Waves don't display any repulsion towards each other. Therefore two wave patterns can overlap in the same region of the space without affecting each other. When two waves overlap their displacements add vectorially. This additive rule is referred to as the principle of superposition of waves.

When two or more waves travelling through a medium arrive at a point of medium simultaneously, each wave produces its own displacement at that point independent of the others. Hence the resultant displacement at that point is equal to the vector sum of the displacements due to all the waves. The phenomenon of superposition will be discussed in detail in XIIth standard.

8.7 Echo, reverberation and acoustics:

Sound waves obey the same laws of reflection as those of light.

8.7.1 Echo:

An echo is the repetition of the original sound because of reflection from some rigid surface at a distance from the source of sound. If we shout in a hilly region, we are likely to hear echo.

Why can't we hear an echo at every place? At 22°C , the velocity of sound in air is 344 m/s . Our brain retains sound for 0.1 second. Thus for us to hear a distinct echo, the sound should take more than 0.1s after starting from the source (i.e., from us) to get reflected and come back to us.

$$\begin{aligned} \text{distance} &= \text{speed} \times \text{time} \\ &= 344 \times 0.1 \\ &= 34.4 \text{ m.} \end{aligned}$$

To be able to hear a distinct echo, the reflecting surface should be at a minimum

distance of half of the above distance i.e 17.2 m. As velocity depends on the temperature of air, this distance will change with temperature.

Example 8.4: A man shouts loudly close to a high wall. He hears an echo. If the man is at 40 m from the wall, how long after the shout will the echo be heard ? (speed of sound in air = 330 m/s)

solution: The distance travelled by the sound wave

$$\begin{aligned} &= 2 \times \text{distance from man to wall.} \\ &= 2 \times 40 \\ &= 80 \text{ m.} \end{aligned}$$

$$\begin{aligned} \therefore \text{Time taken to travel the distance} &= \frac{\text{distance}}{\text{speed}} \\ &= \frac{80 \text{ m}}{330 \text{ m/s}} \\ &= 0.24 \text{ s} \end{aligned}$$

$\therefore$ The man will hear the echo 0.24 s after he shouts.

8.7.2 Reverberation:

If the reflecting surface is nearer than 15 m from the source of sound, the echo joins up with the original sound which then seems to be prolonged. Sound waves get reflected multiple times from the walls and roof of a closed room which are nearer than 15 m. This causes a single sound to be heard not just once but continuously. This is called reverberation. It is this the persistence of sound after the source has switched off, as a result of repeated reflection from walls, ceilings and other surfaces. Reverberation characteristics are important in the design of concert halls, theatres etc.

If the time between successive reflections of a particular sound wave reaching us is small, the reflected sound gets mixed up and produces a continuous sound of increased loudness which can't be heard clearly.

Reverberation can be decreased by making the walls and roofs rough and by using curtains in the hall to avoid reflection of sound. Chairs and wall surfaces are covered with sound absorbing materials. Porous cardboard sheets, perforated acoustic tiles, gypsum boards, thick curtains etc. at the ceilings and at the walls are most convenient to reduce reverberation.

8.7.3 Acoustics:

The branch of physics which deals with the study of production, transmission and reception of sound is called acoustics. This is useful during the construction of theaters and auditorium. While designing an auditorium, proper care for the absorption and reflection of sound should be taken. Otherwise audience will not be able to hear the sound clearly.

For proper acoustics in an auditorium the following conditions must be satisfied.

- 1) The sound should be heard sufficiently loudly at all the points in the auditorium. The surface behind the speaker should be parabolic with the speaker at its focus; so that the distribution of sound is uniform in the auditorium. Reflection of sound is helpful in maintaining good loudness through the entire auditorium.
- 2) Echoes and reverberation must be eliminated or reduced. Echoes can be reduced by making the reflecting surfaces more absorptive. Echo will be less if the auditorium is full.
- 3) Unnecessary focusing of sound should be avoided and there should not be any zone of poor audibility or region of silence. For that purpose curved surface of the wall or ceiling should be avoided.
- 4) Echelon effect : It is due to the mixing of sound produced in the hall by the echoes of sound produced in front of regular structure like the stairs. To avoid this, stair type construction must be avoided in the hall.
- 5) The auditorium should be sound-proof when closed, so that stray sound can not enter from outside.
- 6) For proper acoustics no sound should be produced from the inside fittings, seats, etc. Instead of fans, air conditioners may be used. Soft action door closers should be used.

Acoustics observed in nature

The importance of acoustic principles goes far beyond human hearing. Several animals use

sound for navigation.

- (a) Bats depends on sound rather than light to locate objects. So they can fly in total darkness of caves. They emit short ultrasonic pulses of frequency 30 kHz to 150 kHz. The resulting echoes give them information about location of the obstacle.
- (b) Dolphins use an analogous system for underwater navigation. The frequencies are subsonic about 100 Hz. They can sense an object of about the size of a wavelength i.e., 1.4 m or larger.

Medical applications of acoustics

- (a) Shock waves which are high pressure high amplitude waves are used to split kidney stones into smaller pieces without invasive surgery. A shock wave is produced outside the body and is then focused by a reflector or acoustic lens so that as much of its energy as possible converges on the stone. When the resulting stresses in the stone exceeds its tensile strength, it breaks into small pieces which can be removed easily.
- (b) Reflection of ultrasonic waves from regions in the interior of body is used for ultrasonic imaging. It is used for prenatal (before the birth) examination, detection of anomalous conditions like tumour etc and the study of heart valve action.
- (c) At very high power level, ultrasound is selective destroyer of pathological tissues in treatment of arthritis and certain type of cancer.

Other applications of acoustics

- (a) SONAR is an acronym for Sound Navigational Ranging. This is a technique for locating objects underwater by transmitting a pulse of ultrasonic sound and detecting the reflected pulse. The time delay between transmission of a pulse and the reception of reflected pulse indicates the depth of the object. This system is useful to measure motion and position of the submerged objects like submarine.
- (b) Acoustic principle has important application to environmental problems like noise control. The design of quiet-

mass transit vehicle involves the study of generation and propagation of sound in the motor's wheels and supporting structures.

- (c) We can study properties of the Earth by measuring the reflected and refracted elastic waves passing through its interior. It is useful for geological studies to detect local anomalies like oil deposits etc.

8.8 Qualities of sound:

Audible sound or human response to sound:

Whenever we talk about *audible* sound, what matters is how we *perceive* it. This is purely a subjective attribute of sound waves.

Major qualities of sound that are of our interest are (i) Pitch, (ii) Timbre or quality and (iii) Loudness.

(i) Pitch:

This aspect refers to sharpness or shrillness of the sound. If the frequency of sound is increased, what we perceive is the increase in the pitch or we feel the sound to be sharper. *Tone* refers to the single frequency of that wave while a note may contain one or more than one tones. We use the words high pitch or high tones if frequency is higher. As sharpness is a subjective term, sentences like “sound of double frequency is doubly sharp” make no sense. Also, a high pitch sound *need not* be louder. Tones of guitar are sharper than that of a base guitar, sound of tabla is sharper than that of a daga, (in general) female sound is sharper than that of a male sound and so on.

For a sound amplifier (or equaliser) when we raise the *treble* knob (or *treble* Button), high frequencies are boosted and if we raise *bass* knob, low frequencies are boosted.

(ii) Timbre (sound quality)

During telephonic conversation with a friend, (mostly) you are able to know who is speaking at the other end even if you are not told about who is speaking. Quite often we say, “Couldn't you recognise the voice?” The sound quality in this context is called *timbre*. Same song played on a guitar, a violin, a harmonium or a piano feel significantly different and we can easily identify that instrument. Quality of sound of any sound instrument (including

our vocal organ) depends upon the mixture of tones and overtones in the sound generated by that instrument. Even our own sound quality during morning (after we get up) and in the evening is different. It is drastically affected if we are suffering from *cold* or *cough*. Concept of overtones will be discussed during XIIth standard.

(iii) Loudness:

Intensity of a wave is a measurable quantity which is proportional to square of the amplitude ($I \propto A^2$) and is measured in the (SI) unit of W/m^2 . Human perception of intensity of sound is *loudness*. Obviously, if intensity is more, loudness is more. The human response to intensity is not linear, i.e., a sound of double intensity is louder but not doubly loud. This is also valid for brightness of light. In both cases, the response is approximately logarithmic. Using this property, the loudness (and brightness) can be measured.

Under ideal conditions, for a perfectly healthy human ear, the least audible intensity is $I_0 = 10^{-12} \text{ W/m}^2$. Loudness of a sound of intensity I , measured in the unit **bel** is given by

$$L_{\text{bel}} = \log_{10} \left(\frac{I}{I_0} \right) \quad \text{--- (8.17)}$$

Popular or commonly used unit for loudness is *decibel*. We know, 1 *decimetre* or 1 dm = 0.1m. Similarly, 1 *decibel* or 1 db = 0.1 bel. $\therefore$ 1 bel = 10 db. Thus, loudness expressed in db is 10 times the loudness in bel

$$\therefore L_{\text{db}} = 10 L_{\text{bel}} = 10 \log_{10} \left(\frac{I}{I_0} \right)$$

For sound of least audible intensity I_0 ,

$$L_{\text{db}} = 10 \log_{10} \left(\frac{I_0}{I_0} \right) = 10 \log_{10} (1) = 0 \quad \text{--- (8.18)}$$

This corresponds to threshold of hearing

For sound of 10 db,

$$10 = 10 \log_{10} \left(\frac{I}{I_0} \right) \therefore \left(\frac{I}{I_0} \right) = 10^1 \text{ or } I = 10 I_0$$

For sound of 20 db,

$$20 = 10 \log_{10} \left(\frac{I}{I_0} \right) \therefore \left(\frac{I}{I_0} \right) = 10^2 \text{ or } I = 100 I_0$$

and so on.

Hence, loudness of 20 db sound is *felt* double that of 10 db, but its intensity is 10 times that of the 10 db sound. Now, we feel 40 db sound twice as loud as 20 db sound but its intensity is 100 times as that of 20 db sound and 10000 times that of 10 db sound. This is the power of logarithmic or exponential scale.

If we move away from a (practically) point source, the intensity of its sound varies inversely

with square of the distance, i.e., $I \propto \frac{1}{r^2}$.

Whenever you are using earphones or jam your mobile at your ear, the distance from the source is too small. Obviously, such a habit for a long time can affect your normal hearing.

Example 8.5: When heard independently, two sound waves produce sensations of 60 db and 55 db respectively. How much will the sensation be if those are sounded together, perfectly in phase?

Solution:

$$L_1 = 60 \text{ db} = 10 \log_{10} \frac{I_1}{I_0} \therefore \frac{I_1}{I_0} = 10^6 \text{ or } I_1 = 10^6 I_0$$

$$\text{Similarly, } I_2 = 10^{5.5} I_0$$

As the waves combine perfectly in phase, the vector addition of their amplitudes will be given by $A^2 = (A_1 + A_2)^2 = A_1^2 + A_2^2 + 2A_1A_2$

As intensity is proportional to square of the amplitude.

$$\begin{aligned} \therefore I &= I_1 + I_2 + 2\sqrt{I_1 I_2} \\ &= 10^5 I_0 \left(10^1 + 10^{0.5} + 2\sqrt{10^{1.5}} \right) \\ &= 10^5 I_0 \left(10 + 3.1623 + 2 \times 10^{0.75} \right) \\ &= 24.41 \times 10^5 I_0 = 2.441 \times 10^6 I_0 \end{aligned}$$

$$\begin{aligned} \therefore L &= 10 \log_{10} \left(\frac{I}{I_0} \right) = 10 \log_{10} (2.441 \times 10^6) \\ &= 10 \left[\log_{10} (2.441) + \log_{10} (10^6) \right] \\ &= 10 (0.3876 + 6) \\ &= 63.876 \text{ db} \approx 64 \text{ db} \end{aligned}$$

It is interesting to note that there is only a marginal increase in the loudness.

Table 8.2: Approximate Decibel Ratings of Some Audible Sounds

Source or description of noise	Loudness, L_{db}	Effect
Extremely loud	160	Immediate ear damage
Jet aeroplane, near 25 m	150	Rupture of eardrum
Auto horn, within a metre, Aircraft take off, 60 m	110	Strongly painful
Diesel train, 30 m, Average factory	80	
Highway traffic, 8 m	70	Uncomfortable
Conversation at a restaurant	60	
Conversation at home	50	
Quiet urban background sound	40	
Quiet rural area	30	Virtual silence
Whispering of leaves, 5 m	20	
Normal breathing	10	
Threshold of hearing	0	

8.9 Doppler Effect:

Have you ever heard an approaching train and noticed distinct change in the pitch of the sound of its whistle, when it passes away? Same thing similar happens when a listener moves towards or away from the stationary source of sound. Such a phenomenon was first identified in 1842 by Austrian physicist Christian Doppler (1803-1853) and is known as Doppler effect.

When a source of sound and a listener are in motion relative to each other the frequency of sound heard by listener is not the same as the frequency emitted by the source.

Doppler effect is the apparent change in frequency of sound due to relative motion between the source and listener. Doppler effect is a wave phenomenon. It holds for sound waves and also for EM waves. But here we shall consider it for sound waves only.

The changes in frequency can be studied under 3 different conditions:

- 1) When listener is stationary but source is moving.
- 2) When listener is moving but source is stationary.
- 3) When listener and source both are moving.



Do you know ?

According to the world health organisation a billion young people could be at risk of hearing loss due to unsafe listening practices. Among teenagers and young adults aged 12-35 years (i) about 50% are exposed to unsafe levels of sound from use of personal audio devices and (ii) about 40% are exposed to potentially damaging sound levels at clubs, discotheques and bars.

8.9.1 Source Moving and Listener Stationary:

Consider a source of sound S , moving away from a stationary listener L (called relative recede) with velocity v_s . Speed of sound waves with respect to the medium is v which is always positive. Suppose the listener uses a detector for counting each wave crest that reaches it.

Initially (at $t = 0$), source which is at point S_1 emits a crest when at distance d from the listener see Fig. 8.2 (a). This crest reaches the listener at time $t_1 = d/v$. Let T_0 be the time period at which the waves are emitted. Thus, at $t = T_0$ the source moves the distance $= v_s T_0$ and reaches the point S_2 . Distance of S_2 from the listener is $(d + v_s T_0)$. when at S_2 , the source emits second crest. This crest reaches the listener at

$$t_2 = T_0 + \left[\frac{d + v_s T_0}{v} \right] \quad \text{--- (8.19)}$$

Similarly at time pT_0 , the source emits its $(p+1)^{\text{th}}$ crest (where, p is an integer, $p = 1, 2, 3, \dots$). It reaches the listener at time

$$t_{p+1} = pT_0 + \left[\frac{d + pv_s T_0}{v} \right]$$

Hence the listener's detector counts p crests in the time interval

$$t_{p+1} - t_1 = pT_0 + \left[\frac{d + pv_s T_0}{v} \right] - \frac{d}{v}$$

Hence the period of wave as recorded by the listener is

$$T = \frac{(t_{p+1} - t_1)}{p} \quad \text{or}$$

$$T = \frac{\left[pT_0 + \frac{d + pv_s T_0}{v} - \frac{d}{v} \right]}{p}$$

$$T = T_0 + \frac{v_s T_0}{v}$$

$$T = T_0 \left[1 + \frac{v_s}{v} \right]$$

$$T = T_0 \left[\frac{v + v_s}{v} \right]$$

$$\therefore \frac{1}{T} = \frac{1}{T_0} \left[\frac{v}{v + v_s} \right]$$

$$\therefore n = n_0 \left[\frac{v}{v + v_s} \right] \quad \text{--- (8.20)}$$

where n is the frequency recorded by the listener and n_0 is the frequency emitted by the source.

If source of sound is moving towards the listener with speed v_s (called relative approach), the second term from Eq. (8.18) onwards, will be negative (or will be subtracted).

Thus, in this case,

$$n = n_0 \left[\frac{v}{v - v_s} \right] \quad \text{--- (8.21)}$$

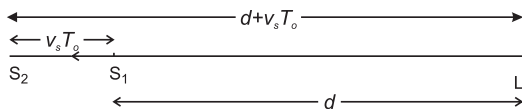


Fig. 8.2 (a): Doppler effect detected when the source is moving and listener is at rest in the medium.

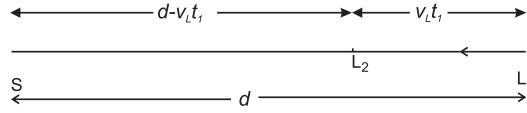


Fig. 8.2 (b): Doppler effect detected when the listener is moving and source is at rest in the medium.

8.9.2 Listener Approaching a Stationary Source with Velocity v_L :

Consider a listener approaching with velocity v_L towards a stationary source S as shown in Fig. 8.2 (b). Let the first wave be emitted by the source at $t = 0$, when the listener was at L_1 at an initial distance d from the source. Let t_1 be the instant when the listener receives this (wave), his position being L_2 . During time t_1 , the listener travels distance $v_L t_1$ towards the stationary source. In this time, the sound wave travels distance $(d - v_L t_1)$ with speed v .

$$\therefore t_1 = \frac{d - v_L t_1}{v} \quad \therefore t_1 = \frac{d}{v + v_L}$$

Second wave is emitted by the source at $t = T_0$ = the time period of the waves emitted by the source. Let t_2 be the instant when the listener receives second wave. During time t_2 , the distance travelled by the listener is $v_L t_2$. Thus, the distance to be travelled by the sound to reach the listener is then $d - v_L t_2$.

$\therefore$ Sound (second wave) travels this distance with speed v in time $= \frac{d - v_L t_2}{v}$

However, this time should be counted after T_0 , as the second wave was emitted at $t = T_0$.

$$\therefore t_2 = T_0 + \frac{d - v_L t_2}{v} \quad \therefore t_2 = \frac{vT_0 + d}{v + v_L}$$

$$\text{Similarly, } t_3 = 2T_0 + \frac{d - v_L t_3}{v} \quad \therefore t_3 = \frac{2vT_0 + d}{v + v_L}$$

Extending this argument to $(p+1)^{\text{th}}$ wave, we can write,

$$t_{p+1} = pT_0 + \frac{d - v_L t_{p+1}}{v} \quad \therefore t_{p+1} = \frac{pvT_0 + d}{v + v_L}$$

Time duration between instances of receiving successive waves is the *observed* or *recorded* period T .

$$\therefore pT = t_{p+1} - t_1 = \frac{pvT_0 + d}{v + v_L} - \frac{d}{v + v_L} = \frac{pvT_0}{v + v_L}$$

$$\therefore T = T_0 \left(\frac{v}{v + v_L} \right) \quad \text{--- (8.22)}$$

$$\therefore \frac{1}{T} = \frac{1}{T_0} \left(\frac{v + v_L}{v} \right)$$

$$\therefore n = n_0 \left(\frac{v + v_L}{v} \right) \quad \text{--- (8.23)}$$

8.9.3 Both Source and Listener are Moving:

In general when both the source and listener are in motion, we can write the observed frequency

$$n = n_0 \left[\frac{v \pm v_L}{v \mp v_s} \right] \quad \text{--- (8.24)}$$

Where the upper signs (in both numerator and denominator) should be chosen during relative approach while lower signs should be chosen during relative recede. It must be remembered that 'when you are deciding the sign for any one of these, the other should be considered to be at rest'.

Illustration:

Consider an observer or listener and a source moving with respective velocities v_L and v_s along the same direction. In this case, listener is approaching the source with v_L (irrespective of whether source is moving or not). Thus, the upper, i.e., positive sign, should be chosen for numerator. However, the source is moving with v_s away from the listener irrespective of listener's motion. Thus the lower sign in the denominator which is positive has to be chosen.

$$\therefore n = n_0 \left[\frac{v + v_L}{v + v_s} \right] \quad \text{--- (8.25)}$$

Case (I) If $|v_L| = |v_s|$, $n = n_0$. Thus there is no Doppler shift as there is no relative motion, even if both are moving.

Case (II) If $|v_L| > |v_s|$, numerator will be greater, $n > n_0$. This is because there is relative approach as the listener approaches the source faster and the source is receding at a slower rate.

Case (III) If $|v_L| < |v_s|$, $n < n_0$ as now there is relative recede (source recedes faster, listener approaches slowly).

8.9.4 Common Properties between Doppler Effect of Sound and Light:

- Wherever there is relative motion between listener (or observer) and source (of sound or light waves), the recorded frequency is different than the emitted frequency.
- Recorded frequency is higher (than emitted frequency), if there is relative approach.
- Recorded frequency is lower, if there is relative recede.
- If v_L or v_s are much smaller than wave speed (speed of sound or light) we can use v_r as relative velocity. In this case, using Eq. (8.24)

$$\frac{\Delta n}{n} \simeq \frac{v_r}{v} \simeq \frac{\Delta \lambda}{\lambda} \quad \text{--- (8.26)}$$

where Δn is Doppler shift or change in the recorded frequency, i.e., $|n - n_0|$ and $\Delta \lambda$ is the recorded change in wavelength.

$$\therefore \frac{|n - n_0|}{n} \simeq \frac{v_r}{v}$$

$$\therefore n = n_0 \left(1 \pm \frac{v_r}{v} \right) \quad \text{--- (8.27)}$$

Once again upper sign is to be used during relative approach while lower sign is to be used during relative recede.

- If velocities of source and observer (listener) are not along the same line their respective components along the line joining them should be chosen for longitudinal Doppler effect and the same mathematical treatment is applicable.

8.9.5 Major Differences between Doppler Effects of Sound and Light:

- As the speed of light is absolute, only relative velocity between the observer and the source matters, i.e., who is in motion is not relevant.
- Classical and relativistic Doppler effects are different in the case of light, while in case of sound, it is only classical.

- C) For obtaining exact Doppler shift for sound waves, it is absolutely important to know who is in motion.
- D) If wind is present, its velocity alters the speed of sound and hence affects the Doppler shift. In this case, component of the wind velocity (v_w) is chosen along the line joining source and observer. This is to be algebraically added with the velocity of sound. Hence ' v ' is to be replaced by $(v \pm v_w)$ in all the above expressions. Positive sign to be used if v and v_w are along the same direction (remember that v is always positive and always from source to listener). Negative sign is to be used if v and v_w are oppositely directed.

Example 8.6: A rocket is moving at a speed of 220 m/s towards a stationary target. It emits a wave of frequency 1200 Hz. Some of the sound reaching the target gets reflected back to the rocket as an echo. Calculate (1) The frequency of sound detected by the target and (2) The frequency of echo detected by rocket (velocity of sound = 330 m/s.)

Solution: Given, target stationary, i.e.,

$$v_L = 0, v_s = 220 \text{ m/s}, v = 330 \text{ m/s}$$

$$n_0 = 1200 \text{ Hz}$$

To find the frequency of sound detected by the target we have to use Eq. (8.25)

$$n = n_0 \left[\frac{v}{v - v_s} \right]$$

$$n = 1200 \left[\frac{330}{330 - 220} \right]$$

$$n = 3600 \text{ Hz}$$

The frequency of sound detected by the target = 3600 Hz.

When echo is heard by rocket's detector, target is considered as source

$$\therefore v_s = 0$$

The frequency of sound emitted by the source (i.e. target) is $n_0 = 3600 \text{ Hz}$, and the frequency detected by rocket is n' . Now listener is approaching the source and so we have to use.

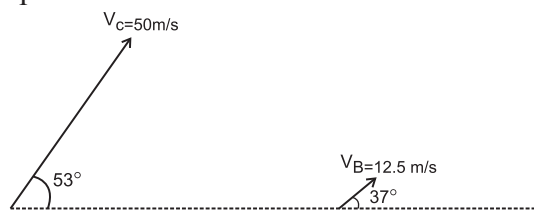
$$n' = n \left[\frac{v + v_L}{v} \right]$$

$$n = 3600 \left[\frac{330 + 220}{330} \right]$$

$$n = 6000 \text{ Hz}$$

The frequency of echo detected by rocket = 6000 Hz

Example 8.7: A bat, flying at velocity $V_B = 12.5 \text{ m/s}$, is followed by a car running at velocity $V_C = 50 \text{ m/s}$. Actual directions of the velocities of the car and the bat are as shown in the figure below, both being in the same horizontal plane (the plane of the figure). To detect the car, the bat radiates ultrasonic waves of frequency 36 kHz. Speed of sound at surrounding temperature is 350 m/s.



There is an ultrasonic frequency detector fitted in the car. Calculate the frequency recorded by this detector.

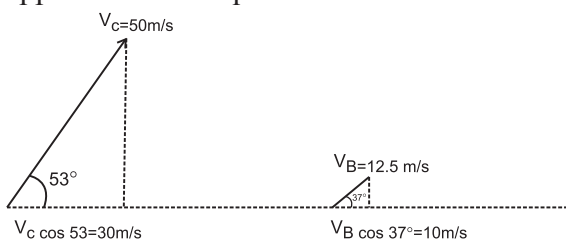
The ultrasonic waves radiated by the bat are reflected by the car. The bat detects these waves and from the detected frequency, it knows about the speed of the car. Calculate the frequency of the reflected waves as detected by the bat. ($\sin 37^\circ = \cos 53^\circ \approx 0.6$, $\sin 53^\circ = \cos 37^\circ \approx 0.8$)

Solution: As shown in the figure below, the components of velocities of the bat and the car, along the line joining them, are

$$V_C \cos 53^\circ \approx 50 \times 0.6 = 30 \text{ m s}^{-1} \text{ and}$$

$$V_B \cos 37^\circ \approx 12.5 \times 0.8 = 10 \text{ m s}^{-1}.$$

These should be used while calculating the doppler shifted frequencies.



Doppler shifted frequency, $n = n_0 \left(\frac{v \pm v_L}{v \mp v_s} \right)$; upper signs to be used during approach, lower signs during recede.

Part I: Frequency radiated by the bat $n_0 = 36 \times 10^3$ Hz, Frequency detected by the detector in the car $= n = ?$

In this case, bat is the source which is moving away from the car (receding) while the detector in the car is the listener, who is approaching the source (bat). $v_s = V_B \cos 37^\circ = 10$ m/s and

$$v_L = V_C \cos 53^\circ = 30 \text{ m/s}$$

The source (bat) is receding, while the listener

(car) is approaching $\therefore n = n_0 \left(\frac{v + v_L}{v + v_s} \right)$

$$\therefore n = 36 \times 10^3 \left(\frac{350 + 30}{350 + 10} \right)$$

$$= 38 \times 10^3 \text{ Hz} = 38 \text{ kHz}$$

Part II: Reflected frequency, as detected by the bat: Frequency reflected by the car is the Doppler shifted frequency as detected at the car. Thus, this time, the car is the source with

emitted frequency $n_0 = 38 \times 10^3$ Hz, $n = ?$

Car, the source, is approaching the listener (bat).

$$\text{Thus, } v_s = v_C \cos 53^\circ = 30 \text{ m/s}$$

$$\text{Thus, } v_L = v_B \cos 37^\circ = 10 \text{ m/s}$$

Now bat-the listener is receding while car the

source is approaching $\therefore n = n_0 \left(\frac{v - v_L}{v - v_s} \right)$

$$\therefore n = 38 \times 10^3 \left(\frac{350 - 10}{350 - 30} \right)$$

$$= 38 \times 10^3 \times \frac{34}{32}$$

$$= 40.375 \times 10^3 \text{ Hz}$$

$$= 40.375 \text{ kHz}$$



Internet my friend

<https://hyperphysics.phys-astr.gsu.edu/hbase/hframe.html>



Exercises

1. Choose the correct alternatives

- A sound carried by air from a sitar to a listener is a wave of following type.
 - Longitudinal stationary
 - Transverse progressive
 - Transverse stationary
 - Longitudinal progressive
- When sound waves travel from air to water, which of these remains constant?
 - Velocity
 - Frequency
 - Wavelength
 - All of above
- The Laplace's correction in the expression for velocity of sound given by Newton is needed because sound waves
 - are longitudinal
 - propagate isothermally
 - propagate adiabatically
 - are of long wavelength
- Speed of sound is maximum in
 - air
 - water
 - vacuum
 - solid
- The walls of the hall built for music

concerns should

- amplify sound
- reflect sound
- transmit sound
- absorb sound

2. Answer briefly.

- Wave motion is doubly periodic. Explain.
- What is Doppler effect?
- Describe a transverse wave.
- Define a longitudinal wave.
- State Newton's formula for velocity of sound.
- What is the effect of pressure on velocity of sound?
- What is the effect of humidity of air on velocity of sound?
- What do you mean by an echo?
- State any two applications of acoustics.
- Define amplitude and wavelength of a wave.
- Draw a wave and indicate points which are (i) in phase (ii) out of phase (iii) have a phase difference of $\pi/2$.

- xii) Define the relation between velocity, wavelength and frequency of wave.
- xiii) State and explain principle of superposition of waves.
- xiv) State the expression for apparent frequency when source of sound and listener are
 - i) moving towards each other
 - ii) moving away from each other
- xv) State the expression for apparent frequency when source is stationary and listener is
 - 1) moving towards the source
 - 2) moving away from the source
- xvi) State the expression for apparent frequency when listener is stationary and source is.
 - i) moving towards the listener
 - ii) moving away from the listener
- xvii) Explain what is meant by phase of a wave.
- xviii) Define progressive wave. State any four properties.
- xix) Distinguish between transverse waves and longitudinal waves.
- xx) Explain Newton's formula for velocity of sound. What is its limitation?

3. Solve the following problems.

- i) A certain sound wave in air has a speed 340 m/s and wavelength 1.7 m for this wave, calculate
 - a) the frequency
 - b) the period.

[Ans a) 200 Hz, b) 0.005s]
- ii) A tuning fork of frequency 170 Hz produces sound waves of wavelength 2 m. Calculate speed of sound.

[Ans: 340 m/s]
- iii) An echo-sounder in a fishing boat receives an echo from a shoal of fish 0.45 s after it was sent. If the speed of sound in water is 1500 m/s, how deep is the shoal?

[Ans : 337.5 m]
- iv) A girl stands 170 m away from a high wall and claps her hands at a steady rate so that each clap coincides with the echo of the one before.
 - a) If she makes 60 claps in 1 minute,

what value should be the speed of sound in air?

b) Now, she moves to another location and finds that she should now make 45 claps in 1 minute to coincide with successive echoes. Calculate her distance for the new position from the wall.

[Ans: a) 340 m/s b) 255 m]

- v) Sound wave A has period 0.015 s, sound wave B has period 0.025. Which sound has greater frequency?

[Ans : A]

- vii) At what temperature will the speed of sound in air be 1.75 times its speed at N.T.P?

[Ans: 836.06 K = 563.06 °C]

- viii) A man standing between 2 parallel cliffs fires a gun. He hears two echoes one after 3 seconds and other after 5 seconds. The separation between the two cliffs is 1360 m, what is the speed of sound?

[Ans: 340 m/s]

- ix) If the velocity of sound in air at a given place on two different days of a given week are in the ratio of 1:1.1. Assuming the temperatures on the two days to be same what quantitative conclusion can you draw about the condition on the two days?

[Ans: Air is moist on one day

and $\rho_{\text{dry}} = 1.1^2 \rho_{\text{dry}} = 1.21 \rho_{\text{moist}}$]

- x) A police car travels towards a stationary observer at a speed of 15 m/s. The siren on the car emits a sound of frequency 250 Hz. Calculate the recorded frequency. The speed of sound is 340 m/s.

[Ans : 261.54 Hz]

- xi) The sound emitted from the siren of an ambulance has frequency of 1500 Hz. The speed of sound is 340 m/s. Calculate the difference in frequencies heard by a stationary observer if the ambulance initially travels towards and then away from the observer at a speed of 30 m/s.

[Ans : 266.79 Hz]



Can you recall?

1. What are laws of reflection and refraction?
2. What is dispersion of light?
3. What is refractive index?
4. What is total internal reflection?
5. How does light refract at a curved surface?
6. How does a rainbow form?

9.1 Introduction:

“See it to believe it” is a popular saying. In order to see, we need light. What exactly is light and how are we able to see anything? We will explore it in this and next standard. We know that acoustics is the term used for science of sound. Similarly, optics is the term used for science of light. There is a difference in the nature of sound waves and light waves which you have seen in chapter 8 and will learn in chapter 13.

9.2 Nature of light:

Earlier, light was considered to be that form of radiant energy which makes objects visible due to stimulation of retina of the eye. It is a form of energy that propagates in the presence or absence of a medium, which we now call waves. At the beginning of the 20th century, it was proved that these are electromagnetic (EM) waves. Later, using quantum theory, particle nature of light was established. According to this, photons are energy carrier particles. By an experiment using countable number of photons, it is now an established fact that light possesses dual nature. In simple words we can say that light consists of energy carrier photons guided by the rules of EM waves. In vacuum, these waves (or photons) travel with a speed of

$c = 299792458 \text{ m s}^{-1}$ According to Einstein’s special theory of relativity, this is the maximum possible speed for any object. For practical purposes we write it as $c = 3 \times 10^8 \text{ m s}^{-1}$.

Commonly observed phenomena concerning light can be broadly split into three categories.

- (I) **Ray optics or geometrical optics:** A particular direction of propagation of energy from a source of light is called a ray of light. We use ray optics for understanding phenomena like reflection, refraction, double refraction, total internal reflection, etc.
- (II) **Wave optics or physical optics:** For explaining phenomena like interference, diffraction, polarization, Doppler effect, etc., we consider light energy to be in the form of EM waves. Wave theory will be further discussed in XIIth standard.
- (III) **Particle nature of light:** Phenomena like photoelectric effect, emission of spectral lines, Compton effect, etc. cannot be explained by using classical wave theory. These involve the interaction of light with matter. For such phenomena we have to use quantum nature of light. Quantum nature of light will be discussed in XIIth standard.

9.3 Ray optics or geometrical optics:

In geometrical optics, we mainly study image formation by mirrors, lenses and prisms. It is based on four fundamental laws/ principles which you have learnt in earlier classes.

- (i) Light travels in a straight line in a homogeneous and isotropic medium. Homogeneous means that the properties of the medium are same every where in the medium and isotropic means that the

In a material medium, the speed of EM waves is given by $c = \sqrt{\frac{1}{\epsilon \mu}}$,

where permittivity ϵ and permeability μ are constants which depend on the electric and magnetic properties of the medium.

The ratio $n = \frac{c}{v}$ is called the absolute refractive index and is the property of the medium.

properties are the same in all directions.

- (ii) Two or more rays can intersect at a point without affecting their paths beyond that point.

(iii) **Laws of reflection:**

- (a) Reflected ray lies in the plane formed by incident ray and the normal drawn at the point of incidence; and the two rays are on either side of the normal.

- (b) Angles of incidence and reflection are equal.

- (iv) **Laws of refraction:** These apply at the boundary between two media

- (a) Refracted ray lies in the plane formed by incident ray and the normal drawn at the point of incidence; and the two rays are on either side of the normal.

- (b) Angle of incidence (θ_1 in a medium of refractive index n_1) and angle of refraction (θ_2 in medium of refractive index n_2) are related by Snell's law, given by

$$(n_1)\sin\theta_1 = (n_2)\sin\theta_2.$$



Do you know ?

Interestingly, all the four laws stated above can be derived from a single principle called Fermat's (pronounced "Ferma") principle. It says that "While travelling from one point to another by one or more reflections or refractions, a ray of light always chooses the path of least time".

Ideally it is the path of extreme time, i.e., path of minimum or maximum time. We strongly recommend you to go through a suitable reference book that will give you the proof of $i = r$ during reflection and Snell's law during refraction using Fermat's principle.

Example 9.1: Thickness of the glass of a spectacle is 2 mm and refractive index of its glass is 1.5. Calculate time taken by light to cross this thickness. Express your answer with the most convenient prefix attached to the unit 'second'.

Solution:

Speed of light in vacuum, $c = 3 \times 10^8 \text{ m/s}$

$$n_{\text{glass}} = 1.5$$

$\therefore$ Speed of light in glass =

$$\frac{c}{n_{\text{glass}}} = \frac{3 \times 10^8}{1.5} = 2 \times 10^8 \text{ m/s}$$

Distance to be travelled by light in glass,

$$s = 2 \text{ mm} = 2 \times 10^{-3} \text{ m}$$

$\therefore$ Time t required by light to travel this distance,

$$t = \frac{s}{v_{\text{glass}}} = \frac{2 \times 10^{-3}}{2 \times 10^8} = 10^{-11} \text{ s}$$

Most convenient prefix to express this small time is pico (p) = 10^{-12}

$$\therefore t = 10 \times 10^{-12} = 10 \text{ ps}$$

9.3.1 Cartesian sign convention:

While using geometrical optics it is necessary to use some sign convention. The relation between only the numerical values of u , v and f for a spherical mirror (or for a lens) will be different for different positions of the object and the type of mirror. Here u and v are the distances of object and image respectively from the optical center, and f is the focal length. Properly used suitable sign convention enables us to use the same formula for all different particular cases. Thus, while deriving a formula and also while using the formula it is necessary to use the same sign convention. Most convenient sign convention is Cartesian sign convention as it is analogous to coordinate geometry. According to this sign convention, (Fig. 9.1):

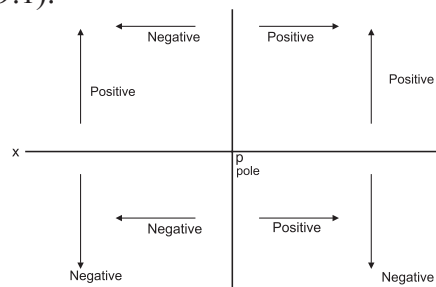


Fig. 9.1 Cartesian sign convention.

- i) All distances are measured from the optical center or pole. For most of the optical objects such as spherical mirrors, thin lenses, etc., the optical centers coincides with their geometrical centers.

- ii) Figures should be drawn in such a way that the incident rays travel from left to right. A diverging beam of incident rays corresponds to a real point object (Fig. 9.2 (a)), a converging beam of incident rays corresponds to a virtual object (Fig. 9.2 (b)) and a parallel beam corresponds to an object at infinity. Thus, a real object should be shown to the left of pole (Fig. 9.2 (a)) and virtual object or image to the right of pole. (Fig. 9.2 (b))

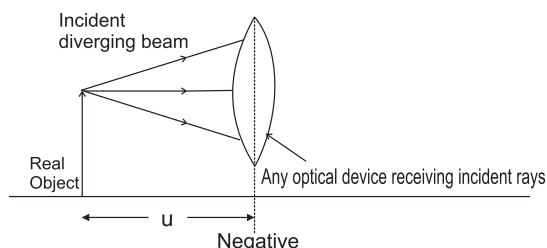


Fig. 9.2: (a) Diverging beam from a real object

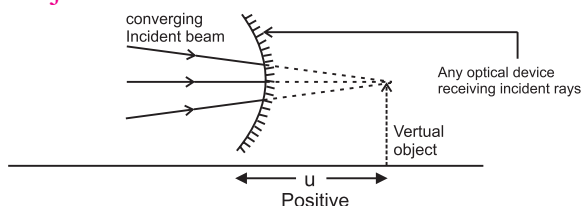


Fig. 9.2: (b) Converging beam towards a virtual object.

- iii) x-axis can be conveniently chosen as the principal axis with origin at the pole.
 iv) Distances to the left of the pole are negative and those to the right of the pole are positive.
 v) Distances above the principal axis (x-axis) are positive while those below it are negative.

Unless specially mentioned, we shall always consider objects to be real for further discussion.

9.4 Reflection:

9.4.1 Reflection from a plane surface:

- a) If the object is in front of a plane reflecting surface, the image is virtual and laterally inverted. It is of the same size as that of the object and at the same distance as that of object but on the other side of the reflecting surface.

- b) If we are standing on the bank of a still water body and look for our image formed by water (or if we are standing on a plane mirror and look for our image formed by the mirror), the image is laterally reversed, of the same size and on the other side.
 c) If an object is kept between two plane mirrors inclined at an angle θ (like in a kaleidoscope), a number of images are formed due to multiple reflections from both the mirrors. Exact number of images depends upon the angle between the mirrors and where exactly the object is kept. It can be obtained as follows (Table 9.1):

$$\text{Calculate } n = \frac{360}{\theta}$$

Let N be the number of images seen.

- (I) If n is an even integer, $N = (n - 1)$, irrespective of where the object is.
 (II) If n is an odd integer and object is exactly on the angle bisector, $N = (n - 1)$.
 (III) If n is an odd integer and object is off the angle bisector, $N = n$
 (IV) If n is not an integer, $N = m$, where m is integral part of n .

Table 9.1

Angle θ°	$n = \frac{360}{\theta}$	Position of the object	N
120	3	On angle bisector	2
120	3	Off angle bisector	3
110	3.28	Anywhere	3
90	4	Anywhere	3
80	4.5	Anywhere	4
72	5	On angle bisector	4
72	5	Off angle bisector	5
60	6	Anywhere	5
50	7.2	Anywhere	7

Example 9.2: A small object is kept symmetrically between two plane mirrors inclined at 38° . This angle is now gradually increased to 41° , the object being symmetrical all the time. Determine the number of images visible during the process.

Solution: According to the convention used in the table above,

$$\theta = 38^\circ \therefore n = \frac{360}{38} = 9.47$$

$\therefore N = 9$. This is valid till the angle is 40° as the object is kept symmetrically

Beyond 40° , $n < 9$ and it decreases upto

$$\frac{360}{41} = 8.78$$

Hence now onwards there will be 8 images till 41° .

9.4.2 Reflection from curved mirrors:

In order to focus a parallel or divergent beam by reflection, we need curved mirrors. You might have noticed that reflecting mirrors for a torch or headlights, rear view mirrors of vehicles are not plane but concave or convex. Mirrors for a search light are parabolic. We shall restrict ourselves to spherical mirrors only which can be studied using simple mathematics. Such mirrors are parts of a sphere polished from outside (convex) or from inside (concave).

Radius of the sphere of which a mirror is a part is called as radius of curvature (R) of the mirror. Only for spherical mirrors, half of radius of curvature is focal length of the mirror $\left(f = \frac{R}{2}\right)$. For a concave mirror it is the distance at which parallel incident rays converge. For a convex mirror, it is the distance from where parallel rays appear to be diverging after reflection. According to sign convention, the incident rays are from left to right and they should face the polished surface of the mirror. Thus, focal length of a convex mirror is positive (Fig 9.3 (a)) while that of a concave mirror is negative (Fig. 9.3 (b)).

Relation between f , u and v :

For a point object or for a small finite object, the focal length of a *small* spherical

(concave or convex) mirror is related to object distance and image distance as

$$\frac{1}{f} = \frac{1}{v} + \frac{1}{u} \quad \text{--- (9.1)}$$

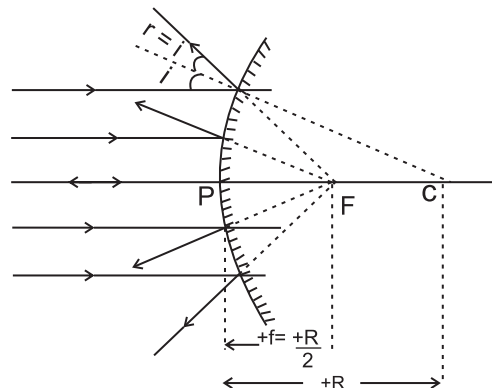


Fig. 9.3 (a): Parallel rays incident from left appear to be diverging from F, lying on the positive side of origin (pole).

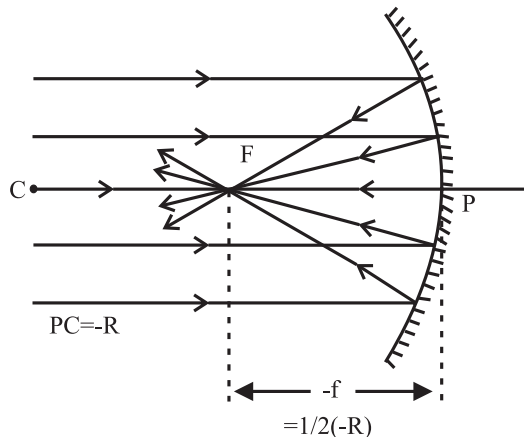


Fig. 9.3 (b): Parallel rays incident from left appear to converge at F, lying on the negative side of origin (pole).

By a *small* mirror we mean its aperture (diameter) is much smaller (at least one tenth) than the values of u , v and f .

Focal power: Converging or diverging ability of a lens or of a mirror is defined as its focal

power. It is measured as $P = \frac{1}{f}$.

In SI units, it is measured as diopter.
 $\therefore 1 \text{ dioptre (D)} = 1 \text{ m}^{-1}$

Lateral magnification: Ratio of linear size of an image to that of the object, measured perpendicular to the principal axis, is defined as

the lateral magnification $m = \frac{v}{u}$

For any position of the object, a convex mirror

always forms virtual, erect and diminished image, $m < 1$. In the case of a concave mirror it depends upon the position of the object. Following Table 9.2 will help you refresh your knowledge.

Table 9.2

Concave mirror (f negative)			
Position of object	Position of image	Real(R) or Virtual (V)	Lateral magnification
$u = \infty$	$v = f$	R	$m = 0$
$u > 2f$	$2f > v > f$	R	$m < 1$
$u = 2f$	$v = 2f$	R	$m = 1$
$2f > u > f$	$v > 2f$	R	$m > 1$
$u = f$	$v = \infty$	R	$m = \infty$
$u < f$	$ v > u $	V	$m > 1$

Example 9.3: A thin pencil of length 20 cm is kept along the principal axis of a concave mirror of curvature 30 cm. Nearest end of the pencil is 20 cm from the pole of the mirror. What will be the size of image of the pencil?

Solution: $R = 30$ cm

$$f = R/2 = -15 \text{ cm} \dots (\text{Concave mirror})$$

$$\frac{1}{f} = \frac{1}{v} + \frac{1}{u}$$

For nearest end, $u = u_1 = -20$ cm. Let the image distance be v_1

$$\therefore \frac{1}{-15} = \frac{1}{v_1} + \frac{1}{-20} \quad \therefore v_1 = -60 \text{ cm}$$

Nearest end is at 20 cm and pencil itself is 20 cm long. Hence farthest end is $20 + 20 = 40$ cm $= -u_2$

Let the image distance be v_2

$$\therefore \frac{1}{-15} = \frac{1}{v_2} + \frac{1}{-40} \quad \therefore v_2 = -24 \text{ cm}$$

$\therefore$ Length of the image $= 60 - 24 = 36$ cm.

Defects or aberration of images: The theory of image formation by mirrors or lenses, and the formulae that we have used such as

$$f = \frac{R}{2} \text{ or } \frac{1}{f} = \frac{1}{v} + \frac{1}{u} \text{ etc.}$$

are based on the following assumptions: (i) Objects and images are situated close to the principal axis.

(ii) Rays diverging from the objects are confined

to a cone of very small angle.

(iii) If there is a parallel beam of rays, it is paraxial, i.e., parallel and close to the principal axis.

However, in reality, these assumptions do not always hold good. This results into distorted or defective image. Commonly occurring defects are spherical aberration, coma, astigmatism, curvature, distortion. Except spherical aberration, all the other arise due to beams of rays inclined to principal axis. These are not discussed here.

Spherical aberration: As mentioned earlier, the relation $\left(f = \frac{R}{2}\right)$ giving a single

point focus is applicable only for small aperture spherical mirrors and for paraxial rays. In reality, when the rays are farther from the principal axis, the focus gradually shifts towards pole (Fig. 9.4). This phenomenon (defect) arises due to spherical shape of the reflecting surface, hence called as spherical aberration. It results into a unsharp (fuzzy) image with unclear boundaries.

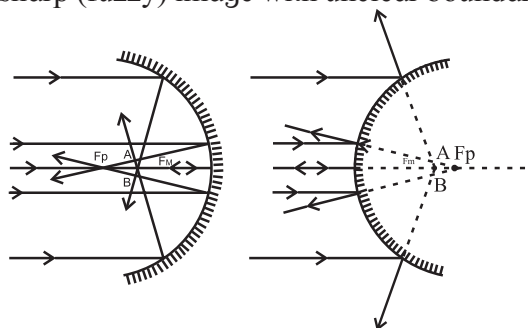


Fig. 9.4: Spherical aberration for curved mirrors.

The distance between F_M and F_P (Fig. 9.4) is measured as the longitudinal spherical aberration. If there is no spherical aberration, we get a single point image on a screen placed perpendicular to the principal axis at that location, for a beam of incident rays parallel to the axis. In the presence of spherical aberration, no such point is possible at any position of the screen and the image is always a circle. At a particular location of the screen, the diameter of this circle is minimum. This is called the circle of least confusion. In the figures it is across AB. Radius of this circle is transverse spherical aberration.

In the case of curved mirrors, this defect can be completely eliminated by using a parabolic mirror. Hence surfaces of mirrors used in a search light, torch, headlight of a car, telescopes, etc., are parabolic and not spherical.



Do you know ?

Why does a parabolic mirror not have spherical aberration?

Parabola is a geometrical shape drawn in such a way that every point on it is equidistant from a straight line and from a point. Figure 9.5 shows a parabola. Points A, B, C, ... on it are equidistant from line RS (called directrix) and point F (called focus). Hence $A'A = AF$, $B'B = BF$, $C'C = CF$,

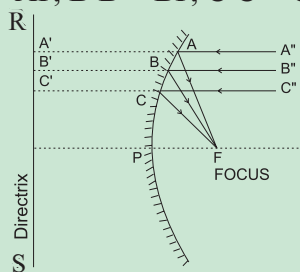


Fig. 9.5: Single focus for parabolic mirror.

If rays of equal optical path converge at a point, that point is the location of real image corresponding to that beam of rays.

Paths $A''A'$, $B''B'$, $C''C'$, etc., are equal paths in the absence of mirror. If the parabola ABC... is a mirror then the respective optical paths will be $A''AF$, $B''BF$, $C''CF$, ... and from the definition of parabola, these are also equal. Thus, F is the single point focus for entire beam parallel to the axis with NO spherical aberration.

9.5 Refraction:

Being an EM wave, the properties of light (speed, wavelength, direction of propagation, etc.) depend upon the medium through which it is traveling. If a ray of light comes to an interface between two media and enters into another medium of different refractive index, it changes itself suitable to that medium. This phenomenon is defined as refraction of light. The extent to which these properties change is decided by the index of refraction, ' n '.

Absolute refractive index:

Absolute refractive index of a medium is defined as the ratio of speed of light in vacuum to that in the given medium.

$n = \frac{c}{v}$ where c and v are respective speeds of light in vacuum and in the medium. As n is the ratio of same physical quantities, it is a unitless and dimensionless physical quantity.

For any material medium (including air) $n > 1$, i.e., light travels fastest in vacuum than in any material medium. Medium having greater value of n is called optically denser. An optically denser medium need not be physically denser, e.g., many oils are optically denser than water but water is physically denser than them.

Relative refractive index:

Refractive index of medium 2 with respect to medium 1 is defined as the ratio of speed of light v_1 in medium 1 to its speed v_2 in medium

$$2. \text{ Thus, } {}^1n_2 = \frac{n_2}{n_1} = \frac{v_1}{v_2}$$



Do you know ?

- Logic behind the convention 1n_2 : Letter n is the symbol for refractive index, n_2 corresponds to refractive index of medium 2 and 1n_2 indicates that it is with respect to medium 1. In this case, light travels from medium 1 to 2 so we need to discuss medium 2 in context to medium 1.
- Dictionary meaning of the word refract is to change the path. However, in context of Physics, we should be more specific. We use the word *deviate* for changing the path. During refraction at normal incidence, there is no change in path. Thus, there is *refraction* but no *deviation*. Deviation is associated with refraction *only* during oblique incidence. Deviation or changing the path or bending is associated with many phenomena such as reflection, diffraction, scattering, gravitational bending due to a massive object, etc.

Illustrations of refraction: 1) When seen from outside, the bottom of a water body appears to be raised. This is due to refraction at the plane surface of water. In this case,

$$n_{\text{water}} \cong \frac{\text{Real depth}}{\text{apparent depth}}$$

This relation holds good for a plane parallel transparent slab also as shown below.

Figure 9.6 shows a plane parallel slab of a transparent medium of refractive index n . A point object O at real depth R appears to be at I at apparent depth A , when seen from outside (air). Incident rays OA (traveling undeviated) and OB (deviating along BC) are used to locate the image.

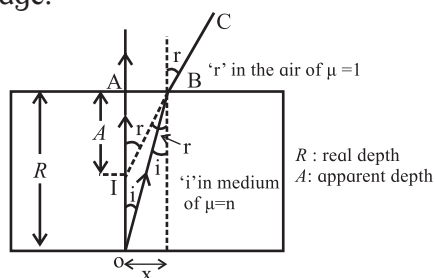


Fig. 9.6: Real and apparent depth.

By considering i and r to be small, we can write,

$$\tan(r) = \frac{x}{A} \cong \sin(r) \text{ and } \tan(i) = \frac{x}{R} \cong \sin(i)$$

$$\therefore n = \frac{\sin(r)}{\sin(i)} \cong \frac{\left(\frac{x}{A}\right)}{\left(\frac{x}{R}\right)} = \frac{R}{A} = \frac{\text{Real depth}}{\text{Apparent depth}}$$

2) A stick or pencil kept obliquely in a glass containing water appears broken as its part in water appears to be raised.

Small angle approximation: For small angles, expressed in radian, $\sin \theta \cong \theta \cong \tan \theta$. For example, for

$$\theta = 30^\circ = \left(\frac{\pi}{6}\right)^c = 0.5236^c,$$

we have $\sin \theta = 0.5$

In this case the error is $0.5236 - 0.5 = 0.0236$ in 0.5, which is 4.72 %.

For practical purposes we consider angles less than 10° where the error in using $\sin \theta \cong \theta$ is less than 0.51 %. (Even for 60° , it is still 15.7 %) It is left to you to verify that this is almost equally valid for $\tan \theta$ till 20° only.

Example 9. 4: A crane flying 6 m above a still, clear water lake sees a fish underwater. For the crane, the fish appears to be 6 cm below the water surface. How much deep should the crane immerse its beak to pick that fish?

For the fish, how much above the water surface does the crane appear? Refractive index of water = $4/3$.

Solution: For crane, apparent depth of the fish is 6 cm and real depth is to be determined.

For fish, real depth (height, in this case) of the crane is 6 m and apparent depth (height) is to be determined.

$$n = \frac{R}{A} = \frac{\text{Real depth}}{\text{Apparent depth}}$$

For crane, it is water with respect to air as real depth is in water and apparent depth is as seen from air

$$\therefore n = \frac{4}{3} = \frac{R}{A} = \frac{R}{6} \quad \therefore R = 8 \text{ cm}$$

For fish, it is air with respect to water as the real height is in air and seen from water.

$$\therefore n = \frac{3}{4} = \frac{R}{A} = \frac{6}{A} \quad \therefore A = 8 \text{ m}$$

9.6 Total internal reflection:

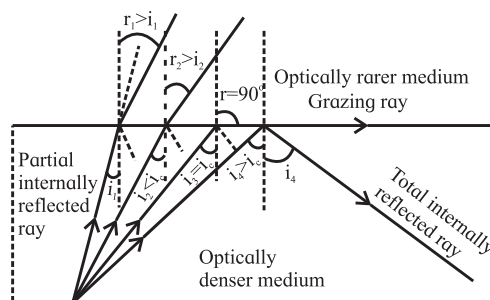


Fig. 9.7: Total internal reflection.

Figure 9.7 shows refraction of light emerging from a denser medium into a rarer medium for various angles of incidence. The angles of refraction in the rarer medium are larger than the corresponding angles of incidence. At a particular angle of incidence i_c in the denser medium, the corresponding angle of refraction in the rarer medium is 90° . For angles of incidence greater than i_c , the angle of refraction become larger than 90° and the ray does not enter into rarer medium at all but is reflected totally into the denser medium. This is

called total internal reflection. In general, there is always partial reflection and partial refraction at the interface. During total internal reflection TIR, it is total reflection and no refraction. The corresponding angle of incidence in the denser medium is greater than or equal to the critical angle.

Critical angle for a pair of refracting media can be defined as that angle of incidence in the denser medium for which the angle of refraction in the rarer medium is 90° .



Do you know ?

In Physics the word *critical* is used when certain phenomena are not applicable or more than one phenomenon are applicable. Some examples are as follows.

- (i) In case of total internal reflection, the phenomenon of reversibility of light is not applicable at critical angle and refraction is possible only for angles of incidence in the denser medium smaller than the critical angle.
- (ii) At the *critical* temperature, a substance coexists into all the three states; solid, liquid and gas. At all the other temperatures, only two states are simultaneously possible.
- (iii) For liquids, streamline flow is possible till *critical* velocity is achieved. At critical velocity it can be either streamline or turbulent.

Let μ be the relative refractive index of denser medium with respect to the rarer. Applying Snell's law at the critical angle of incidence, i_c , we can write $\sin(i_c) = \frac{1}{\mu}$ as,

$$(\mu) \sin(i_c) = (1) \sin 90^\circ$$

For commonly used glasses of

$\mu = 1.5$, $i_c = 41^\circ 49' \cong 42^\circ$ and for water of

$\mu = \frac{4}{3}$, $i_c = 48^\circ 35'$ (Both, with respect to air)

9.6.1 Applications of total internal reflection:

(i) Optical fibre: Though little costly for initial set up, optic fibre communication is undoubtedly the most effective way of telecommunication by way of EM waves.

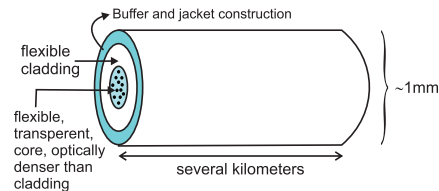


Fig. 9.8 (a): Optical fibre construction.

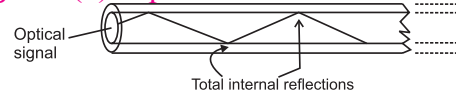


Fig. 9.8 (b): Optical fibre working.

An optical fibre essentially consists of an extremely thin (slightly thicker than a human hair), transparent, flexible core surrounded by optically rarer (smaller refractive index), flexible cover called cladding. This system is coated by a buffer and a jacket for protection. Entire thickness of the fibre is less than half a mm. (Fig. 9.8(a)). Number of such fibres may be packed together in an outer cover.

An optical signal (ray) entering the core suffers multiple total internal reflections (Fig. 9.8 (b)) and emerges after several kilometers with extremely low loss travelling with highest possible speed in that material ($\sim 2,00,000$ km/s for glass). Some of the advantages of optic fibre communication are listed below.

- (a) Broad bandwidth (frequency range): For TV signals, a single optical fibre can carry over 90000 channels (independent signals).
- (b) Immune to EM interference: Being electrically non-conductive, it is not able to pick up nearby EM signals.
- (c) Low attenuation loss: The loss is lower than 0.2 dB/km so that a single long cable can be used for several kilometers.
- (d) Electrical insulator: No issue with ground loops of metal wires or lightning.
- (e) Theft prevention: It does not use copper or other expensive material.
- (f) Security of information: Internal damage is most unlikely.

(ii) Prism binoculars: Binoculars using only two cylinders have a limitation of field of view as the distance between the two cylinders can't be greater than that between the two eyes. This limitation can be overcome

by using two right angled glass prisms ($i_c \sim 42^\circ$) used for total internal reflection as shown in the Fig. 9.9. Total internal reflections occur inside isosceles, right angled prisms.

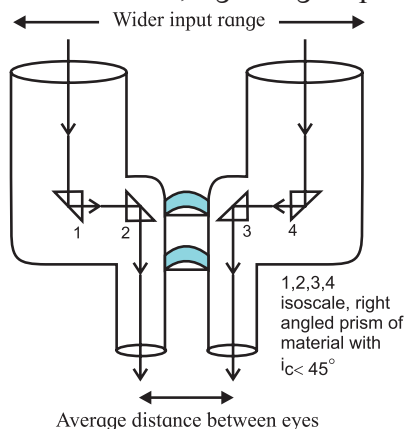


Fig. 9.9: Prism binoculars

(iii) **Periscope:** It is used to see the objects on the surface of a water body from inside water. The rays of light should be reflected twice through right angle. Reflections are similar to those in the binoculars (Fig 9.10) and total internal reflections occur inside isosceles, right angled prisms.

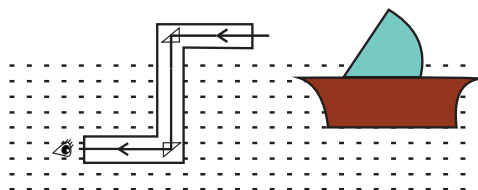
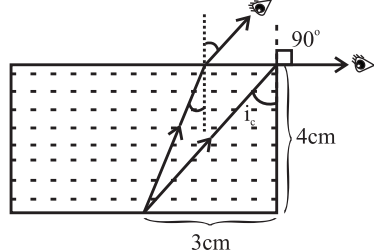


Fig. 9.10: Periscope.

Example 9.5: There is a tiny LED bulb at the center of the bottom of a cylindrical vessel of diameter 6 cm. Height of the vessel is 4 cm. The beaker is filled completely with an optically dense liquid. The bulb is visible from any inclined position but just visible if seen along the edge of the beaker. Determine refractive index of the liquid.

Solution: As seen from the accompanying figure, if the bulb is just visible from the edge, angle of incidence in the liquid (at the edge) must be the critical angle of incidence, i_c



From the dimensions given,

$$\tan(i_c) = \frac{3}{4} \therefore \sin(i_c) = \frac{3}{5} \therefore n_{\text{liquid}} = \frac{\sin 90^\circ}{\sin(i_c)} = \frac{5}{3}$$

9.7 Refraction at a spherical surface and lenses:

In the section 9.5 we saw that due to refraction, the bottom of a water body appears to be raised and $n_{\text{water}} = \frac{\text{Real depth}}{\text{apparent depth}}$.

However, this is valid only if we are dealing with refraction at a plane surface. In many cases such as liquid drops, lenses, ellipsoid paper weights, etc, curved surfaces are present and the formula mentioned above may not be true. In such cases we need to consider refraction at one or more spherical surfaces. This will involve parameters including the curvature such as radius of curvature, in addition to refractive indices.

Lenses: Commonly used lenses can be visualized to be consisting of intersection of two spheres of radii of curvature R_1 and R_2 or of one sphere and a plane surface ($R = \infty$). A lens is said to be convex if it is thicker in the middle and narrowing towards the periphery. A lens is concave if it is thicker at periphery and narrows down towards center. Convex lens is visualized to be internal cross section of two spheres (or one sphere and a plane surface) while concave lens is their external cross section (Figs. 9.11-a to 9.11-f). Concavo-convex and convexo-concave lenses are commonly used for spectacles of positive and negative numbers, respectively.

For lenses of material optically denser than the medium in which those are kept, convex lenses have positive focal length [according to Cartesian sign convention] and converge the incident beam while concave lenses have negative focal length and diverge the incident beam.

For most of the applications of lenses, maximum thickness of lens is negligible (at least 50 times smaller) compared with all the other distances such as R_1 and R_2 , u , v , f , etc. Such a lens is called as a thin lens and physical

center of such a lens can be assumed to be the common pole (or optical center) for both its refracting surfaces.

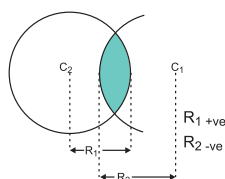


Fig. 9.11 (a): Convex lens as internal cross section of two spheres.

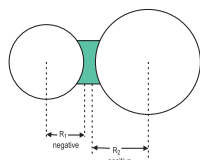


Fig. 9.11 (b): Concave lens as external cross section of two spheres.

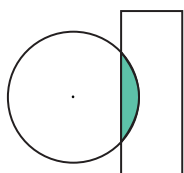


Fig. 9.11 (c): Plano convex lens

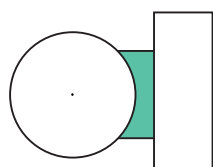


Fig. 9.11 (d): Plano concave lens

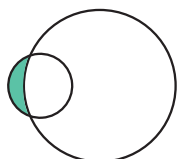


Fig. 9.11 (e): concave-convex lens

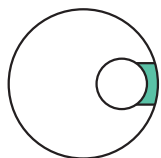


Fig. 9.11 (f) convex-concave lens

For any thin lens, $\frac{1}{f} = \frac{1}{v} - \frac{1}{u}$ --- (9.2)

If necessary, we can have a number of thin lenses in contact with each other having common principal axis. Focal power of such combination is given by the algebraic addition (by considering $\pm$ signs) of individual focal powers.

$$\begin{aligned} \therefore \frac{1}{f} &= \sum \left(\frac{1}{f_i} \right) = \frac{1}{f_1} + \frac{1}{f_2} + \frac{1}{f_3} + \dots \\ &= P_1 + P_2 + P_3 + \dots = \sum P_i = P \end{aligned}$$

For only two thin lenses, separated in air by distance d ,

$$\frac{1}{f} = \frac{1}{f_1} + \frac{1}{f_2} - \frac{d}{f_1 f_2} = P_1 + P_2 - d P_1 P_2 = P$$

For lenses, the relations between u , v , R and f depend also upon the refractive index n of the material of the lens. The relation $\left(f = \frac{R}{2} \right)$ does NOT hold good for lenses. Below we shall derive the necessary relation by considering refraction at the two surfaces of a lens independently.

Unless mentioned specifically, we assume lenses to be made up of optically denser material compared to the medium in which those are kept, e.g., glass lenses in air or in water, etc. As special cases we may consider lenses of rarer medium such as an air lens in water or inside a glass. A spherical hole inside a glass slab is also a lens of rarer medium. In such case, physically (or geometrically or shape-wise) convex lens diverges the incident beam while concave lens converges the incident beam.

Refraction at a single spherical surface:

Consider a spherical surface YPY' of radius of curvature R , separating two transparent media of refractive indices n_1 and n_2 respectively with $n_1 < n_2$. P is the pole and X'PX is the principal axis. A point object O is at an object distance $-u$ from the pole, in the medium of refractive index n_1 . Convexity or concavity of a surface is always with respect to the incident rays, i.e., with respect to a real object. Hence in this case the surface is convex (Fig. 9.12).

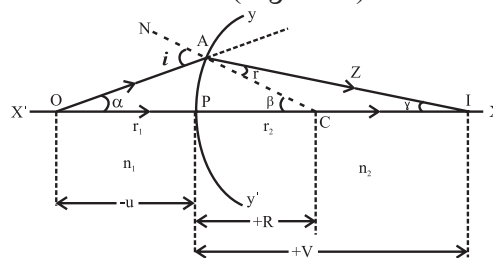


Fig. 9.12: Refraction at a single refracting surface.

To locate its image and in order to minimize spherical aberration, we consider two paraxial rays. The ray OP along the principal axis travels undeviated along PX. Another ray OA strikes the surface at A. CAN is the normal from center of curvature C of the surface at A. Angle of incidence in the medium n_1 at A is i .

As $n_1 < n_2$, the ray deviates towards the normal, travels along AZ and cuts the principal axis at I. Thus, real image of point object O is formed at I. Angle of refraction in medium n_2 is r . According to Snell's law,

$$n_1 \sin(i) = n_2 \sin(r) \quad \text{--- (9.3)}$$

Let α, β and γ be the angles subtended by incident ray, normal and refracted ray with the principal axis.

$$\therefore i = \alpha + \beta \text{ and } r = \beta - \gamma$$

For paraxial rays, all these angles are small and PA can be considered as an arc for α, β and γ .

$$\therefore \sin(i) \cong i \text{ and } \sin(r) \cong r$$

$$\text{Also, } \alpha \cong \frac{\text{arc } AP}{PO} = \frac{\text{arc } AP}{-u},$$

$$\beta = \frac{\text{arc } AP}{PC} = \frac{\text{arc } AP}{R} \text{ and}$$

$$\gamma \cong \frac{\text{arc } AP}{PI} = \frac{\text{arc } AP}{v}$$

$$\therefore n_1 i = n_2 r$$

$$\therefore n_1 (\alpha + \beta) = n_2 (\beta - \gamma)$$

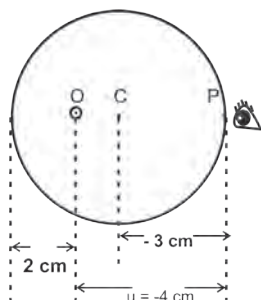
$$\therefore (n_2 - n_1) \beta = n_2 \gamma + n_1 \alpha$$

Substituting α, β and γ and canceling 'arc AP', we get

$$\frac{n_2 - n_1}{R} = \frac{n_2}{v} - \frac{n_1}{u} \quad \text{--- (9.4)}$$

Example 6: A glass paper-weight ($n = 1.5$) of radius 3 cm has a tiny air bubble trapped inside it. Closest distance of the bubble from the surface is 2 cm. Where will it appear when seen from the other end (from where it is farthest)?

Solution: Accompanying Figure below illustrates the location of the bubble.



According to the symbols used in the Eq. (9.4), we get,

n_1 = refractive index of the medium of real object (medium of incident rays) = 1.5

n_2 = refractive index of the other medium = 1

$$u = -4 \text{ cm}$$

$$v = ?$$

$$R = -3 \text{ cm}$$

$$\frac{n_2 - n_1}{R} = \frac{n_2}{v} - \frac{n_1}{u}$$

$$\therefore \frac{1 - 1.5}{-3} = \frac{1}{v} - \frac{1.5}{-4} \therefore \frac{1}{6} = \frac{1}{v} + \frac{3}{8} \therefore v = -4.8 \text{ cm}$$

In this case apparent depth is NOT less than real depth. This is due to curvature of the refracting surface.

In this case (Fig. 9.12) we had considered the object placed in rarer medium, real image in denser medium and the surface facing the object to be convex. However, while deriving the relation, all the symbolic values (which could be numeric also) were substituted as per the Cartesian sign convention (e.g. 'u' as negative, etc.). Hence the final expression (Eq. 9.4) is applicable to any surface separating any two media, and real or virtual image provided you substitute your values (symbolic or numerical) as per Cartesian sign convention. The *only restriction* is that n_1 is for medium of real object and n_2 is the *other* medium (*not necessarily* the medium of image). Only in the case of real image, it will be in medium n_2 . If virtual, it will be in the medium n_1 (with image distance negative how do you justify this?).

We strongly suggest you to do the derivations yourself for any other special case such as object placed in the denser medium, virtual image, concave surface, etc. It must be remembered that in any case you will land up with the same expression as in Eq. (9.4).

Lens makers' equation: Relation between refractive index (n), focal length (f) and radii of curvature R_1 and R_2 for a thin lens.

Consider a lens of radii of curvature R_1 and R_2 kept in a medium such that n is refractive

index of material of the lens with respect to the outside medium. Assuming the lens to be thin, P is the common pole for both the surfaces. O is a point object on the principal axis at a distance u from P. First refracting surface of the lens of radius of curvature R_1 faces the object (Fig 9.13).

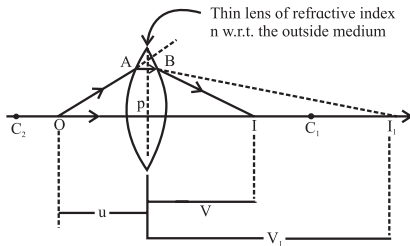


Fig. 9.13: Lens maker's equation.

Axial ray OP travels undeviated. Paraxial ray OA deviates towards normal and would intersect axis at I_1 , in the absence of second refracting surface. $PI_1 = v_1$ is the image distance for intermediate image I_1 .

Thus, the symbols to be used in Eq. (9.4) are

$$n_2 = n, \quad n_1 = 1, \quad R = R_1, \quad u = u, \quad v = v_1$$

$$\therefore \frac{n-1}{R_1} = \frac{n}{v_1} - \frac{1}{u} \quad \text{--- (9.5)}$$

(Not that, in this case, we are not substituting the algebraic values but just using different symbols.)

Before reaching I_1 , the ray PI_1 is intercepted at B by the second refracting surface. In this case, the incident rays AB and OP are in the medium of refractive index n and converging towards I_1 . Thus, I_1 acts as *virtual* object for second surface of radius of curvature (R_2) and object distance is ($u = v_1$). As the incident rays are in the medium of refractive index n , this is the medium of (virtual) object $\therefore n_1 = n$ and refractive index of the other medium is $n_2 = 1$.

After refraction, the ray bends away from the normal and intersects the principal axis at I which is the real image of object O formed due to the lens. $\therefore PI = v$.

Substituting all these symbols in Eq. (9.4), we get

$$\frac{1-n}{R_2} = \frac{n-1}{-R_2} = \frac{1}{v} - \frac{n}{(v_1)} \quad \text{--- (9.6)}$$

Adding Eq. (9.5) and (9.6), we get,

$$(n-1) \left(\frac{1}{R_1} - \frac{1}{R_2} \right) = \frac{1}{v} - \frac{1}{u}$$

For $u = \infty, v = f$

$$\therefore \frac{1}{f} = (n-1) \left(\frac{1}{R_1} - \frac{1}{R_2} \right) \quad \text{--- (9.7)}$$

For preparing spectacles, it is necessary to grind the glass (or acrylic, etc.) for having the desired radii of curvature. Equation (9.7) can be used to calculate the radii of curvature for the lens, hence it is called the lens makers' equation. (It should be remembered that while solving problems when you are using equations 9.1, 9.2, 9.4, 9.7, etc., we will be substituting the values of the corresponding quantities. Hence this time it is algebraic substitution, i.e., with

Special cases:

Most popular and most common special case is the one in which we have a thin, symmetric, double lens. In this case, R_1 and R_2 are numerically equal.

(A) Thin, symmetric, double convex lens: R_1 is positive, R_2 is negative and numerically equal. Let $|R_1| = |R_2| = R$.

$$\therefore \frac{1}{f} = (n-1) \left(\frac{1}{R} - \frac{1}{-R} \right) = \frac{2(n-1)}{R}$$

Further, for popular variety of glasses, $n \cong 1.5$. In such a case, $f = R$.

(B) Thin, symmetric, double concave lens: R_1 is negative, R_2 is positive and numerically equal. Let $|R_1| = |R_2| = R$.

$$\therefore \frac{1}{f} = (n-1) \left(\frac{1}{-R} - \frac{1}{R} \right) = \frac{2(n-1)}{-R}$$

Further if $n \cong 1.5, f = -R$

(C) Thin, planoconvex lenses: One radius is R and the other is ∞ . $\therefore \frac{1}{f} = \frac{n-1}{R}$

Further if $n \cong 1.5, f = 2|R|$

proper $\pm$ sign)

Example 7: A dense glass double convex lens ($n = 2$) designed to reduce spherical aberration has $|R_1|:|R_2|=1:5$. If a point object is kept 15 cm in front of this lens, it produces its real image at

7.5 cm. Determine R_1 and R_2 .

Solution: $u = -15$ cm, $v = +7.5$ cm (real image is on opposite side).

$$\frac{1}{f} = \frac{1}{v} - \frac{1}{u} \therefore \frac{1}{f} = \frac{1}{7.5} - \frac{1}{-15} \therefore f = +5 \text{ cm}$$

The lens is double convex. Hence, R_1 is positive and R_2 is negative. Also, $|R_2| = 5|R_1|$ and $n = 2$.

$$(n-1) \left(\frac{1}{R_1} - \frac{1}{R_2} \right) = \frac{1}{f}$$

$$\therefore (2-1) \left(\frac{1}{R_1} - \frac{1}{-(5R_1)} \right) = \frac{1}{5}$$

$$\therefore (1) \left(\frac{6}{5R_1} \right) = \frac{1}{5} \therefore R_1 = 6 \text{ cm} \therefore R_2 = 30 \text{ cm}$$

9.8 Dispersion of light and prisms:

The colour of light that we see depends upon the frequency of that ray (wave). The refractive index of a material also depends upon the frequency of the wave and increases with frequency. Obviously refractive index of light is different for different colours. As a result, for an obliquely incident ray, the angles of refraction are different for each colour and they separate (disperse) as they travel along different directions. This phenomenon is called angular dispersion Fig 9.14.

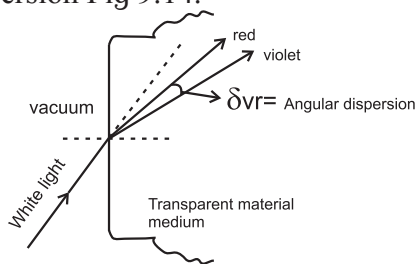


Fig. 9.14: Angular dispersion at a single surface.

If a polychromatic beam of light (bundle of rays of different colours) is obliquely incident upon a plane parallel transparent slab, emergent beam consists of all component colours separated out. However, in this case all those are parallel to each other and also parallel to initial direction. This is lateral dispersion which is measured as the perpendicular distance between the direction of incident ray and respective directions of dispersed emergent rays (L_R and L_V) Fig 9.15. For it to be easily detectable, the

parallel surfaces must be separated over very large distance and i should be large.

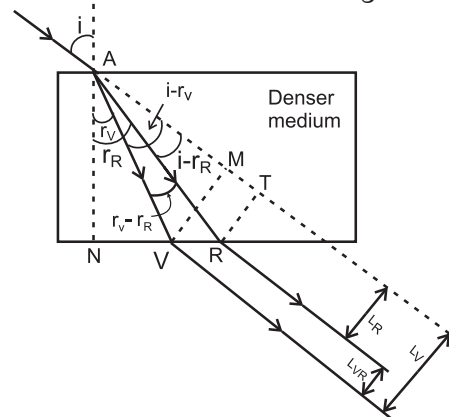


Fig. 9.15: Lateral dispersion due to plane parallel slab.

Example 8: A fine beam of white light is incident upon the longer side of a plane parallel glass slab of breadth 5 cm at angle of incidence 60° . Calculate angular deviation of red and violet rays within the slab and lateral dispersion between them as they emerge from the opposite side. Refractive indices of the glass for red and violet are 1.51 and 1.53 respectively.

Solution: As shown in the Fig. 9.15 above, $VM = L_V$ and $RT = L_R$ give respective lateral deviations for violet and red colours and $L_{VR} = L_V - L_R$ is the lateral dispersion between these colours. $n_R = 1.51$, $n_V = 1.53$ and $i = 60^\circ$

$$\therefore \sin r_R = \frac{\sin i}{n_R} = \frac{\sin 60^\circ}{1.51} = 0.5735$$

$$\sin r_V = \frac{\sin i}{n_V} = \frac{\sin 60^\circ}{1.53} = 0.566$$

$$\therefore r_R = 35^\circ \text{ and } r_V = 34^\circ 28' \therefore \delta_{RV} = r_R - r_V = 32'$$

$$\therefore i - r_R = 25^\circ, i - r_V = 25^\circ 32'$$

$$\therefore AR = \frac{AN}{\cos r_R} = 6.104 \text{ cm}$$

$$AV = \frac{AN}{\cos r_V} = 6.063 \text{ cm}$$

$$\therefore L_R = RT = AR (\sin [i - r_R]) = 2.58 \text{ cm}$$

$$L_V = VM = AV (\sin [i - r_V]) = 2.58 \text{ cm}$$

$$\therefore L_{VR} = L_V - L_R = 0.033 \text{ cm} = 0.33 \text{ mm}$$

It shows that the lateral dispersion is too small to detect.

In order to have appreciable and observable dispersion, two parallel surfaces are not useful. In such case we use prisms, in which two refracting surfaces inclined at an angle are used. Popular variety of prisms are having three rectangular surfaces forming a triangle. At a time two of these are taking part in the refraction. The one, not involved in refraction is called base of the prism. Fig 9.16.

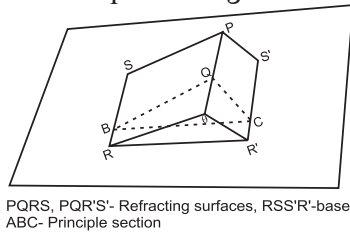


Fig. 9.16: Prism consisting of three plane surfaces.

Any section of prism perpendicular to the base is called principal section of the prism. Usually we consider all the rays in this plane. Fig 9.17 a and 9.17 b show refraction through a prism for monochromatic and white beams respectively. Angular dispersion is shown for white beam.

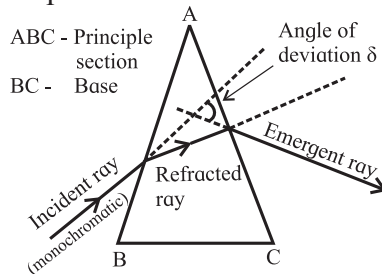


Fig. 9.17 (a): Refraction through a prism (monochromatic light).

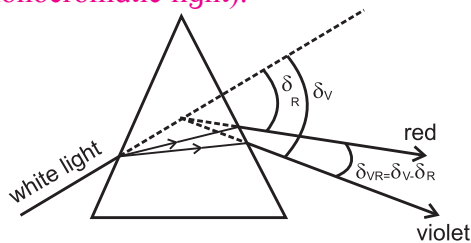


Fig. 9.17 (b): Angular dispersion through a prism. (white light).

Relations between the angles involved: Figure 9.18 shows principal section ABC of a prism of absolute refractive index n kept in air. Refracting surfaces AB and AC are inclined at angle A , which is refracting angle of prism or simply 'angle of prism'. Surface BC is the base. A monochromatic ray PQ obliquely strikes first

reflecting surface AB. Normal passing through the point of incidence Q is MQN. Angle of incidence at Q is i . After refraction at Q, the ray deviates towards the normal and strikes second refracting surface AC at R which is the point of emergence. MRN is the normal through R. Angles of refraction at Q and R are r_1 and r_2 respectively.

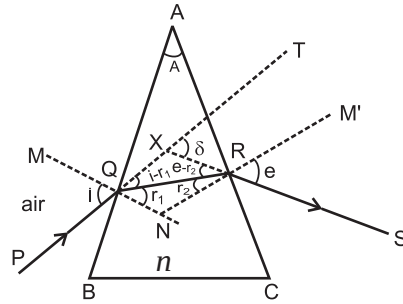


Fig. 9.18: Deviation through a prism.

After R, the ray deviates away from normal and finally emerges along RS making e as the angle of emergence. Incident ray PQ is extended as QT. Emergent ray RS meets QT at X if traced backward. Angle TXS is angle of deviation δ .

$\angle AQN = \angle ARN = 90^\circ$ (Angles at normal)

$\therefore$ From quadrilateral AQNR,

$$A + \angle QNR = 180^\circ \quad \text{--- (9.8)}$$

$$\text{From } \Delta QNR, r_1 + r_2 + \angle QNR = 180^\circ \quad \text{--- (9.9)}$$

$\therefore$ From Eqs. (9.8) and (9.9),

$$A = r_1 + r_2 \quad \text{--- (9.10)}$$

Angle δ is exterior angle for triangle XQR.

$$\therefore \angle XQR + \angle XRQ = \delta$$

$$\therefore (i - r_1) + (e - r_2) = \delta$$

$$\therefore (i + e) - (r_1 + r_2) = \delta$$

Hence, using Eq. (9.10), $(i + e) - (A) = \delta$

$$\therefore i + e = A + \delta \quad \text{--- (9.11)}$$

Deviation curve, minimum deviation and prism formula: From the relations (9.10) and (9.11), it is clear that δ, e, r_1 and r_2 depend upon i, A and n . After a certain minimum value of angle of incidence $i_{\min}$, the emergent ray is possible. This is because of the fact that for $i < i_{\min}$, $r_2 > i_c$ and there is total internal reflection at the second surface and there is no emergent ray. This will be shown later. Then onwards,

as i increases, r_1 increases as $\frac{\sin i}{\sin r_1} = n$ but r_2 and e decrease. However, variation in δ with increasing i is different. It is as plotted in the Fig. 9.19.

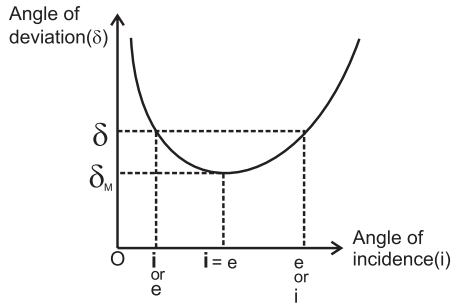


Fig. 9.19: Deviation curve for a prism.

It shows that, with increasing values of i , the angle of deviation δ decreases initially to a certain minimum (δ_m) and then increases. It should also be noted that the curve is not a symmetric parabola, but the slope in the part after is less. It is clear that except at $\delta = \delta_m$, (Angle of minimum deviation) there are two values of i for any given δ . By applying the principle of reversibility of light to path PQRS it is obvious that if one of these values is i , the other must be e and vice versa. Thus at $\delta = \delta_m$, we have $i = e$. Also, in this case, $r_1 = r_2$

$$\text{and } A = r_1 + r_2 = 2r \therefore r = \frac{A}{2}$$

Only in this case QR is parallel to base BC and the figure is symmetric.

Using these in Eq. (9.11), we get,

$$i + i = A + \delta_m \therefore i = \frac{(A + \delta_m)}{2}$$

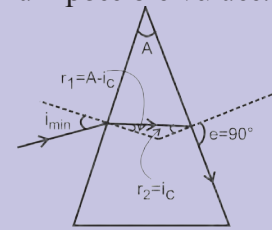
According to Snell's law,

$$\therefore n = \frac{\sin\left(\frac{A + \delta_m}{2}\right)}{\sin\left(\frac{A}{2}\right)} \quad \text{--- (9.12)}$$

Equation (9.12) is called prism formula.

Example 9.9: For a glass ($n=1.5$) prism having refracting angle 60° , determine the range of angle of incidence for which emergent ray is possible from the opposite surface and the corresponding angles of emergence. Also calculate the angle of incidence for which $i = e$. How much is the corresponding angle of minimum deviation?

(I) Grazing emergence and minimum angle of incidence: At the point of emergence, the ray travels from a denser medium into rarer (popular prisms are of denser material, kept in rarer). Thus if $r_2 = \sin^{-1}\left(\frac{1}{n}\right)$ is the critical angle, the angle of emergence $e = 90^\circ$. This is called grazing emergence or we say that the ray just emerges. Angle of prism A is constant for a given prism and $A = r_1 + r_2$. Hence the corresponding r_1 and i will have their minimum possible values.



Grazing emergence or the ray just emerges

(II) For commonly used glass prisms,

$$n = 1.5, \sin^{-1}\left(\frac{1}{n}\right) = \sin^{-1}\left(\frac{1}{1.5}\right) = 41^\circ 49' = (r_2)_{max}$$

If prism is symmetric (equilateral),

$$A = 60^\circ \therefore r_1 = 60^\circ - 41^\circ 49' = 18^\circ 11'$$

$$\therefore n = 1.5 = \frac{\sin(i_{min})}{\sin 18^\circ 11'} \therefore i_{min} = 27^\circ 55' \cong 28^\circ$$

(III) For a symmetric (equilateral) prism, the prism formula can be written as

$$n = \frac{\sin\left(\frac{60 + \delta_m}{2}\right)}{\sin\left(\frac{60}{2}\right)} = \frac{\sin\left(30 + \frac{\delta_m}{2}\right)}{\sin(30)} = 2\sin\left(30 + \frac{\delta_m}{2}\right)$$

(IV) For a prism of denser material, kept in a rarer medium, the incident ray deviates towards the normal during the first refraction and away from the normal during second refraction. However, during both the refractions it deviates towards the base only.

Solution: As shown in the box above, $i_{\min} = 27^\circ 55'$. Angle of emergence for this is $e_{\max} = 90^\circ$.

From the principle of reversibility of light, $i_{\max} = 90^\circ$ and $e_{\min} = 27^\circ 55'$

Also, from the box above,

$$n = \frac{\sin\left(\frac{60 + \delta_m}{2}\right)}{\sin\left(\frac{60}{2}\right)}$$

$$= \frac{\sin\left(30 + \frac{\delta_m}{2}\right)}{\sin(30)} = 2\sin\left(30 + \frac{\delta_m}{2}\right)$$

$$\therefore 1.5 = 2\sin\left(30 + \frac{\delta_m}{2}\right) \therefore 0.75 = \sin\left(30 + \frac{\delta_m}{2}\right)$$

$$\therefore \left(30 + \frac{\delta_m}{2}\right) = 48^\circ 35'$$

$$\therefore \frac{\delta_m}{2} = 18^\circ 35' \therefore \delta_m = 37^\circ 10'$$

$$i + e = A + \delta \quad \text{and} \quad i = e \text{ for } \delta = \delta_m$$

$$\therefore i + i = 60 + 37^\circ 10' = 97^\circ 10' \therefore i = 48^\circ 35'$$

Thin prisms: Prisms having refracting angle less than 10° ($A < 10^\circ$) are called thin prisms. For such prisms we can comfortably use $\sin \theta \cong \theta$. For such prisms to deviate the incident ray towards the base during both refractions, it is essential that i should also be less than 10° so that all the other angles will also be small.

Thus

$$n = \frac{\sin i}{\sin r_1} \cong \frac{i}{r_1} \text{ and } n = \frac{\sin e}{\sin r_2} \cong \frac{e}{r_2}$$

$$\therefore i \cong nr_1 \text{ and } e \cong nr_2$$

Using these in Eq. (9.11), we get,

$$i + e = nr_1 + nr_2 = n(r_1 + r_2) = nA = A + \delta$$

$$\therefore \delta = A(n-1) \quad \text{--- (9.13)}$$

A and n are constant for a given prism. Thus, for a thin prism, for small angles of incidences, angle of deviation is constant (independent of angle of incidence).

Angular dispersion and mean deviation:

As discussed earlier, if a polychromatic beam is incident upon a prism, the emergent beam consists of all the individual colours angularly

separated. This is angular dispersion (Fig. 9.20).

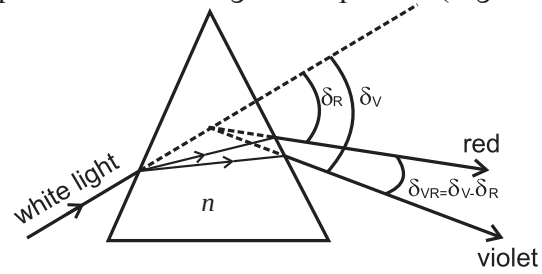


Fig. 9.20: Angular dispersion through a prism.

It is measured for any two component colours.

$$\therefore \delta_{21} = \delta_2 - \delta_1$$

Normally we do it for extreme colours.

For white light, violet and red are the extreme colours.

$$\therefore \delta_{VR} = \delta_V - \delta_R$$

Using deviation for thin prism (Eq. 9.13), we can write

$$\therefore \delta_{21} = \delta_2 - \delta_1 = A(n_2 - 1) - A(n_1 - 1)$$

$$= A(n_2 - n_1)$$

where n_1 and n_2 are refractive indices for the two colours.

Also,

$$\delta_{VR} = \delta_V - \delta_R = A(n_V - 1) - A(n_R - 1)$$

$$= A(n_V - n_R) \quad \text{--- (9.14)}$$

Yellow is practically chosen to be the mean colour for violet and red.

This gives mean deviation

$$\delta_{VR} = \frac{\delta_V + \delta_R}{2} \cong \delta_Y = A(n_Y - 1) \quad \text{--- (9.15)}$$



Do you know ?

- (i) If you see a rainbow widthwise, yellow appears to be centrally located. Hence angular deviation of yellow is average for the entire colour span. This may be the reason for choosing yellow as the mean colour. Remember, red band is widest and violet is much thinner than blue.
- (ii) While obtaining the expression for ω , we have used thin prism formula for δ . However, the expression for ω (equation 9.16) is valid as well for equilateral prisms or right-angled prisms.

Dispersive power: Ability of an optical material to disperse constituent colours is its dispersive power. It is measured for any two colours as the ratio of angular dispersion to the mean deviation for those two colours. Thus, for the extreme colours of white light, dispersive power is given by

$$\omega = \frac{[\delta_V - \delta_R]}{\left[\frac{\delta_V + \delta_R}{2} \right]} \approx \frac{\delta_V - \delta_R}{\delta_Y}$$

$$= \frac{A(n_V - n_R)}{A(n_Y - 1)} = \frac{n_V - n_R}{n_Y - 1} \quad \text{--- (9.16)}$$

As ω is the ratio of same physical quantities, it is unitless and dimensionless quantity. From the expression in terms of refractive indices it should be understood that dispersive power depends only upon refractive index (hence material only) and not upon the dimensions of prism. For commonly used glasses it is around 0.03.

Example 10: For a dense flint glass prism of refracting angle 10° , obtain angular deviation for extreme colours and dispersive power of dense flint glass. ($n_{red} = 1.712$, $n_{violet} = 1.792$)

$$\delta_V = A(n_V - 1) = 10(1.792 - 1) = (7.92)^\circ$$

$$\delta_R = A(n_R - 1) = 10(1.712 - 1) = (7.12)^\circ$$

$$\therefore \text{Angular dispersion, } \delta_{VR} = \delta_V - \delta_R = (0.8)^\circ$$

dispersive power, $\omega =$

$$= \frac{\delta_V - \delta_R}{\left(\frac{\delta_V + \delta_R}{2} \right)}$$

$$= 2 \left(\frac{7.92 - 7.12}{7.92 + 7.12} \right)$$

$$= \frac{2 \times 0.8}{15.04} = 0.1064$$

(This is much higher than popular crown glass)

9.9 Some natural phenomena due to Sunlight:

Mirage: On a hot clear Sunny day, along a level road, a pond of water appears to be there ahead. However, if we physically reach the spot, there is nothing but the dry road and water pond again appears ahead. This illusion

is called a mirage (Fig. 9.21).

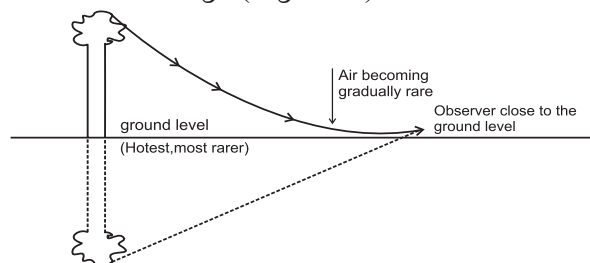


Fig. 9.21: The Mirage.

On a hot day the air in contact with the road is hottest and as we go up, it gets gradually cooler. The refractive index of air thus increases with height. As shown in the figure, due to this gradual change in the refractive index, the ray of light coming from the top of an object becomes more and more horizontal as it almost touches the road. For some reason (mentioned later) it bends above. Then onwards, upward bending continues due to denser air. As a result, for an observer, it appears to be coming from below thereby giving an illusion of reflection from an (imaginary) water surface.

Rainbow: Undoubtedly, rainbow is an eye-catching phenomenon occurring due to rains and Sunlight. It is most popular because it is observable from anywhere on the Earth. A few lucky persons might have observed two rainbows simultaneously one above the other. Some might have seen a complete circular rainbow from an aeroplane (Of course, this time it's not a bow!). Optical phenomena discussed till now are sufficient to explain the formation of a rainbow.

The facts to be explained are:

- (i) It is seen during rains and on the opposite side of the Sun.
- (ii) It is seen only during mornings and evenings and not throughout the day.
- (iii) In the commonly seen rainbow red arch is outside and violet is inside.
- (iv) In the rarely occurring concentric secondary rainbow, violet arch is outside and red is inside.
- (v) It is in the form of arc of a circle.
- (vi) Complete circle can be seen from a higher altitude, i.e., from an aeroplane.

- (vii) Total internal reflection is not possible in this case.

Conditions necessary for formation of a rainbow: Light shower with relatively large raindrops, morning or evening time and enough Sunlight.

Optical phenomena involved: During the formation of a rainbow, the rays of Sunlight incident on water drops, deviate and disperse during refraction, internally (NOT total internally) reflect once (for primary rainbow) or twice (for secondary rainbow) and finally refract again into air. At all stages there is angular dispersion which leads to clear separation of the colours.

Primary rainbow: Figure 9.22 (a) shows the optical phenomena involved in the formation of a primary rainbow due to a spherical water drop.



Do you know ?

Possible reasons for the upward bending at the road during mirage could be:

- Angle of incidence at the road is glancing. At glancing incidence, the reflection coefficient is very large which causes reflection.
- Air almost in contact with the road is not steady. The non-uniform motion of the air bends the ray upwards and once it has bent upwards, it continues to do so.
- Using Maxwell's equations for EM waves, correct explanation is possible for the reflection.

It may be pointed out that total internal reflection is NEVER possible here because the relative refractive index is just less than 1 and hence the critical angle (discussed in the article 9.6) is also approaching 90° .

White ray AB from the Sun strikes from upper portion of a water drop at an incident angle i . On entering into water, it deviates and disperses into constituent colours. Extreme colours violet(V) and red(R) are shown. Refracted rays BV and BR strike the opposite inner surface of water drop and suffer internal (NOT total internal) reflection. These reflected rays finally

emerge from V' and R' and can be seen by an observer on the ground. For the observer they appear to be coming from opposite side of the Sun. Minimum deviation rays of red and violet colour are inclined to the ground level at $\theta_R = 42.8^\circ \cong 43^\circ$ and $\theta_V = 40.8^\circ \cong 41^\circ$ respectively. As a result, in the 'bow' or arch, the red is above or outer and violet is lower or inner.

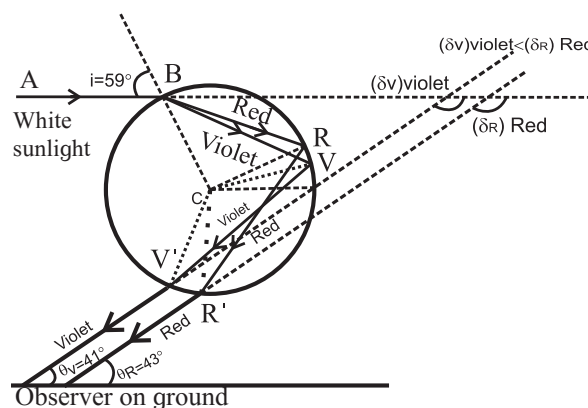


Fig. 9.22 (a): Formation of primary rainbow.

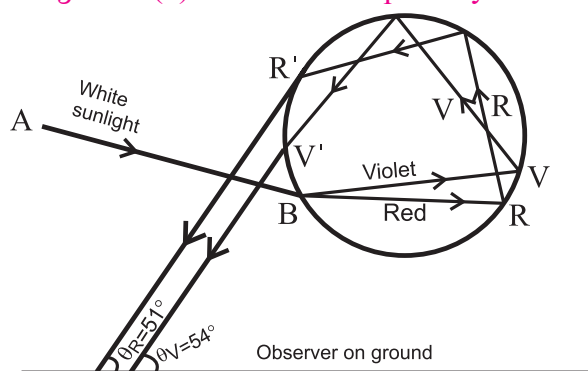


Fig. 9.22 (b): Formation of secondary rainbow.

Secondary rainbow: Figure 9.22 (b) shows some optical phenomena involved in the formation of a secondary rainbow due to a spherical water drop. White ray AB from the Sun strikes from lower portion of a water drop at an incident angle i . On entering into water, it deviates and disperses into constituent colours. Extreme colours violet(V) and red(R) are shown. Refracted rays BV and BR finally emerge the drop from V' and R' after suffering two internal reflections and can be seen by an observer on the ground. Minimum deviation rays of red and violet colour are inclined to the ground level at $\theta_R \cong 51^\circ$ and $\theta_V \cong 53^\circ$ respectively. As a result, in the 'bow' or arch, the violet is above or outer and red is lower or inner.



Do you know ?

(I) Why total internal reflection is not possible during formation of a rainbow?

Angle of incidence i in air, at the water drop, can't be greater than 90° . As a result, angle of refraction r in water will *always less than the critical angle*. From Fig a and b and by simple geometry, it is clear that this r itself

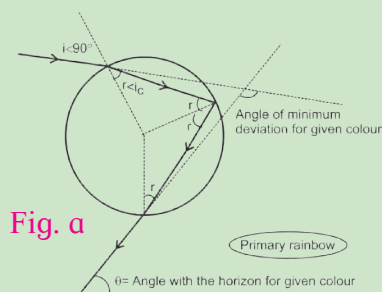


Fig. a

is the angle of incidence at any point for one or more internal reflections. Obviously, total internal

reflection is not possible.

(II) Why is rainbow seen only for a definite angle range with respect to the ground?

For clear visibility we must have a beam of enough intensity. From the deviation curve (Fig 9.19) it is clear that near minimum deviation the curve is almost parallel to x-axis, i.e., for majority of angles of incidence in this range, the angle of deviation is nearly the same

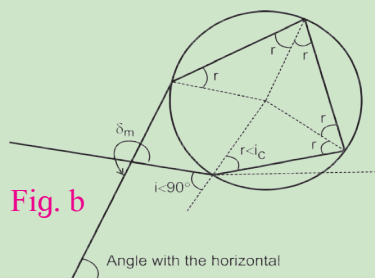


Fig. b

and those are almost parallel forming a beam of enough intensity. Thus, the rays in the near vicinity of minimum

deviation are almost parallel to each other. Rays beyond this range suffer wide angular dispersion and thus will not have enough intensity for visibility.

By using simple geometry for Figs. a and b it can be shown that the angle of deviation between final emergent ray and the incident ray is $\delta = \pi + 2i - 4r$ during primary rainbow, and $\delta = 2\pi + 2i - 6r$ during secondary rainbow. Using these relations and Snell's law $\sin i = n \sin r$, we can obtain derivatives of δ . Second derivative of δ comes out to be negative, which shows that it is the minima condition. Equating first derivative to zero we can obtain corresponding values

of i and r . Again, by using Figs. a and b, we can obtain the corresponding angles θ_R and θ_V at the horizontal, which is the visible angular position for the rainbow.

(III) Why is the rainbow a bow or an arch? Can we see a complete circular rainbow?

Figure c illustrates formation of primary and secondary rainbows with their common centre O is the point where the line joining the sun and the observer meets the Earth when extended. P is location of the observer. Different colours of rainbows are seen on arches of cones of respective angles described earlier.

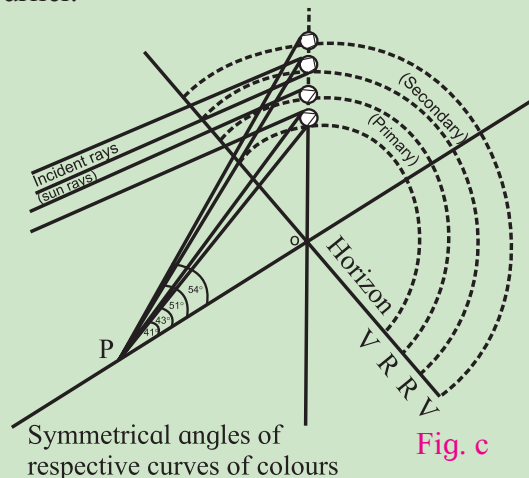


Fig. c

Smallest half angle refers to the cone of violet colour of primary rainbow, which is 41° . As the Sun rises, the common centre of the rainbows moves down. Hence as the Sun comes up, smaller and smaller part of the rainbows will be seen. If the Sun is above 41° , violet arch of primary rainbow cannot be seen. Obviously beyond 53° , nothing will be seen. That is why rainbows are visible only during mornings and evenings.

However, if observer moves up (may be in an aeroplane), the line PO itself moves up making lower part of the arches visible. After a certain minimum elevation, entire circle for all the cones can be visible.

(IV) Size of water drops convenient for rainbow: Water drops responsible for the formation of a rainbow should not be too small. For too small drops the phenomenon of diffraction (redistribution of energy due to obstacles, discussed in XIIth standard) dominates and clear rainbow can't be seen.

9.10 Defects of lenses (aberrations of optical images):

As mentioned in the section 9.4 for aberration for curved mirrors, while deriving various relations, we assume most of the rays to be paraxial by using lenses of small aperture. In reality, we have objects of finite sizes. Also, we need optical devices of large apertures (lenses and/or mirrors of size few meters for telescopes, etc.). In such cases the beam of rays is no more paraxial, quite often not parallel also. As a result, the spherical aberration discussed for spherical mirrors can occur for lenses also. Only one defect is mentioned corresponding to monochromatic beam of light.

Chromatic aberration: In case of mirrors there is no dispersion of light due to refractive index. However, lenses are prepared by using a transparent material medium having different refractive index for different colours. Hence angular dispersion is present. A convex lens can be approximated to two thin prisms connected base to base and for a concave lens those are vertex to vertex. (Fig. 9.23 (a) and 9.23 (b))

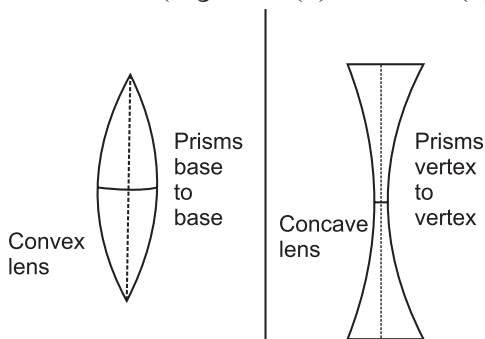


Fig. 9.23: (a) Convex lens (b) Concave lens

If the lens is thick, this will result into notably different foci corresponding to each colour for a polychromatic beam, like a white light. This defect is called chromatic aberration, violet being focused closest to pole as it has maximum deviation. (Fig 9.24 (a) and 9.24 (b)) Longitudinal chromatic aberration, transverse chromatic aberration and circle of least confusion are defined in the same manner as that of spherical aberration for spherical mirrors.

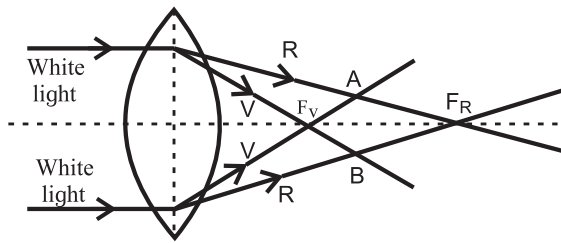


Fig. 9.24: Chromatic aberration: (a) Convex lens.

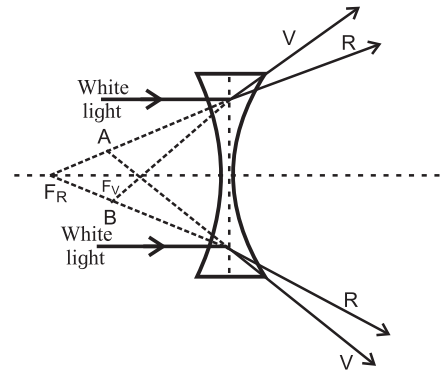


Fig. 9.24: Chromatic aberration: (b) Concave lens

Reducing/eliminating chromatic aberration:

Eliminating chromatic aberration simultaneously for all the colours is impossible. We try to eliminate it for extreme colours which reduces it for other colours. Convenient methods to do it use either a convex and a concave lens in contact or two thin convex lenses with proper separation. Such a combination is called achromatic combination.

Achromatic combination of two lenses in contact: Let ω_1 and ω_2 be the dispersive powers of materials of the two component lenses used in contact for an achromatic combination. Their focal lengths f for violet, red and yellow (assumed to be the mean colour) are suffixed by respective letters V, R and Y.

Also, let $K_1 = \left(\frac{1}{R_1} - \frac{1}{R_2} \right)_1$ for lens 1 and

$K_2 = \left(\frac{1}{R_1} - \frac{1}{R_2} \right)_2$ for lens 2.

For two thin lenses in contact,

$$\frac{1}{f} = \frac{1}{f_1} + \frac{1}{f_2} \dots\dots$$

To be used separately for respective colours.

For the combination to be achromatic, the

resultant focal length of the combination must be the same for both the colours, i.e.,

$$f_V = f_R \text{ or } \frac{1}{f_V} = \frac{1}{f_R}$$

$$\therefore \frac{1}{f_{1V}} + \frac{1}{f_{2V}} = \frac{1}{f_{1R}} + \frac{1}{f_{2R}}$$

$(n_{1V}-1)K_1 + (n_{2V}-1)K_2 = (n_{1R}-1)K_1 + (n_{2R}-1)K_2$
..... using lens makers' Eq. (9.7)

$$\therefore \frac{K_1}{K_2} = \frac{n_{2V} - n_{2R}}{n_{1V} - n_{1R}} \quad \text{--- (9.17)}$$

For mean colour yellow,

$$\frac{1}{f_Y} = \frac{1}{f_{1Y}} + \frac{1}{f_{2Y}}$$

with $\frac{1}{f_{1Y}} = (n_{1Y} - 1)K_1$

and $\frac{1}{f_{2Y}} = (n_{2Y} - 1)K_2$

$$\therefore \frac{K_1}{K_2} = \left(\frac{n_{2Y} - 1}{n_{1Y} - 1} \right) \left(\frac{f_{2Y}}{f_{1Y}} \right) \quad \text{--- (9.18)}$$

Equating R.H.S. of (9.17) and (9.18) and rearranging, we can write

$$\frac{f_{2Y}}{f_{1Y}} = - \left(\frac{n_{2V} - n_{2R}}{n_{2Y} - 1} \right) \div \left(\frac{n_{1R} - n_{1R}}{n_{1Y} - 1} \right)$$

$$= - \frac{\omega_2}{\omega_1} \quad \text{--- (9.19)}$$

Equation (9.19) is the condition for achromatic combination of two lenses, in contact.

Dispersive power ω is always positive. Thus, one of the lenses must be convex and the other concave.

If second lens is concave, f_{2Y} is negative.

$$\therefore \frac{1}{f_Y} = \frac{1}{f_{1Y}} - \frac{1}{f_{2Y}}$$

For this combination to be converging, f_Y should be positive.

Hence, $f_{1Y} < f_{2Y}$ and $\omega_1 < \omega_2$

Thus, for an achromatic combination if there is a choice between flint glass ($n = 1.655$) and crown glass ($n = 1.517$), the convergent (convex) lens must be of crown glass and the divergent (concave) lens of flint glass.

Example 9.11: After Cataract operation, a person is recommended with concavo-convex spectacles of curvatures 10 cm and 50 cm. Crown glass of refractive indices 1.51 for red and 1.53 for violet colours is used for this. Calculate the lateral chromatic aberration occurring due to these glasses.

Solution: For a concavo-convex lens, both the radii of curvature are either positive or both negative. If convex shape faces object, both will be positive. See the accompanying figure.

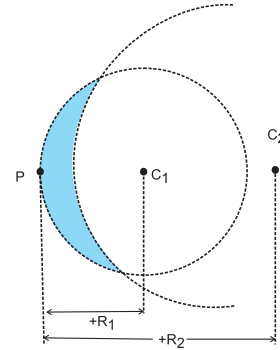


Fig. Concavo-convex lens with convex face receiving incident rays

$$\therefore R_1 = 10 \text{ cm and } R_2 = 50 \text{ cm}$$

$$\therefore \left(\frac{1}{R_1} - \frac{1}{R_2} \right) = \left(\frac{1}{+10} - \frac{1}{+50} \right) = 0.08 \text{ cm}^{-1}$$

$$\therefore \frac{1}{f_R} = (n_R - 1) \left(\frac{1}{R_1} - \frac{1}{R_2} \right)$$

$$= (1.51 - 1) \times 0.08 = 0.0408$$

$$\therefore f_R = 25.51 \text{ cm}$$

$$\text{and } \frac{1}{f_V} = (n_V - 1) \left(\frac{1}{R_1} - \frac{1}{R_2} \right)$$

$$= (1.53 - 1) \times 0.08 = 0.0424$$

$$\therefore f_V = 23.58 \text{ cm}$$

$$\therefore \text{Longitudinal chromatic aberration}$$

$$= f_V - f_R = 25.51 - 23.58$$

$$= 1.93 \text{ cm, ... (quite appreciable!)}$$

Verify that you get the same answer even if you consider the concave surface facing the incident rays.

Spherical aberration: Longitudinal spherical aberration, transverse spherical aberration and circle of least confusion are defined in the same manner as that for spherical mirrors. (Fig 9.25 (a) and 9.25 (b))

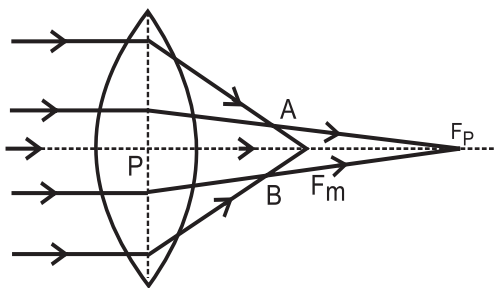


Fig. 9.25 (a): Spherical aberration, Convex lens.

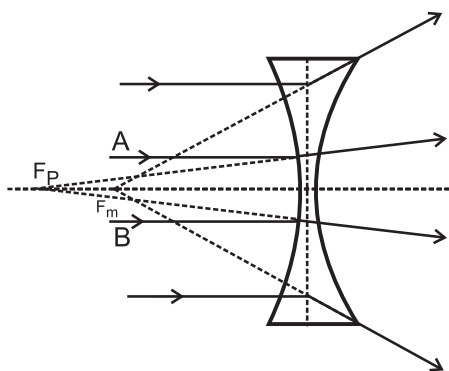


Fig. 9.25 (b): Spherical aberration, Concave lens

Methods to reduce/eliminate spherical aberration of lenses:

- (i) Cheapest method to reduce the spherical aberration is to use a planoconvex or planoconcave lens with curved side facing the incident rays (real object). Reversing it increases the aberration appreciably.
- (ii) Certain ratio of radii of curvature for a given refractive index almost eliminates the spherical aberration. For $n = 1.5$, the ratio is $\frac{R_1}{R_2} = \frac{1}{6}$ and for $n = 2$, it is $\frac{1}{5}$
- (iii) Use of two thin converging lenses separated by distance equal to difference between their focal lengths with lens of larger focal length facing the incident rays considerably reduces spherical aberration.
- (iv) Spherical aberration of a convex lens is positive (for real image), while that of a concave lens is negative. Thus, a suitable combination of them (preferably a double convex lens of smaller focal length and a planoconcave lens of greater focal length) can completely eliminate spherical aberration.

9.11 Optical instruments:

Introduction: Whether an object appears bigger or not does not necessarily depend upon its own size. Huge mountains far off may appear smaller than a small tree close to us. This is because the angle subtended by the mountain at the eye from that distance (called the visual angle) is smaller than that subtended by the tree from its position. Hence, apparent size of an object depends upon the visual angle subtended by the object from its position. Obviously, for an object to appear bigger, we must bring it closer to us or we should go closer to it.

However, due to the limitation for focusing the eye lens it is not possible to take an object closer than a certain distance. This distance is called least distance of distance vision D . For a normal, unaided human eye $D = 25\text{cm}$. If an object is brought closer than this, we cannot see it clearly. If an object is too small (like the legs of an ant), the corresponding visual angle from 25 cm is not enough to see it and if we bring it closer than that, its image on the retina is blurred. Also, the visual angle made by cosmic objects far away from us (such as stars) is too small to make out minor details and we cannot bring those closer. In such cases we need optical instruments such as a microscope in the former case and a telescope in the latter. It means that microscopes and telescopes help us in increasing the visual angle. This is called angular magnification or magnifying power.

Magnifying power: Angular magnification or magnifying power of an optical instrument is defined as the ratio of the visual angle made by the image formed by that optical instrument (β) to the visual angle subtended by the object when kept at the least distance of distinct vision (α). (Figure 9.26 (a) and 9.26 (b)) In the case of telescopes, α is the angle subtended by the object from its own position as it is not possible to get it closer.

Simple microscope or a reading glass: In order to read very small letters in a newspaper, sometimes we use a convex lens. You might have seen watch-makers using a special type of small convex lens while looking at very tiny

parts of a wrist watch. Convex lens used for this purpose is a simple microscope.

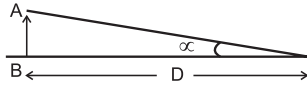


Fig. 9.26: (a) Visual Angle α .

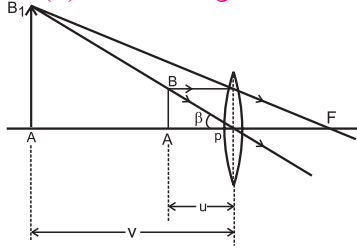


Fig. 9.26: (b) Visual Angle β .

Figure 9.26 (a) shows visual angle α made by an object, when kept at the least distance of distinct vision D . Without an optical instrument this is the greatest possible visual angle as we cannot get the object closer than this. Figure 9.26 (b) shows a convex lens forming erect, virtual and magnified image of the same object, when placed within the focus. The visual angle β of the object and the image in this case are the same. However, this time the viewer is looking at the image which is not closer than D . Hence the same object is now at a distance smaller than D . It makes β greater than α and the same object appears bigger.

Angular magnification or magnifying power, in this case, is given by

$$M = \frac{\text{Visual angle of the image}}{\text{Visual angle of the object, when at } D} = \frac{\beta}{\alpha}$$

For small angles α and β , we can write,

$$M = \frac{\beta}{\alpha} \cong \frac{\tan(\beta)}{\tan(\alpha)} = \frac{\left(\frac{BA}{PA}\right)}{\left(\frac{BA}{D}\right)} = \frac{D}{u} \quad \text{-(numerically)}$$

Limiting cases:

- (i) For maximum magnifying power, the image should be nearest possible, i.e., at D .

$$\text{For a thin lens, } \frac{1}{f} = \frac{1}{v} - \frac{1}{u}$$

In this case, $f = +f$, $v = v_{\min} = -D$, $u = -u$

and $M = M_{\max}$

$$\therefore \frac{1}{f} = \frac{1}{-D} - \frac{1}{-u} \quad \therefore \frac{D}{f} = \frac{D}{-D} + \frac{D}{u}$$

$$\therefore M_{\max} = \frac{D}{u} = 1 + \frac{D}{f}$$

- (ii) For minimum magnifying power, $v = \infty$, i.e., $u = f$ (numerically)

$$\therefore M_{\min} = \frac{D}{u} = \frac{D}{f}$$

Thus the angular magnification by a lens of focal length f is between $\left(\frac{D}{f}\right)$ and $\left(1 + \frac{D}{f}\right)$ only.

For common human eyesight, $D = 25$ cm. Thus, if $f = 5$ cm,

$$M_{\min} = \left(\frac{D}{f}\right) = 5 \text{ and } M_{\max} = \left(1 + \frac{D}{f}\right) = 6.$$

Hence image appears to be only 5 to 6 times bigger for a lens of focal length 5 cm.

For $M_{\min} = \left(\frac{D}{f}\right) = 5$, $v = \infty$. $\therefore m = \frac{v}{u} = \infty$. Thus, the image size is infinite times that of the object, but appears only 5 times bigger.

For

$$M_{\max} = 1 + \left(\frac{D}{f}\right) = 6,$$

$$v = -25 \text{ cm. Corresponding } u = \frac{-25}{6} \text{ cm}$$

$\therefore m = \frac{v}{u} = 6$. Thus, image size is 6 times that of the object, and appears also 6 times larger.

Example 9.12: A magnifying glass of focal length 10 cm is used to read letters of thickness 0.5 mm held 8 cm away from the lens. Calculate the image size. How big will the letters appear? Can you read the letters if held 5 cm away from the lens? If yes, of what size would the letters appear? If no, why not?

$$f = +10 \text{ cm}, u = -8 \text{ cm}, v = ?$$

$$\frac{1}{f} = \frac{1}{v} - \frac{1}{u} \quad \therefore \frac{1}{10} = \frac{1}{v} - \frac{1}{-8} \quad \therefore v = -40 \text{ cm}$$

$$m = \frac{v}{u} = \frac{\text{Image size } h_i}{\text{Object size } h_o} \quad \therefore \frac{40}{8} = \frac{h_i}{0.5}$$

$$\therefore h_i = 2.5 \text{ cm (5 times that of the object)}$$

$$M = \frac{D}{u} = \frac{25}{8} = 3.125$$

$\therefore$ Image will appear to be 3.125 times bigger.
i.e., $3.125 \times 0.5 = 1.5625$ cm

For $\mu = -5$ cm, v will be -10 cm.

For an average human being to see clearly, the image must be at or beyond 25 cm. Thus, it will not possible to read the letters if held 5 cm away from the lens.

Compound microscope: As seen above, the magnifying power of a simple microscope is inversely proportional to its focal length. However, if we need focal length to be smaller and smaller, the corresponding lens becomes thicker and thicker. For such a lens both spherical as well as chromatic aberrations are dominant. Thus, if higher magnifying power is needed, we go for using more than one lenses. The instrument is then called a compound microscope. It is used to view very small objects (sizes $\sim 10^{-1}$ mm to 10^{-3} mm). Also, whether the image is erect or inverted is immaterial.

A compound microscope essentially uses two convex lenses of suitable focal lengths fit into a cylindrical tube with some adjustment possible for its length. The smaller lens (~ 4 mm to 6 mm aperture) facing the object is called the objective. Other lens with which the observer jams her/his eye is litter larger and called as the eye lens. (Fig 9.27) During this discussion we consider the eye lens to be a single lens, but in practice it is an eyepiece, itself consisting of two planoconvex lenses.

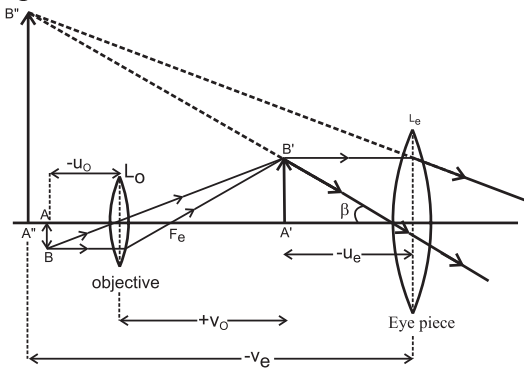


Fig. 9.27: Compound Microscope.

As shown in the Fig. 9.27, a tiny object AB is placed between f and $2f$ of the objective which produces its real, inverted and magnified image $A' B'$ in front of the eye lens. Position of the eye lens is so adjusted that the (inter-

mediate) image $A' B'$ is within its focus. Hence, for this object $A' B'$, the eye lens behaves as a simple microscope and produces its virtual and magnified image $A'' B''$, which is inverted with respect to original object AB.

Magnifying power of a compound microscope with two lenses: From its position, the final image $A'' B''$ makes a visual angle β at the eye (jammed at the eye lens). Visual angle made by the object from distance D is α .

$$\therefore \tan \beta = \frac{A'' B''}{v_e} = \frac{A' B'}{u_e}$$

$$\tan \alpha = \frac{AB}{D} \quad (\text{Fig. 9.29 (a)})$$

$\therefore$ Angular magnification or magnifying power,

$$M = \frac{\beta}{\alpha} \cong \frac{\tan \beta}{\tan \alpha} = \left(\frac{A' B'}{u_e} \right) \times \left(\frac{D}{AB} \right)$$

$$= \left(\frac{A' B'}{AB} \right) \times \left(\frac{D}{u_e} \right)$$

$$\therefore M = m_o \times M_e$$

Where, $\left(\frac{A' B'}{AB} \right) = m_o = \frac{v_o}{u_o}$ is the linear (lateral) magnification of the objective and

$\left(\frac{D}{u_e} \right) = M_e$ is the angular magnification or magnifying power of the eye lens. Length of the compound microscope then becomes $L = \text{distance between the two lenses } v_o + u_e$.

Remarks:

(i) In order to increase m_o , we need to decrease u_o . Thereby, the object comes closer and closer to the focus of the objective. This increases v_o and hence length of the microscope. Thus m_o can be increased only within the limitation of length of the microscope.

(ii) Minimum value of M_e is $\left(\frac{D}{f_e} \right)$ for final image at infinity and maximum value of M_e is $\left(1 + \frac{D}{f_e} \right)$ for final image at D

respectively. M_e and m_o together decide the minimum and maximum magnifying power of the microscope.

Example 9.13: The pocket microscope used by a student consists of eye lens of focal length 6.25 cm and objective of focal length 2 cm. At microscope length 15 cm, the final image appears biggest. Estimate distance of the object from the objective and magnifying power of the microscope.

Solution:

$$\begin{aligned}
 f_e &= 6.25 \text{ cm}, f_o = 2 \text{ cm}, L = |v_o| + |u_e| \\
 &= 15 \text{ cm}, v_e = 25 \text{ cm (Image appears largest)} \\
 \frac{1}{f_e} &= \frac{1}{v_e} - \frac{1}{u_e} \therefore \frac{1}{6.25} = \frac{4}{25} \\
 &= \frac{1}{-25} - \frac{1}{u_e} \therefore |u_e| = 5 \text{ cm} \\
 \therefore |v_o| &= L - |u_e| = 15 - 5 = 10 \text{ cm} \\
 \frac{1}{f_o} &= \frac{1}{v_o} - \frac{1}{u_o} \therefore \frac{1}{2} = \frac{1}{+10} - \frac{1}{u_o} \therefore |u_o| = 2.5 \text{ cm} \\
 M &= m_o \times M_e = \left(\frac{v_o}{u_o} \right) \left(\frac{D}{u_e} \right) \\
 &= \left(\frac{10}{2.5} \right) \left(\frac{25}{5} \right) = 4 \times 5 = 20
 \end{aligned}$$

Telescope: Telescopes are used to see terrestrial or astronomical bodies. A telescope essentially uses two lenses (or one large parabolic mirror and a lens). The lens facing the object (called objective) is of aperture as large as possible. For Newtonian telescopes, a large parabolic mirror faces the object.

For terrestrial telescopes the objects to be seen are on the Earth, like mountains, trees, players playing a match in a stadium, etc. In such case, the final image must be erect. Eye lens used for this purpose must be concave and such a telescope is popularly called a binocular. A variety of binoculars use three convex lenses with proper separation. The third lens again inverts the second intermediate image and makes final image erect with respect to the object. In this text we shall be discussing astronomical telescope.

For an astronomical telescope, the objects to be seen are planets, stars, galaxies, etc. In this case there is no necessity of erect image.

Such telescopes use convex lens as eye lens. (Fig. 9.27).

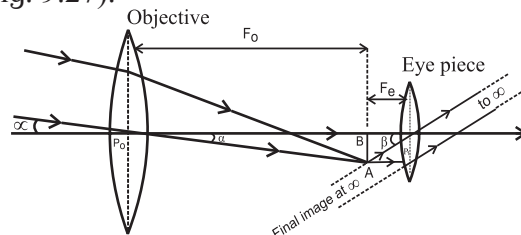


Fig. 9.28: Telescope.

Magnifying power of a telescope: Objects to be seen through a telescope cannot be brought to distance D from the objective, like in microscopes. Hence, for telescopes, α is the visual angle of the object from its own position, which is practically at infinity. Visual angle of the final image is β and its position can be adjusted to be at D . However, under normal adjustments, the final image is also at infinity but making a greater visual angle than that of the object. (If the image is really at infinity, there will not be any parallax at the cross wires). Beam of incident rays is now inclined at an angle α with the principal axis while emergent beam is inclined at a greater angle β with the principal axis causing angular magnification. (Fig. 9.28)

Objective of focal length f_o focusses the parallel incident beam at a distance f_o from the objective giving an inverted image AB. For normal adjustment, the eye lens is so adjusted that the intermediate image AB happens to be at the focus of the eye lens. Rays refracted beyond the eye lens form a parallel beam inclined at an angle β with the principal axis resulting into the image also at infinity.

$\therefore$ Angular magnification or magnifying power,

$$\begin{aligned}
 M &= \frac{\beta}{\alpha} \cong \frac{\tan \beta}{\tan \alpha} = \frac{\left(\frac{BA}{P_e B} \right)}{\left(\frac{BA}{P_o B} \right)} = \frac{\left(\frac{BA}{f_e} \right)}{\left(\frac{BA}{f_o} \right)} \\
 \therefore M &= \frac{f_o}{f_e}
 \end{aligned}$$

Length of the telescope for normal adjustment is $L = f_o + f_e$

Under the allowed limit of length objective of

maximum possible focal length f_o and eye lens of minimum possible focal length f_e can be chosen for maximum magnifying power.

Example 14: Focal length of the objective of an astronomical telescope is 1 m. Under normal adjustment, length of the telescope is 1.05 m. Calculate focal length of the eyepiece and magnifying power under normal adjustment.

Solution: For astronomical telescope,

$$L = f_o + f_e \therefore 1.05 = 1 + f_e \therefore f_e = 0.05 \text{ m} = 5 \text{ cm}$$

Under normal adjustments,

$$M = \frac{f_o}{f_e} = \frac{1}{0.05} = 20$$



Exercises

1. Choose the correct option

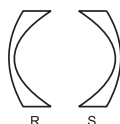
- i. As per recent understanding *light* consists of
 - (A) rays
 - (B) waves
 - (C) corpuscles
 - (D) photons obeying the rules of waves
- ii. Consider optically denser lenses P, Q, R and S drawn below. According to Cartesian sign convention which of these have positive focal length?



P



Q



R



S

- (A) Only P
 - (B) Only P and Q
 - (C) Only P and R
 - (D) Only Q and S
- iii. Two plane mirrors are inclined at angle 40° between them. Number of images seen of a tiny object kept between them is
 - (A) Only 8
 - (B) Only 9
 - (C) 8 or 9
 - (D) 9 or 10
- iv. A concave mirror of curvature 40 cm, used for *shaving purpose* produces image of double size as that of the object. Object distance must be
 - (A) 10 cm only
 - (B) 20 cm only
 - (C) 30 cm only
 - (D) 10 cm or 30 cm

- v. Which of the following aberrations will NOT occur for spherical mirrors?
 - (A) Chromatic aberration
 - (B) Coma
 - (C) Distortion
 - (D) Spherical aberration
- vi. There are different fish, monkeys and water on the habitable planet of the star Proxima b. A fish swimming underwater feels that there is a monkey at 2.5 m on the top of a tree. The same monkey feels that the fish is 1.6 m below the water surface. Interestingly, height of the tree and the depth at which the fish is swimming are exactly same. Refractive index of that water must be
 - (A) 6/5
 - (B) 5/4
 - (C) 4/3
 - (D) 7/5
- vii. Consider following phenomena/applications: P) Mirage, Q) rainbow, R) Optical fibre and S) glittering of a diamond. Total internal reflection is involved in
 - (A) Only R and S
 - (B) Only R
 - (C) Only P, R and S
 - (D) all the four
- viii. A student uses spectacles of number -2 for seeing distant objects. Commonly used lenses for her/his spectacles are
 - (A) bi-concave
 - (B) double concave
 - (C) concavo-convex
 - (D) convexo-concave

- ix. A spherical marble of refractive index 1.5 and curvature 1.5 cm, contains a tiny air bubble at its centre. Where will it appear when seen from outside?
 (A) 1 cm inside (B) at the centre
 (C) $\frac{5}{3}$ cm inside (D) 2 cm inside
- x. Select the WRONG statement.
 (A) Smaller angle of prism is recommended for greater angular dispersion.
 (B) Right angled isosceles glass prism is commonly used for total internal reflection.
 (C) Angle of deviation is practically constant for thin prisms.
 (D) For emergent ray to be possible from the second refracting surface, certain minimum angle of incidence is necessary from the first surface.
- xi. Angles of deviation for extreme colours are given for different prisms. Select the one having maximum dispersive power of its material.
 (A) $7^\circ, 10^\circ$ (B) $8^\circ, 11^\circ$
 (C) $12^\circ, 16^\circ$ (D) $10^\circ, 14^\circ$
- xii. Which of the following is not involved in formation of a rainbow?
 (A) refraction
 (B) angular dispersion
 (C) angular deviation
 (D) total internal reflection
- xiii. Consider following statements regarding a simple microscope:
 (P) It allows us to keep the object within the least distance of distant vision.
 (Q) Image appears to be biggest if the object is at the focus.
 (R) It is simply a convex lens.
 (A) Only (P) is correct
 (B) Only (P) and (Q) are correct
 (C) Only (Q) and (R) are correct
 (D) Only (P) and (R) are correct

2. Answer the following questions.

- i) As per recent development, what is the nature of light? Wave optics and particle nature of light are used to explain which phenomena of light, respectively?
- ii) Which phenomena can be satisfactorily explained using ray optics? State the assumptions on which ray optics is based.
- iii) What is focal power of a spherical mirror or of a lens? What may be the reason for using $P = \frac{1}{f}$ as its expression?
- iv) At which positions of the objects do spherical mirrors produce (i) diminished image, (ii) magnified image?
- v) State the restrictions for having images produced by spherical mirrors to be appreciably clear.
- vi) Explain spherical aberration for spherical mirrors. How can it be minimized? Can it be eliminated by some curved mirrors?
- vii) Define absolute refractive index and relative refractive index. Explain in brief, with an illustration for each.
- viii) Explain 'mirage' as an illustration of refraction.
- ix) Under what conditions is total internal reflection possible? Explain it with a suitable example. Define critical angle of incidence and obtain an expression for it.
- x) Describe construction and working of an optical fibre. What are the advantages of optical fibre communication over electronic communication?
- xi) Why is a prism binoculars preferred over traditional binoculars? Describe its working in brief.
- xii) A spherical surface separates two transparent media. Derive an expression that relates object and image distances with the radius of curvature for a point object. Clearly state the assumptions, if any.

- xiii) Derive lens makers' equation. Why is it called so? Under which conditions focal length f and radii of curvature R are numerically equal for a lens?

2. Answer the following questions in detail.

- i) What are different types of dispersions of light? Why do they occur?
- ii) Define angular dispersion for a prism. Obtain its expression for a thin prism. Relate it with the refractive indices of the material of the prism for corresponding colours.
- iii) Explain and define dispersive power of a transparent material. Obtain its expressions in terms of angles of deviation and refractive indices.
- iv) (i) State the conditions under which a rainbow can be seen.
(ii) Explain the formation of a primary rainbow. For which angular range with the horizontal is it visible?
(iii) Explain the formation of a secondary rainbow. For which angular range with the horizontal is it visible?
(iv) Is it possible to see primary and secondary rainbow simultaneously? Under what conditions?
- v) (i) Explain chromatic aberration for spherical lenses. State a method to minimize or eliminate it.
(ii) What is achromatism? Derive a condition to achieve achromatism for a lens combination. State the conditions for it to be converging.
- vi) Describe spherical aberration for spherical lenses. What are different ways to minimize or eliminate it?
- vii) Define and describe magnifying power of an optical instrument. How does it differ from linear or lateral magnification?
- viii) Derive an expression for magnifying power of a simple microscope. Obtain its minimum and maximum values in terms of its focal length.

- ix) Derive the expressions for the magnifying power and the length of a compound microscope using two convex lenses.

- x) What is a terrestrial telescope and an astronomical telescope?

- xi) Obtain the expressions for magnifying power and the length of an astronomical telescope under normal adjustments.

- xii) What are the limitations in increasing the magnifying powers of (i) simple microscope (ii) compound microscope (iii) astronomical telescope?

3. Solve the following numerical examples

- i) A monochromatic ray of light strike the water ($n = 4/3$) surface in a cylindrical vessel at angle of incidence 53° . Depth of water is 36 cm. After striking the water surface, how long will the light take to reach the bottom of the vessel? [Angles of the most popular Pythagorean triangle of sides in the ratio 3:4:5 are nearly 37° , 53° and 90°]

[Ans: 2 ns]

- ii) Estimate the number of images produced if a tiny object is kept in between two plane mirrors inclined at 35° , 36° , 40° and 45° .

[Ans: 10, 9, 9 or 8, 7 respectively]

- iii) A rectangular sheet of length 30 cm and breadth 3 cm is kept on the principal axis of a concave mirror of focal length 30 cm. Draw the image formed by the mirror on the same ray diagram, as far as possible on scale.

[Ans: Inverted image starts from 50 cm and ends at 90 cm. Its height in the beginning is 2 cm and at the end it is 6 cm. At 60 cm, image height is 3 cm. Thus, outer boundary if the image is a curve]

- iv) A car uses a convex mirror of curvature 1.2 m as its rear-view mirror. A minibus of cross section $2.4 \text{ m} \times 2.4 \text{ m}$ is 6.6 m away from the mirror. Estimate the image size.

[Ans: A square of edge 0.2 m]

- v) A glass slab of thickness 2.5 cm having refractive index $5/3$ is kept on an ink spot. A transparent beaker of very thin bottom, containing water of refractive index $4/3$ up to 8 cm, is kept on the glass block. Calculate apparent depth of the ink spot when seen from the outside air.

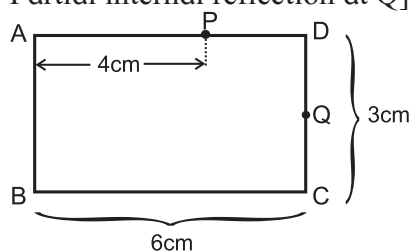
[Ans: 7.5 cm]

- vi) A convex lens held some distance above a 6 cm long pencil produces its image of SOME size. On shifting the lens by a distance equal to its focal length, it again produces the image of the SAME size as earlier. Determine the image size.

[Ans: 12 cm]

- vii) Figure below shows the section ABCD of a transparent slab. There is a tiny green LED light source at the bottom left corner B. A certain ray of light from B suffers total internal reflection at nearest point P on the surface AD and strikes the surface CD at point Q. Determine refractive index of the material of the slab and distance DQ. At Q, the ray PQ will suffer partial or total internal reflection? [You may use the approximation given in Q 1 above].

[Ans: $n = 5/4$, $DQ = 1.5$ cm, Partial internal reflection at Q]



- viii) A point object is kept 10 cm away from one of the surfaces of a thick double convex lens of refractive index 1.5 and radii of curvature 10 cm and 8 cm. Central thickness of the lens is 2 cm. Determine location of the final image considering paraxial rays only.

Hint : Single spherical surface formula to be used twice.

[Ans: 64 cm away from the other surface]

- ix) A monochromatic ray of light is incident at 37° on an equilateral prism of refractive index $3/2$. Determine angle of emergence and angle of deviation. If angle of prism is adjustable, what should its value be for emergent ray to be just possible for the same angle of incidence.

[Ans: $e = 63^\circ$, $\delta = 40^\circ$, $A = 65^\circ 24'$ for $e = 90^\circ$ (just emerges)]

- x) From the given data set, determine angular dispersion by the prism and dispersive power of its material for extreme colours. $n_R = 1.62$, $n_V = 1.66$, $\delta_R = 3.1^\circ$

[Ans: $\delta_{VR} = 0.2^\circ$, $\omega_{VR} = \frac{1}{16} = 0.0625$]

- xi) Refractive index of a flint glass varies from 1.60 to 1.66 for visible range. Radii of curvature of a thin convex lens are 10 cm and 15 cm. Calculate the chromatic aberration between extreme colours.

[Ans: 10/11 cm]

- xii) A person uses spectacles of 'number' 2.00 for reading. Determine the range of magnifying power (angular magnification) possible. It is a concavo-convex lens ($n = 1.5$) having curvature of one of its surfaces to be 10 cm. Estimate that of the other.

[Ans: $M_{\min} = 0.5$, $M_{\max} = 1.5$, $R_2 = 50/3$ cm]

- xiii) Focal power of the eye lens of a compound microscope is 6 dioptre. The microscope is to be used for maximum magnifying power (angular magnification) of at least 12.5. The packing instructions demand that length of the microscope should be 25 cm. Determine minimum focal power of the objective. How much will its radius of curvature be if it is a biconvex lens of $n = 1.5$.

[Ans: 40 dioptre, 2.5 cm]



Can you recall?

1. Have you experienced a shock while getting up from a plastic chair and shaking hand with your friend?
2. Ever heard a crackling sound while taking out your sweater in winter?
3. Have you seen the lightning striking during pre-monsoon weather?

10.1 Introduction:

Electrostatics deals with static electric charges, the forces between them and the effects produced in the form of electric fields and electric potentials. We have already studied some aspects of electrostatics in earlier standards. In this Chapter we will review some of them and then go on to study some aspects in details.

Current electricity, which plays a major role in our day to day life, is produced by moving charges. Charges are present everywhere around us though their presence can only be felt under special circumstances. For example, when we remove our sweater in winter on a dry day, we hear some crackling sound and the sweater appears to stick to our body. This is because of the electric charges produced due to friction between our body and the sweater. Similarly, the lightening that we see in the sky is also due to the flow of large amount of electric charges that develop on the clouds due to friction.

10.2 Electric Charges:

Historically, opposite electric charges were known to the Greeks in the 600 BC. They realized that equal and opposite charges develop on amber and fur when rubbed against each other. Now we know that electric charge is a basic property of elementary particles of which matter is made of. These elementary particles are proton, neutron, and electron. Atoms are made of these particles and matter is made from atoms. A proton is considered to be positively charged and electron to be negatively charged. Neutron is electrically neutral, i.e., it has no charge. An atomic nucleus is made up of protons and neutrons and hence is positively charged. Negatively charged electrons surround

the nucleus so as to make an atom electrically neutral. Thus, most matter around us is electrically neutral.

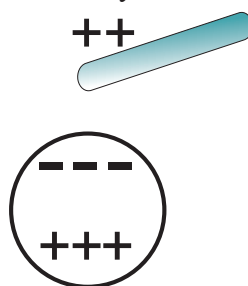


Fig. 10.1 a

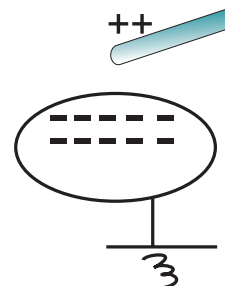


Fig. 10.1 b

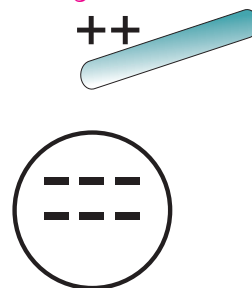


Fig. 10.1 c

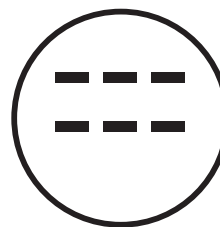


Fig. 10.1 d

Fig. 10.1 (a): Insulated conductor

Fig. 10.1 (b): +ve charge is neutralized by electron from Earth

Fig. 10.1 (c): Earthing is removed -ve charge still stays on the conductor due to +ve charged rod

Fig. 10.1 (d): Rod removed -ve charge is distributed over the surface of the conductor

When certain dissimilar substances, like fur and amber or comb and dry hair are rubbed against each other, electrons get transferred to the other substance making them charged. The substance receiving electrons develops a negative charge while the other is left with an equal amount of positive charge. This can be called charging by conduction as charges are transferred from one body to another. Charges can be separated by other means as well, like

chemical reactions (in cells), convection (in clouds), diffusion (in living cells) etc.

If an uncharged conductor is brought near a charged body, (not in physical contact) the nearer side of the conductor develops opposite charge to that on the charged body and the far side of the conductor develops charge similar to that on the charged body. This is called induction. This happens because the electrons in a conductor are free and can move easily in presence of a charged body. This can be seen from Fig. 10.1.

A charged body attracts or repels electrons in a conductor depending on whether the charge on the body is positive or negative respectively. Positive and negative charges are redistributed and are accumulated at the ends of the conductor near and away from the charged body. From the above discussion it can be inferred that there are only two types of charges found in nature, namely, positive and negative charges. In induction, there is no transfer of charges between the charged body and the conductor. So when the charged body is moved away from the conductor, the charges in the conductor are free again.



Can you tell?

1. When a petrol or a diesel tanker is emptied in a tank, it is grounded.
2. A thick chain hangs from a petrol or a diesel tanker and it is in contact with ground when the tanker is moving.

10.3 Basic Properties of Electric Charge:

10.3.1 Additive Nature of Charge:

Electric charge is additive, similar to mass. The total electric charge on an object is equal to the algebraic sum of all the electric charges distributed on different parts of the object.

It may be pointed out that while taking the algebraic sum, the sign (positive or negative) of the electric charges must be taken into account. Thus if two bodies have equal and opposite charges, the net charge on the system of the two bodies is zero. This is similar to that in case of atoms where the nucleus is positively charged and this charge is equal to the negative charge

Gold Leaf Electroscope:

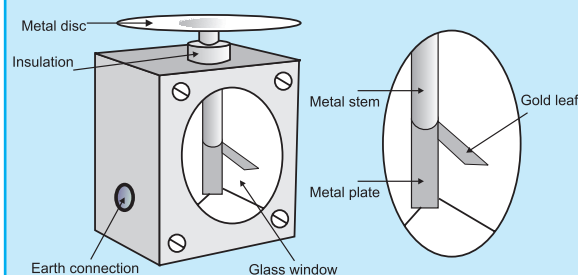
This is a classic instrument for detecting presences of electric charge. A metal disc is connected to one end of a narrow metal rod and a thin piece of gold leaf is fixed to the other end. The whole of this part of the electroscope is insulated from the body of the instrument. A glass front prevents air draughts but allows to observe the effect of charge on the leaf.

When a charge is put on the disc at the top it spreads down to the plate and leaf moves away from the plate. This happens because similar charges repel. The more the charge on the disc, more is the separation of the leaf from the plate.

The leaf can be made to fall again by touching the disc. This is done by earthing the electroscope. An earth terminal prevents the case from accumulating any stray charge. The electroscope can be charged in two ways.

- (a) by contact- a charged rod is brought in contact with the disc and charge is transferred to the electroscope. This method gives the gold leaf the same charge as that on the conductor. This is not a very effective method of charging the electroscope.
- (b) by induction- a charged rod is brought close to the disc (not touching it) and the electroscope is earthed. The rod is then removed. This method give the gold leaf opposite charges.

The following diagrams show how the charges spread to the gold leaf and lift it.



of the electrons making the atoms electrically neutral.

It is interesting to compare the additive property of charge with that of mass.

- 1) The masses of the particles constituting an object are always positive, whereas the charges distributed on different parts of the object may be positive or negative.
- 2) The total mass of an object is always positive whereas, the total charge on the object may be positive, zero or negative.

10.3.2 Quantization of Charge:

The minimum value of the charge on an electron as determined by the Milikan's oil drop experiment is $e = 1.6 \times 10^{-19} \text{ C}$. This is called the elementary charge. Here, C stands for coulomb which is the unit of charge in SI system. Unit of charge is defined in article 10.4.3. Since protons (+ve) and electrons (-ve) are the charged particles constituting matter, the charge on an object must be an integral multiple of $\pm e$. $q = \pm ne$, where n is an integer.

Further, charge on an object can be increased or decreased in multiples of e . It is because, during the charging process an integral number of electrons can be transferred from one body to the other body. This is known as **quantization of charge** or discrete nature of charge.

The discrete nature of electric charge is usually not observable in practice. It is because the magnitude of the elementary electric charge, e , is extremely small. Due to this, the number of elementary charges involved in charging an object becomes extremely large. Suppose, for example, when a glass rod is rubbed with silk, a charge of the order of one μC (10^{-6} C) appears on the glass rod or silk. Since elementary charge $e = 1.6 \times 10^{-19} \text{ C}$, the number of elementary charges on the glass rod (or silk) is given by

$$n = \frac{10^{-6} \text{ C}}{1.6 \times 10^{-19} \text{ C}} = 6.25 \times 10^{12}$$

Since it is a tremendously large number, the quantization of charge is not observed and one usually observes a continuous variation of charge.

Example 10.1: How much positive and negative charge is present in 1gm of water? How many electrons are present in it? Given, molecular mass of water is 18.0 g.

Solution: Molecular mass of water is 18.0 gm, that means the number of molecules in 18.0 gm of water is 6.02×10^{23} .

$\therefore$ Number of molecules in 1gm of water $= 6.02 \times 10^{23} / 18$. One molecule of water (H_2O) contains two hydrogen atoms and one oxygen atom. Thus the number of electrons in H_2O is sum of the number of electrons in H_2 and oxygen. There are 2 electrons in H_2 and 8 electrons oxygen.

$\therefore$ Number of electrons in $\text{H}_2\text{O} = 2 + 8 = 10$.

Total number of protons / electrons in 1.0 gm of water $= \frac{6.02 \times 10^{23}}{18} \times 10 = 3.34 \times 10^{23}$

Total positive charge $= 3.34 \times 10^{23} \times \text{charge on a proton}$

$$= 3.34 \times 10^{23} \times 1.6 \times 10^{-19} \text{ C} = 5.35 \times 10^4 \text{ C}$$

This positive charge is balanced by equal amount of negative charge so that the water molecule is electrically neutral.



Do you know ?

According to recent advancement in physics, it is now believed that protons and neutrons are themselves built out of more elementary units called quarks. They are of six types, having fractional charge $(-1/3)e$ or $(+2/3)e$. A proton or a neutron consists of a combination of three quarks. It may be clearly understood that even in the quark model, quantization of charge is not affected. It is only the step size of the charge that decreases from e to $e/3$. Quarks are always present in bound states and no free quarks are known to exist. In modern day experiments it is possible to observe the discrete nature of charge in very sensitive devices such as single electron transistor.

10.3.3 Conservation of Charge:

We know that when a glass rod is rubbed with silk, it becomes positively charged and silk becomes negatively charged. The amount

of positive charge on glass rod is found to be exactly the same as negative charge on silk. Thus, the systems of glass rod and silk together possesses zero net charge after rubbing.

Result and conclusion of this experiment can be generalized and we can say that "in any given physical process, charge may get transferred from one part of the system to another, but the total charge in the system remains constant" or, **for an isolated system total charge cannot be created nor destroyed**. In simple words, the total charge of an isolated system is always conserved.

10.3.4 Forces between Charges:

It was observed in carefully conducted experiments with charged objects that they experience force when brought close (not touching) to each other. This force can be attractive or repulsive. **Like charges repel each other and unlike charges attract each other**. Figure 10.2 describes this schematically. This is the reason for charging by induction as described in section 10.2 and Fig. 10.1.

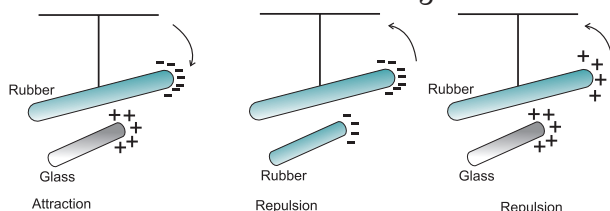


Fig. 10.2: Attractive and repulsive force.

10.4 Coulomb's Law:

The electric interaction between two charged bodies can be expressed in terms of the forces they exert on each other. Coulomb (1736-1806) made the first quantitative investigation of the force between electric charges. He used point charges at rest to study the interaction. A point charge is a charge whose dimensions are negligibly small compared to its distance from another bodies. **Coulomb's law is a fundamental law governing interaction between charges at rest.**

10.4.1 Scalar form of Coulomb's Law:

Statement : *The force of attraction or repulsion between two point charges at rest is directly proportional to the product of the magnitude of the charges and inversely*

proportional to the square of the distance between them. This force acts along the line joining the two charges.

Let q_1 and q_2 be two point charges at rest with respect to each other and separated by a distance r . The magnitude F of the force between them is given by,

$$F \propto \frac{q_1 q_2}{r^2}$$

$$F = K \frac{q_1 q_2}{r^2} \quad \text{--- (10.1)}$$

where K is the constant of proportionality. Its magnitude depends on the units in which F , q_1 , q_2 and r are expressed and also on the properties of the medium around the charges.

The force between the two charges will be attractive if they are unlike (one positive and one negative). The force will be repulsive if charges are similar (both positive or both negative). Figure 10.3 describes this schematically.

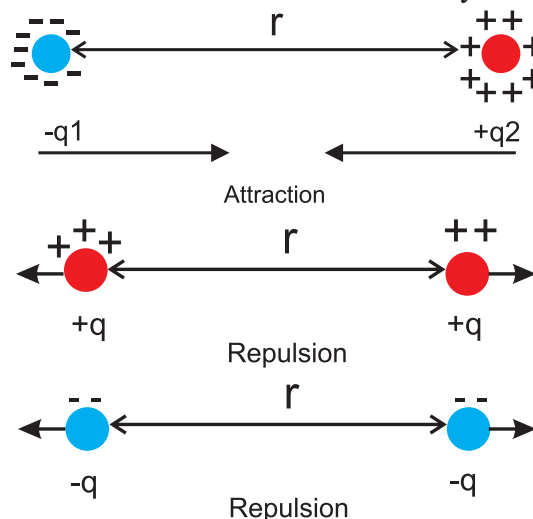


Fig. 10.3: Coulomb's law.

10.4.2 Relative Permittivity or Dielectric Constant:

While discussing the coulomb's law it was assumed that the charges are held stationary in vacuum. When the charges are kept in a material medium, such as water, mica or parafined paper, the medium affects the force between the charges. The force between the two charges placed in a medium may be written as,

$$F_{\text{med}} = \frac{1}{4\pi\epsilon} \left(\frac{q_1 q_2}{r^2} \right) \quad \text{--- (10.2)}$$

where ϵ is called the absolute permittivity of the medium. The force between the same two charges placed in free space or vacuum at distance r is given by,

$$F_{\text{vac}} = \frac{1}{4\pi\epsilon_0} \left(\frac{q_1 q_2}{r^2} \right) \quad \text{--- (10.3)}$$

Dividing Eq. (10.3) by (10.2)

$$\frac{F_{\text{vac}}}{F_{\text{med}}} = \frac{\frac{1}{4\pi\epsilon_0} \left(\frac{q_1 q_2}{r^2} \right)}{\frac{1}{4\pi\epsilon} \left(\frac{q_1 q_2}{r^2} \right)} = \frac{\epsilon}{\epsilon_0}$$

The ratio $\frac{\epsilon}{\epsilon_0}$ is the relative permittivity or dielectric constant of the medium and is denoted by ϵ_r or K .

$$K \text{ or } \epsilon_r = \frac{\epsilon}{\epsilon_0} = \frac{F_{\text{vac}}}{F_{\text{med}}} \quad \text{--- (10.4)}$$

Thus,

- (i) ϵ_r is the ratio of absolute permittivity of a medium to the permittivity of free space.
- (ii) ϵ_r is the ratio of the force between two point charges placed a certain distance apart in free space or vacuum to the force between the same two point charges when placed at the same distance in the given medium.

ϵ_r is a dimensionless quantity.

- (iii) ϵ_r is also called specific inductive capacity.

The force between two point charges q_1 and q_2 placed at a distance r in a medium of relative permittivity ϵ_r is given by

$$F = \frac{1}{4\pi\epsilon_0\epsilon_r} \frac{q_1 q_2}{r^2} \quad \text{--- (10.5)}$$

For water, $\epsilon_r = 80$ then from Eq. (10.4)

$$\frac{F_{\text{vac}}}{F_{\text{water}}} = \epsilon_r = 80$$

$$F_{\text{water}} = \frac{F_{\text{vac}}}{80}$$

This means that when two point charges are placed some distance apart in water, the force between them is reduced to $\left(\frac{1}{80}\right)^{\text{th}}$ of the

force between the same two charges placed at the same distance in vacuum.

Thus, a material medium reduces the

force between charges by a factor of ϵ_r , its relative permittivity.

While using Eq. (10.5) we assume that the medium is homogeneous, isotropic and infinitely large.

10.4.3 Definition of Unit Charge from the Coulomb's Law:

The force between two point charges q_1 and q_2 , separated by a distance r in free space, is written by using Eq. (10.2),

$$F = \frac{1}{4\pi\epsilon_0} \times \frac{q_1 q_2}{r^2} = 9 \times 10^9 \times \frac{q_1 q_2}{r^2}$$

$$\epsilon_0 = 8.85 \times 10^{-12} \text{ C}^2 \text{ N}^{-1} \text{ m}^{-2}$$

$$\text{If } q_1 = q_2 = 1 \text{ C and } r = 1.0 \text{ m}$$

$$\text{Then } F = 9.0 \times 10^9 \text{ N}$$

From this, we define, coulomb (C) the unit of charge in SI units.

One coulomb is the amount of charge which, when placed at a distance of one metre from another charge of the same magnitude in vacuum, experiences a force of $9.0 \times 10^9 \text{ N}$. This force is a tremendously large force realisable in practical situations. It is, therefore, necessary to express the charge in smaller units for practical purpose. Subunits of coulomb are used in electrostatics. For example, micro-coulomb (10^{-6} C , μC), nano-coulomb (10^{-9} C , nC) or pico-coulomb. (10^{-12} C , pC) are normally used units.



Do you know ?

Force between two charges of 1.0 C each, separated by a distance of 1.0 m is $9.0 \times 10^9 \text{ N}$ or, about 10 million metric tonne. A normal truck-load is about 10 metric tonne. So, this force is equivalent to about one million truck-loads. A tremendously large force indeed !

Example 10.2: Charge on an electron is $1.6 \times 10^{-19} \text{ C}$. How many electrons are required to accumulate a charge of one coulomb?

Solution: $1.6 \times 10^{-19} \text{ C} = 1 \text{ electron}$

$$\therefore 1 \text{ C} = \frac{1}{1.6 \times 10^{-19}} \text{ electrons}$$

$$= 0.625 \times 10^{19} = 6.25 \times 10^{18} \text{ electrons}$$

6.25×10^{18} electrons are required to accumulate a charge of one coulomb.

It is now possible to measure a very small amount of current in otto-amperes which measures flow of single electron.

10.4.4 Coulomb's Law in Vector Form:

As shown in Fig 10.4, q_1 and q_2 are two similar point charges situated at points A and B. r_{12} is the distance of separation between them.

$\vec{F}_{21}$ denotes the force exerted on q_2 by q_1

$$\vec{F}_{21} = \frac{1}{4\pi\epsilon_0} \times \frac{q_1 q_2}{|\hat{r}_{21}|^2} \times \hat{r}_{21} \quad \text{--- (10.6)}$$

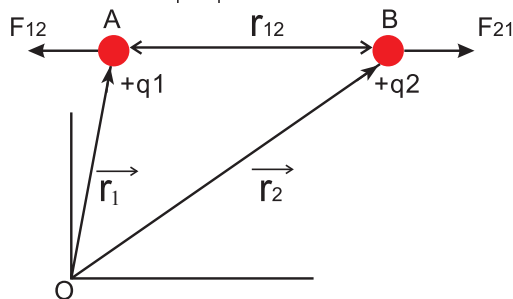


Fig. 10.4: Coulomb's law in vector form

$\hat{r}_{21}$ is the unit vector along $\overline{AB}$, away from B. Similarly, the force $\vec{F}_{12}$ exerted on q_1 by q_2 is given by

$$\vec{F}_{12} = \frac{1}{4\pi\epsilon_0} \times \frac{q_1 q_2}{|\hat{r}_{12}|^2} \times \hat{r}_{12} \quad \text{--- (10.7)}$$

$\hat{r}_{12}$ is the unit vector along $\overline{BA}$, away from A. $\vec{F}_{12}$ acts on q_1 at A and is directed along BA, away from A. The unit vectors $\hat{r}_{12}$ and $\hat{r}_{21}$ are oppositely directed i.e., $\hat{r}_{12} = -\hat{r}_{21}$ hence,

$$\vec{F}_{21} = -\vec{F}_{12}$$

Thus, the two charges experience force of equal magnitude and opposite in direction. These two forces form an action- reaction pair. As $\vec{F}_{21}$ and $\vec{F}_{12}$ act along the line joining the two charges, the electrostatic force is a central force.

Example 10.3: Calculate and compare the electrostatic and gravitational forces between two protons which are 10^{-15} m apart. Value of $G = 6.674 \times 10^{-11} \text{ m}^3 \text{ kg}^{-1} \text{ s}^{-2}$ and mass of the proton is $1.67 \times 10^{-27} \text{ kg}$

Solution: The electrostatic force between the protons is given by $F_e = \frac{1}{4\pi\epsilon_0} \frac{q_1 q_2}{r^2}$

Here, $q_1 = q_2 = +1.6 \times 10^{-19} \text{ C}$, $r = 10^{-15} \text{ m}$

$$\begin{aligned} \therefore F_e &= 9 \times 10^9 \frac{(1.6 \times 10^{-19})(1.6 \times 10^{-19})}{(10^{-15})^2} \\ &= 9 \times 1.6 \times 1.6 \times 10^1 \text{ N} \\ F_e &= 2.3 \times 10^2 \text{ N} \quad \text{--- (10.8.a)} \end{aligned}$$

The gravitational force between the protons is given by

$$\begin{aligned} F_g &= G \frac{m_1 m_2}{r^2} \\ &= \frac{6.674 \times 10^{-11} \times 1.67 \times 10^{-27} \times 1.67 \times 10^{-27}}{(10^{-15})^2} \\ F_g &= 1.86 \times 10^{-34} \text{ N} \quad \text{--- (10.8.b)} \end{aligned}$$

Comparing 10.8. (a) and 10.8.(b)

$$\frac{F_e}{F_g} = \frac{2.30 \times 10^{-2} \text{ N}}{1.86 \times 10^{-34} \text{ N}} = 1.23 \times 10^{36}$$

Thus, the electrostatic force is about 36 orders of magnitude stronger than the gravitational force.

Comparison of gravitational and electrostatic forces:

Similarities

- Both forces obey inverse square law : $F \propto \frac{1}{r^2}$
- Both are central forces : act along the line joining the two objects.

Differences

- Gravitational force between two objects is always attractive while electrostatic force between two charges can be either attractive or repulsive depending on the nature of charges.
- Gravitational force is about 36 orders of magnitude weaker than the electrostatic force.

10.5 Principle of Superposition:

The principle of superposition states that when a number of charges are interacting, the resultant force on a particular charge is given by the vector sum of the forces exerted by individual charges.

Consider a number of point charges q_1, q_2, q_3 ----- kept at points A_1, A_2, A_3 --- as shown in Fig. 10.5. The force exerted on the charge q_1 by q_2 is $\vec{F}_{12}$. The value of $\vec{F}_{12}$ is calculated

by ignoring the presence of other charges. Similarly, we find $\vec{F}_{13}$, $\vec{F}_{14}$ etc, one at a time, using the coulomb's law.

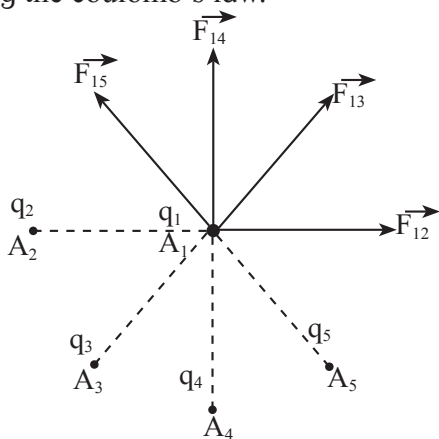


Fig. 10.5: Principle of superposition.

Total force $\vec{F}_1$ on charge q_1 is the vector sum of all such forces.

$$\vec{F}_1 = \vec{F}_{12} + \vec{F}_{13} + \vec{F}_{14} + \dots$$

$$= \frac{1}{4\pi\epsilon_0} \left[\frac{q_1 q_2}{r_{12}^2} \hat{r}_{12} + \frac{q_1 q_3}{r_{13}^2} \hat{r}_{13} + \dots \right],$$

where $\hat{r}_{12}$, $\hat{r}_{13}$ etc., are unit vectors directed to q_1 from q_2 , q_3 etc., and r_{12} , r_{13} , r_{14} , etc., are the distances from q_1 to q_2 , q_3 etc respectively.

Let there be N point charges q_1 , q_2 , q_3 etc., q_N . The force $\vec{F}$ exerted by these charges on a test charge q_0 can be written using the summation notation Σ as follows,

$$\vec{F}_{\text{test}} = \vec{F}_1 + \vec{F}_2 + \vec{F}_3 + \dots + \vec{F}_N \quad \text{--- (10.9)}$$

$$= \sum_{n=1}^N \vec{F}_n = \frac{1}{4\pi\epsilon_0} \sum_{n=1}^N \frac{q_0 q_n}{r_{0n}^2} \hat{r}_{0n} \quad \text{--- (10.10)}$$

Where $\hat{r}_{0n}$ is a unit vector directed from the n^{th} charge to the test charge q_0 and r_{0n} is the separation between them, $\vec{r}_{0n} = r_{0n} \hat{r}_{0n}$



Can you tell?

Three charges, q each, are placed at the vertices of an equilateral triangle. What will be the resultant force on charge q placed at the centroid of the triangle?

Example 10.4: Three charges of $2\mu\text{C}$, $3\mu\text{C}$ and $4\mu\text{C}$ are placed at points A, B and C respectively, as shown in Fig. a. Determine the force on A due to other charges.

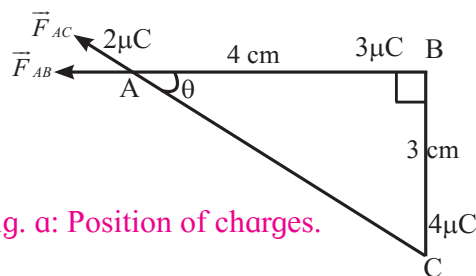


Fig. a: Position of charges.

Solution : Given,

$$AB = 4.0 \text{ cm}, BC = 3.0 \text{ cm}$$

$$\therefore AC = \sqrt{4^2 + 3^2} = 5.0 \text{ cm}$$

Magnitude of force $\vec{F}_{AB}$ on A due to B is,

$$\begin{aligned} \vec{F}_{AB} &= \left(\frac{1}{4\pi\epsilon_0} \right) \frac{2 \times 10^{-6} \times 3 \times 10^{-6}}{(4 \times 10^{-2})^2} \\ &= \frac{9 \times 10^9 \times 6}{16 \times 10^{-4}} \times 10^{-12} \\ &= 3.37 \times 10 \\ &= 33.7 \text{ N} \end{aligned}$$

This force acts at point A and is directed along $\vec{BA}$ (Fig. (b)).

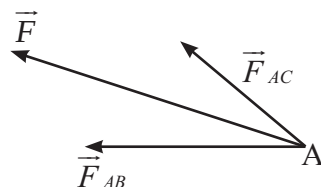


Fig. b: Forces acting at point A.

Magnitude of force $\vec{F}_{AC}$ on A due to C is,

$$\begin{aligned} \vec{F}_{AC} &= \left(\frac{1}{4\pi\epsilon_0} \right) \frac{2 \times 10^{-6} \times 4 \times 10^{-6}}{(5 \times 10^{-2})^2} \\ &= \frac{9 \times 10^9 \times 8.0 \times 10^{-12}}{25 \times 10^{-4}} \\ &= \frac{72}{25} \times 10 = 28.8 \text{ N} \end{aligned}$$

This force acts at point A and is directed along $\vec{CA}$. (Fig. 10.6.(b))

$$\vec{F} = \vec{F}_{AB} + \vec{F}_{AC}$$

Magnitude of resultant force is,

$$\begin{aligned} F &= \left[F_{AC}^2 + F_{AB}^2 + 2F_{AC} \cdot F_{AB} \cdot \cos \theta \right]^{1/2} \\ &= 59.3 \text{ N} \end{aligned}$$

Direction of the resultant force is 16.9° north of west. (Fig. c)

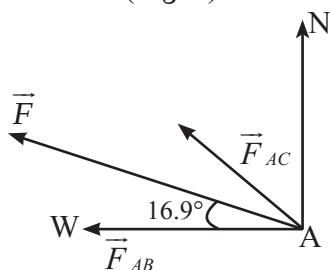


Fig. c: Direction of the resultant force.

10.6 Electric Field:

Space around a charge Q gets modified so that when a test charge is brought in this region, it experiences a coulomb force. *This region around a charged object in which coulomb force is experienced by another charge is called electric field.*

Mathematically, electric field is defined as the force experienced per unit charge. Let Q and q be two charges separated by a distance r .

The coulomb force between them is given by $\vec{F} = \frac{1}{4\pi\epsilon_0} \frac{Qq}{r^2} \hat{r}$, where, $\hat{r}$ is the unit vector

along the line joining Q to q .

Therefore, electric field due to charge Q is given by,

$$\vec{E} = \frac{\vec{F}}{q} = \frac{Q}{4\pi\epsilon_0 r^2} \hat{r} \quad \text{--- (10.11)}$$

The coulomb force acts across an empty space (vacuum) and does not need any intervening medium for its transmission.

The electric field exists around a charge irrespective of the presence of other charges.

Since the coulomb force is a vector, the

A precise definition of electric field is: Electric field is the force experienced by a test charge in presence of the given charge at the given distance from it.

$$E = \lim_{q \rightarrow 0} \frac{\vec{F}}{q}$$

Test charge is a positive charge so small in magnitude that it does not affect the surroundings of the given charge.

electric field of a charge is also a vector and is directed along the direction of the coulomb force, experienced by a test charge.

The magnitude of electric field at a distance r from a point charge Q is same at all points on the surface of a sphere of radius r as shown in Fig. 10.6. Its direction is along the radius of the sphere, pointing away from its centre if the charge is positive.

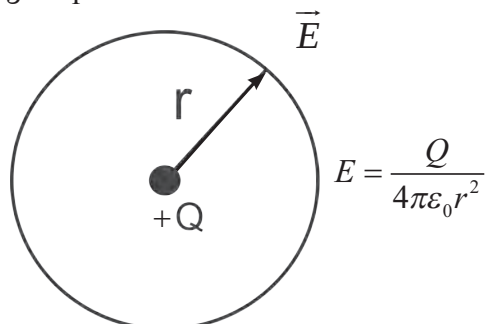


Fig. 10.6: Electric field due to a point charge (+Q).

SI unit of electric intensity is newton per coulomb (NC^{-1}). Practically, electric field is expressed in volt per metre (Vm^{-1}). This is discussed in article 10.6.2.

Dimensional formula of E is,

$$E = \frac{F}{q_0} = \frac{[\text{LMT}^{-2}]}{[\text{IT}]}$$

$$E = [\text{LMT}^{-3} \text{I}^{-1}]$$

10.6.1 Electric Field Intensity due to a Point Charge in a Material Medium:

Consider a point charge q placed at point O in a medium of dielectric constant K as shown in Fig. 10. 7.

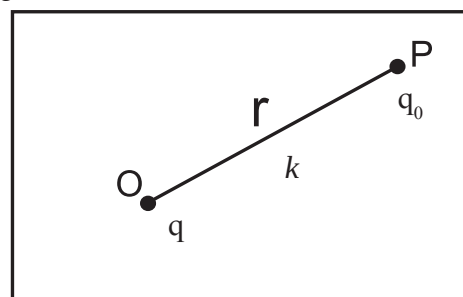


Fig. 10.7: Field in a material medium.

Consider the point P in the electric field of point charge q at distance r from it. A test charge q_0 placed at the point P will experience a force which is given by the Coulomb's law,

$$\vec{F} = \frac{1}{4\pi\epsilon_0 K} \frac{q q_0}{r^2} \hat{r}$$

where $\hat{r}$ is the unit vector in the direction of force i.e., along OP.

By the definition of electric field intensity

$$\vec{E} = \frac{\vec{F}}{q_0} = \frac{1}{4\pi\epsilon_0 K} \frac{q}{r^2} \hat{r}$$

The direction of $\vec{E}$ will be along OP when q is positive and along PO when q is negative.

The magnitude of electric field intensity in a medium is given by

$$E = \frac{1}{4\pi\epsilon_0 K} \frac{q}{r^2} \quad \text{--- (10.12)}$$

For air or vacuum $K = 1$ then

$$E = \frac{1}{4\pi\epsilon_0} \frac{q}{r^2}$$

The coulomb force between two charges and electric field E of a charge both follow the inverse square law, ($F \propto 1/r^2$, $E \propto 1/r^2$) Fig. 10.8.

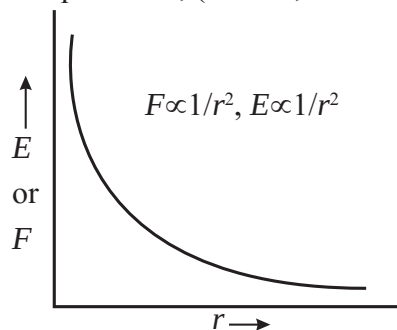


Fig. 10.8: Variation of Coulomb force/ Electric field due to a point charge.

- 1. Uniform electric field:** A uniform electric field is a field whose magnitude and direction is same at all points. For example, field between two parallel plates. Fig 10.9.a
- 2. Non uniform electric field:** A field whose magnitude and direction is not the same at all points. For example, field due to a point charge. In this case, the magnitude of field is same at distance r from the point charge in any direction but the direction of the field is not same. Fig 10.9.b

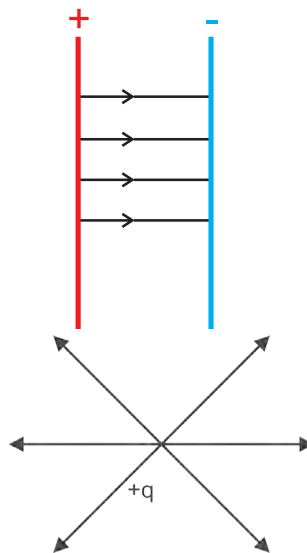


Fig. 10.9 (a): uniform electric field.

Fig. 10.9 (b): non uniform electric field.

10.6.2 Practical Way of Calculating Electric Field

A pair of charged parallel plates is arranged as shown in Fig. 10.10. The electric field between them is uniform. A potential difference V is applied between two parallel plates separated by a distance ' d '. The electric field between them is directed from plate A to plate B as shown.

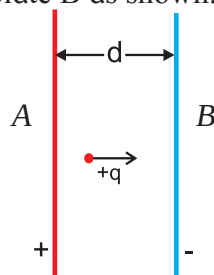


Fig 10.10: Electric field between two parallel plates.

A charge $+q$ placed between the plates experiences a force F due to the electric field. If we have to move the charge against the direction of field, i.e., towards the positive plate, we have to do some work on it. If we move the charge $+q$ from the negative plate B to the positive plate A, the work done against the field is $W = Fd$; where ' d ' is the separation between the plates. The potential difference V between the two plates is given by

$$W = Vq, \text{ but } W = Fd$$

$$\therefore Vq = Fd \therefore F/q = V/d = E$$

$\therefore$ Electric field can be defined as

$$E = V/d \quad \text{--- (10.13)}$$

This is the commonly used definition of electric field.

Example 10.5: Gap between two electrodes of the spark-plug used in an automobile engine is 1.25 mm. If the potential of 20 V is applied across the gap, what will be the magnitude of electric field between the electrodes?

Solution:

$$E = \frac{V}{d}$$

$$E = \frac{20V}{1.25 \times 10^{-3} m} = 16 \times 10^{-3} = 1.6 \times 10^4 \frac{V}{m}$$

This electric field is sufficient to ionize the gaseous mixture of fuel compressed in the cylinder and ignite it.



Can you tell?

Why a small voltage can produce a reasonably large electric field?

Example 10.6: Three point charges are placed at the vertices of a right isosceles triangle as shown in the Fig. a. What is the magnitude and direction of the resultant electric field at point P which is the mid point of its hypotenuse?

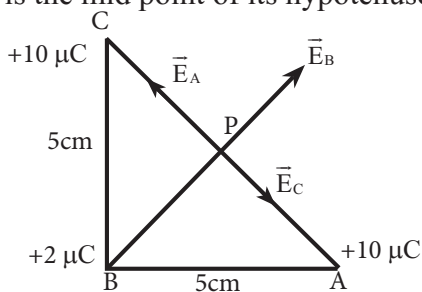


Fig (a): Position of charges.

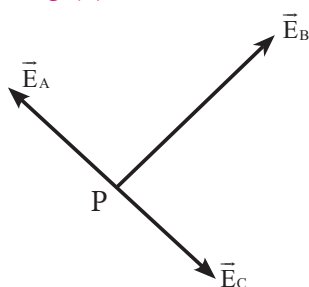


Fig (b): Electric field at point P.

Solution: Electric field is the force on a unit positive charge, the fields at P due to the charges at A, B and C are shown in the Fig. b. $\vec{E}_A$ is the field at P due to charge at A and $\vec{E}_C$ is the field at P due to charge at C. Since P is the midpoint of AC and the fields at A and C are equal in

magnitudes and are opposite in direction, $\vec{E}_A = -\vec{E}_C$. $\vec{E}_A + \vec{E}_C = 0$. Thus, the field at P is only to the charge at B and can be written as

$$\vec{E}_P = \vec{E}_B = \frac{2 \times 10^{-6}}{4\pi\epsilon_0 (BP)^2}$$

$$\begin{aligned} \vec{E}_P &= \frac{2 \times 10^{-6} \times 9 \times 10^9}{(5/\sqrt{2})^2} \\ &= \frac{2 \times 9 \times 10^3 \times 2}{25} \\ &= \frac{36}{25} \times 10^3 \\ &= 1.44 \times 10^{15} \text{ NC}^{-1} \text{ along } \overline{BP} \end{aligned}$$

To calculate BP

$$|\overline{BP}| = (BA) \cos(45^\circ) = \frac{5}{\sqrt{2}}$$

Example 10.7: A simplified model of hydrogen atom consists of an electron revolving about a proton at a distance of $5.3 \times 10^{-11} \text{ m}$. The charge on a proton is $+1.6 \times 10^{-19} \text{ C}$. Calculate the intensity of the electric field due to proton at this distance.

Solution:

$$E = \frac{q}{4\pi\epsilon_0 r^2}$$

$$q = +1.6 \times 10^{-19} \text{ C}$$

$$r = 5.3 \times 10^{-11} \text{ m}$$

$$\frac{1}{4\pi\epsilon_0} = 9.0 \times 10^9 \text{ Nm}^2 \text{C}^{-1}$$

$$E = \frac{1.6 \times 10^{-19}}{(5.3 \times 10^{-11})^2} \times 9.0 \times 10^9 = 5.1 \times 10^{11} \text{ NC}^{-1}$$

The force between electron and proton in hydrogen atom can be calculated by using the

electric field. We have, $E = \frac{F}{q} \therefore F = qE$

$$\begin{aligned} F &= -1.6 \times 10^{-19} \text{ C} \times 5.1 \times 10^{11} \text{ NC}^{-1} \\ &= -8.16 \times 10^8 \text{ N.} \end{aligned}$$

This force is attractive.

Using the Coulomb's law,

$$\begin{aligned} F &= \frac{1}{4\pi\epsilon_0} \frac{q_1 q_2}{r^2} \\ &= 9.0 \times 10^9 \text{ Nm}^2 \text{C}^{-1} \times \frac{(-1.6 \times 10^{-19} \text{ C}) \times (+1.6 \times 10^{-19} \text{ C})}{(5.3 \times 10^{-11} \text{ m})^2} \\ &= -8.6 \times 10^8 \text{ N} \end{aligned}$$

Knowing electric field at a point is useful to estimate the force experienced by a charge at that point.

10.6.3 Electric Lines of Force:

Michael Faraday (1791-1867) introduced the concept of lines of force for visualising electric and magnetic fields. *An electric line of force is an imaginary curve drawn in such a way that the tangent at any given point on this curve gives the direction of the electric field at that point.* See Fig.10.11. If a test charge is placed in an electric field it would be acted upon by a force at every point in the field and will move along a path. **The path along which the unit positive charge moves is called a line of force.**

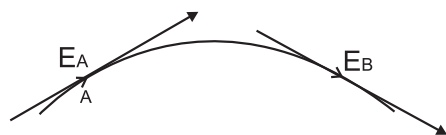


Fig. 10.11: Electric line of force.

A line of force is defined as a curve such that the tangent at any point to this curve gives the direction of the electric field at that point.

The density of field lines indicates the strength of electric fields at the given point in space. Figure 10.12.

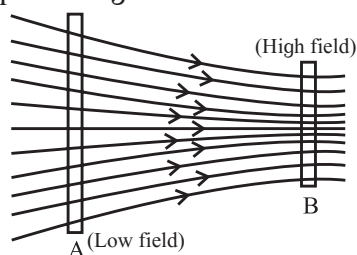


Fig. 10.12: density of field lines and strength of electric field.

Characteristics of electric lines of force

- (1) The lines of force originate from a positively charged object and terminate on a negatively charged object.
- (2) The lines of force neither intersect nor meet each other, as it will mean that electric field has two directions at a single point.
- (3) The lines of force leave or terminate on a conductor normally.
- (4) The lines of force do not pass through conductor i.e. electric field inside a conductor is always zero, but they pass through insulators.
- (5) Magnitude of the electric field intensity is proportional to the number of lines of force per unit area of the surface held perpendicular to the field.

- (6) Electric lines of force are crowded in a region where electric intensity is large.
- (7) Electric lines of force are widely separated from each other in a region where electric intensity is small
- (8) The lines of force of an uniform electric field are parallel to each other and are equally spaced.

The lines of force are purely a geometric construction which help us visualise the nature of electric field in a region. **The lines of force have no physical existence.**

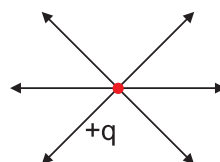


Fig. 10.13 (a): Lines of force due to positive charge.

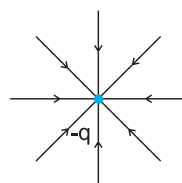


Fig. 10.13 (b): Lines of force due to negative charge.

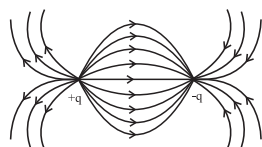


Fig. 10.13 (c): Lines of force due to opposite charge.

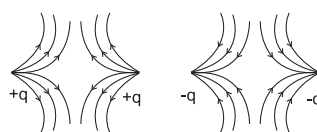


Fig. 10.13 (d): Lines of force due to similar charge.

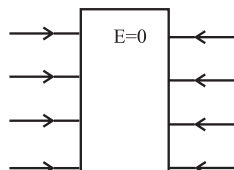


Fig. 10.13 (e): Lines of force terminate on a conductor.

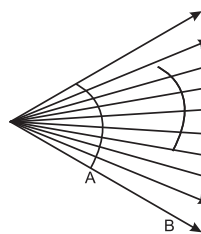


Fig. 10.13 (f): Intensity of a electric field is more at point A and less at B. More lines cross the area at A and less at the same area at B.

Fig. 10.13: The lines of force due to various geometrical arrangement of electrical charges.



Can you tell?

Lines of force are imaginary, can they have any practical use?

10.7 Electric Flux:

As discussed previously, the number of lines of force per unit area is the intensity of the electric field $\vec{E}$.

$$\therefore E = \frac{\text{Number of lines of force}}{\text{Area enclosing the lines of force}} \quad (10.14)$$

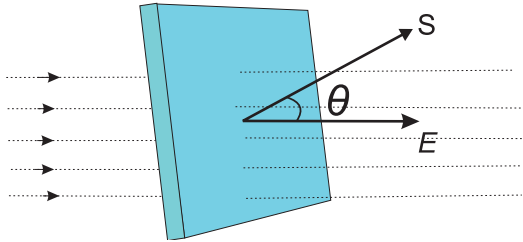


Fig. 10.14: Flux through area S.

Number of lines of force = $(E) \cdot (\text{Area})$
When the area is inclined at an angle θ with the direction of electric field, Fig. 10.14, the electric flux can be calculated as follows.

Let the angle between electric field $\vec{E}$ and area vector $d\vec{S}$ be θ , then the electric flux passing through area dS is given by

$$d\phi = (\text{component of } d\vec{S} \text{ along } \vec{E}) \cdot (\text{area of } d\vec{S})$$

$$d\phi = E (dS \cos \theta)$$

$$d\phi = E dS \cos \theta$$

$$d\phi = \vec{E} \cdot d\vec{S} \quad \text{--- (10.15)}$$

Total flux through the entire surface

$$\Phi = \int d\phi = \int_s \vec{E} \cdot d\vec{S} = \vec{E} \cdot \vec{S} \quad \text{--- (10.16)}$$

The SI unit of electric flux can be calculated using,

$$\Phi = \vec{E} \cdot \vec{S} = (\text{V/m}) \text{ m}^2 = \text{Vm}$$

10.8 Gauss' Law:

Karl Friedrich Gauss (1777-1855) one of the greatest mathematician of all times, formulated a law expressing the relationship between the electric charge and its electric field which is called the Gauss' law. Gauss' law is analogous to Coulomb's law in the sense that it too expresses the relationship between electric field and electric charge. Gauss' law provides equivalent method for finding electric intensity. It relates values of field at a closed surface and the total charges enclosed by that surface.

Consider a closed surface of any shape which encloses number of positive electric charges (Fig. 10.15). To prove Gauss' theorem,

imagine a small charge $+q$ present at a point O inside closed surface. Imagine an infinitesimal area dA of the given irregular closed surface.

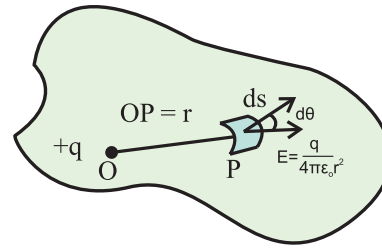


Fig. 10.15: Gauss' law.

The magnitude of electric field intensity at point P on dS due to charge $+q$ at point O is,

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

The direction of E is away from point O. Let θ be the angle subtended by normal drawn to area dS and the direction of E . Electric flux, $d\phi$, passing through area dS , = $E \cos \theta dS$

$$= \frac{q}{4\pi\epsilon_0 r^2} \cos \theta dS$$

$$= \left(\frac{q}{4\pi\epsilon_0} \right) \left(\frac{dS \cos \theta}{r^2} \right)$$

$$d\phi = \left(\frac{q}{4\pi\epsilon_0} \right) d\omega \quad \text{--- (10.17)}$$

where, $d\omega = \frac{dS \cos \theta}{r^2}$ is the solid angle subtended by area dS at point O.

Total electric flux, ϕ_E , crossing the given closed surface can be obtained by integrating Eq. (10.17) over its area. Thus,

$$\Phi_E = \int d\phi = \int_s \vec{E} \cdot d\vec{S} = \int \frac{q}{4\pi\epsilon_0} d\omega = \frac{q}{4\pi\epsilon_0} \int d\omega$$

But $\int d\omega = 4\pi$ = solid angle subtended by entire closed surface at point O

$$\text{Total flux} = \frac{q}{4\pi\epsilon_0} (4\pi)$$

$$\Phi_E = \int_s \vec{E} \cdot d\vec{S} = +q / \epsilon_0 \quad \text{--- (10.18)}$$

This is true for every electric charge enclosed by a given closed surface.

Total flux due to charge q_1 , over the given closed surface = $+ q_1 / \epsilon_0$

Total flux due to charge q_2 , over the given closed surface $= + q_2/\epsilon_0$

Total flux due to charge q_n , over the given closed surface $= + q_n/\epsilon_0$

Positive sign in Eq. (10.18) indicates that the flux is directed outwards, away from the charge. If the charge is negative, the flux will be directed inwards as shown in Fig 10.16 (b). If a charge is outside the closed surface the net flux through it will be zero Fig 10.16 (c).

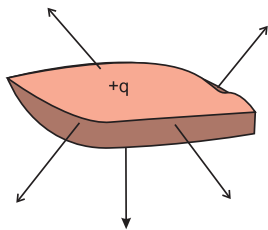


Fig. 10.16 (a): Flux due to positive charge.

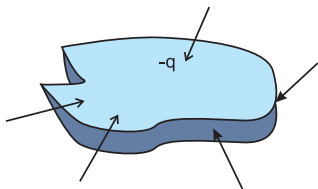


Fig. 10.16 (b): Flux due to negative charge.

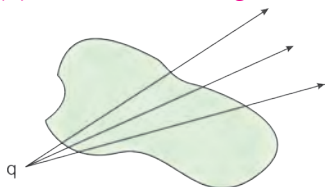


Fig. 10.16 (c): Flux due to charge outside a closed surface is zero.

According to the superposition principle, the total flux ϕ due to all charges enclosed within the given closed surface is

$$\Phi_E = \frac{q_1}{\epsilon_0} + \frac{q_2}{\epsilon_0} + \frac{q_3}{\epsilon_0} + \dots + \frac{q_n}{\epsilon_0} = \sum_{i=1}^{i=n} \frac{q_i}{\epsilon_0} = \frac{Q}{\epsilon_0}$$

Statement of Gauss' law

The flux of the net electric field through a closed surface equals the net charge enclosed by the surface divided by ϵ_0

$$\int \vec{E} \cdot d\vec{S} = \frac{Q}{\epsilon_0}$$

where Q is the total charge within the surface.

Gauss' law is applicable to any closed surface of regular or irregular shape.

Example 10.8: A charge of 5.0 C is kept at the centre of a sphere of radius 1 m. What is the flux passing through the sphere? How will this value change if the radius of the sphere is doubled?

Solution: Flux per unit area is given by Eq. 10.16.

According to Gauss law, the total flux through the sphere $\Phi = \int \vec{E} \cdot d\vec{S}$, where the integration is over the surface of the sphere. As the electric field is same all over the sphere i.e. $|\vec{E}| = \text{constant}$ and the direction of $\vec{E}$ as well as that of $d\vec{S}$ is along the radius, we get

$$\text{flux} = \Phi = |\vec{E}| 4\pi R^2$$

$$E = \frac{q}{4\pi\epsilon_0 r^2} = 9 \times 10^9 \times \frac{5.0 \text{ C}}{(1.0 \text{ m})^2}$$

$$E = 9 \times 10^9 \times 5 = 4.5 \times 10^{10} \text{ NC}^{-1}$$

$$\phi = \vec{E} \cdot \vec{S}$$

$$\therefore \text{flux} = 4.5 \times 10^{10} \times 4\pi(1)^2$$

$$= 5.65 \times 10^{11} \text{ Vm}$$

Thus the total flux is independent of radius.

$E \propto 1/r^2$, and area $\propto r^2$. This can also be seen from Gauss' law, where the net flux crossing a closed surface is equal to q/ϵ_0 where q is the net charge inside the closed surface. As the charge inside the sphere is unchanged, the flux passing through a sphere of any radius is the same. Thus, if the radius of the sphere is increased by a factor of 2, the net flux passing through its surface remains unchanged. As shown in Fig. 10.17, same number of lines of force cross both the surfaces. The total flux is independent of shape of the closed surface because Eq. 10.18 does not involve any radius.

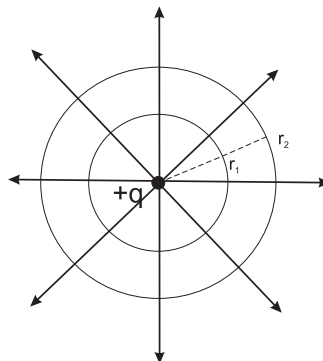


Fig. 10.17: Flux is independent of the shape and size of closed surface.



Do you know ?

Gaussian surface

All the lines of force originating from a point charge penetrate an imaginary three dimensional surface. The total flux $\Phi_E = q/\epsilon_0$. The same number of lines of force will cross the surface of any shape. The total flux through both the surfaces is the same. Calculating flux involves calculating $\int \vec{E} \cdot d\vec{s}$, hence it is convenient to consider a regular surface surrounding the given charge distribution. A surface enclosing the given charge distribution and symmetric about it is a Gaussian surface.

For example, if we have a point charge the Gaussian surface will be a sphere. If the charge distribution is linear, the Gaussian surface would be a cylinder with the charges distributed along its axis. Gaussian surface offers convenience of calculating the integral $\int \vec{E} \cdot d\vec{s}$.

Remember that a Gaussian surface is purely imaginary and does not exist physically.

10.9 Electric Dipole:

A pair of equal and opposite charges separated by a finite distance is called an electric dipole. It is shown in Fig. (10.18).

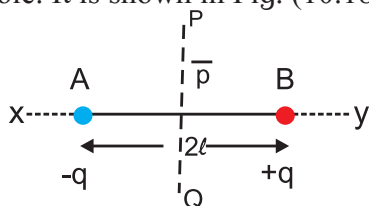


Fig. 10.18: Electric dipole. x-y axial line, P-Q equatorial line.

Line joining the two charges is called the dipole axis. A line passing through the dipole axis is called **axial line**. A line passing through the centre of the dipole and perpendicular to the axial line is called the **equatorial line** as shown in Fig. 10.18.

Strength of a dipole is measured in terms of a quantity called the **dipole moment**. Let q be the magnitude of each charge and $2\vec{l}$ be the distance from negative charge to positive charge. Then the product $q (2\vec{l})$ is called the

dipole moment $\vec{p}$.

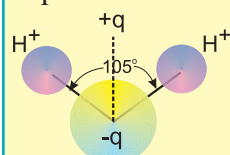
Dipole moment is defined as

$$\vec{p} = q(2\vec{l}) \quad \text{--- (10.19)}$$

A dipole moment is a vector whose magnitude is $q(2l)$ and the direction is from the negative to the positive charge. The unit of dipole moment is Columb-meter (Cm) or Debye (D). $1D = 3.33 \times 10^{-30} \text{ Cm}$. If two charges $+e$ and $-e$ are separated by $1.0\text{\AA}$, the dipole moment is $1.6 \times 10^{-29} \text{ Cm}$ or 4.8 D . For example, a water molecule has a permanent dipole moment of

Natural dipole:

The water molecule is non-linear, i.e., the two hydrogen atoms and one oxygen atom are not in a straight line. The two hydrogen-oxygen bonds in water molecule are at an angle of 105° . The positive charge of a water molecule is effectively concentrated on the hydrogen side and the negative charge on the oxygen side of the molecule. Thus, the positive and negative charges of the water molecule are inherently separated by a small distance. This separation of positive and negative charges gives rise to the permanent dipole moment of a water molecule.



Molecules of water, ammonia, sulphur dioxide, sodium chloride etc. have an inherent separation of centers of positive and negative charges. Such molecules are called **polar molecules**.

Polar molecules are the molecules in which the center of positive charge and the negative charge is naturally separated.

Molecules such as H_2 , Cl_2 , CO_2 , CH_4 and many others have their positive and negative charges effectively centered at the same point and are called **non-polar molecules**.

Non-polar molecules are the molecules in which the center of positive charge and the negative charge is one and the same. They do not have a permanent electric dipole. When an external electric field is applied to such molecules the centers of positive and negative charge are displaced and a dipole is induced.

6.172×10^{-30} Cm or 1.85 D. Its direction is from oxygen to hydrogen. See box on Natural dipole.

10.9.1 Couple Acting on an Electric Dipole in a Uniform Electric Field:

Consider an electric dipole placed in a uniform electric field E . The axis of electric dipole makes an angle θ with the direction of electric field as shown in Fig. 10.19 a.

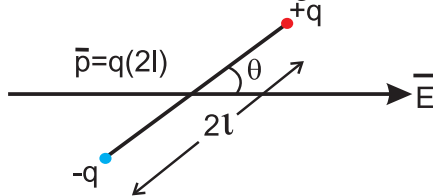


Fig. 10.19 (a): Dipole in uniform electric field.

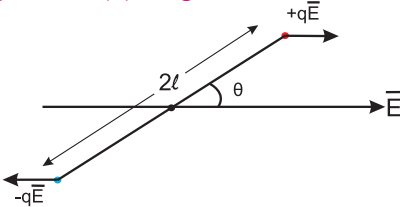


Fig. 10.19 (b): Couple acting on a dipole.

Figure 10.19. b shows the couple acting on an electric dipole in uniform electric field.

The force acting on charge $-q$ at A is

$\vec{F}_A = -q\vec{E}$ in the direction opposite to that of $\vec{E}$ and the force acting on charge $+q$ at B is $\vec{F}_B = +q\vec{E}$ in the direction of $\vec{E}$. Since $\vec{F}_A = -\vec{F}_B$, the two equal and opposite forces separated by a distance form a couple. Moment of the couple is called torque and is defined by $\vec{\tau} = \vec{d} \times \vec{F}$ where, d is the perpendicular distance between the two equal and opposite forces.

$\therefore$ Magnitude of Torque =

Magnitude of force $\times$ Perpendicular distance

$$\therefore \text{Torque on the dipole} = \vec{\tau} = \vec{BP} \times q\vec{E}$$

$$\therefore \tau = qE2l \sin \theta \quad \text{--- (10.20)}$$

$$\text{but } \vec{p} = q \times 2\vec{l}$$

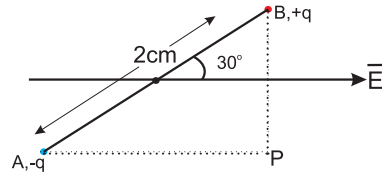
$$\therefore \tau = pE \sin \theta \quad \text{--- (10.21)}$$

$$\text{In vector form } \vec{\tau} = \vec{p} \times \vec{E} \quad \text{--- (10.22)}$$

If $\theta = 90^\circ$ $\sin \theta = 1$, then $\tau = pE$

When the axis of electric dipole is perpendicular to uniform electric field, torque of the couple acting on the electric dipole is maximum, i.e., $\tau = pE$. If $\theta = 0$ then $\tau = 0$, this is the minimum torque on the dipole. Torque tends to align the dipole axis along the direction of electric field.

Example 10.9: An electric dipole of length 2.0 cm is placed with its axis making an angle of 30° with a uniform electric field of 10^5 N/C. as shown in figure. If it experiences a torque of $10\sqrt{3}$ Nm, calculate the magnitude of charge on the dipole.



Solution: Given

$$\tau = 10\sqrt{3} \text{ Nm}, \quad E = 10^5 \text{ N/C},$$

$$2l = 2.0 \times 10^{-2} \text{ m}, \quad \theta = 30^\circ$$

$$\tau = qE2l \sin \theta$$

$$10\sqrt{3} \text{ Nm} = q10^5 \text{ N/C} \cdot 2.0 \times 10^{-2} \text{ m} \left(\frac{1}{2} \right)$$

$$\therefore q = \frac{10\sqrt{3}}{10^3} = \sqrt{3} \times 10^{-2} \text{ C}$$

$$= 1.73 \times 10^{-2} \text{ C}$$

10.9.2 Electric Intensity at a Point due to an Electric Dipole:

Case 1 : At a point on the axis of a dipole.

Consider an electric dipole consisting of two charges $-q$ and $+q$ separated by a distance $2l$ as shown in Fig. 10.20. Let P be a point at a distance r from the centre C of the dipole. The electric intensity $\vec{E}_a$ at P due to the dipole is the vector sum of the field due to the charge $-q$ at A and $+q$ at B.

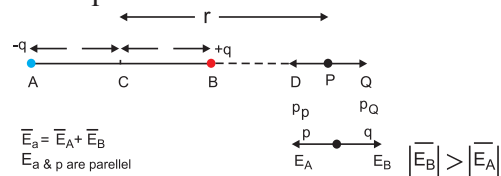


Fig. 10.20: Electric field of a dipole along its axis.

Electric field intensity at P due to the charge $-q$ at A

$$= \vec{E}_A = \frac{1}{4\pi\epsilon_0} \frac{(-q)}{(r+l)^2} \vec{PD}$$

where $\vec{PD}$ is unit vector directed along $\vec{PD}$

Electric intensity at P due to charge $+q$ at B

$$= \vec{E}_B = \frac{1}{4\pi\epsilon_0} \frac{q}{(r-l)^2} \vec{PQ}$$

where $\vec{PQ}$ is a unit vector directed along $\vec{PQ}$.

The magnitude of $\vec{E}_B$ is greater than that of $\vec{E}_A$. (Because $BP < AB$)

Resultant field $\vec{E}_a$ at P on the axis, due to the dipole is

$$\vec{E}_a = \vec{E}_B + \vec{E}_A$$

The magnitude of $\vec{E}_a$ is given by

$$\begin{aligned} |\vec{E}_a| &= \frac{1}{4\pi\epsilon_0} \left[\frac{q}{(r-l)^2} - \frac{q}{(r+l)^2} \right] \\ |\vec{E}_a| &= \frac{q}{4\pi\epsilon_0} \left[\frac{r^2 + l^2 + 2lr - r^2 + 2lr - l^2}{(r^2 - l^2)^2} \right] \\ |\vec{E}_a| &= \frac{2(2lq)r}{4\pi\epsilon_0(r^2 - l^2)^2} \end{aligned}$$

But $2lq = p$, the dipole moment

$$|\vec{E}_a| = \frac{1}{4\pi\epsilon_0} \frac{2pr}{(r^2 - l^2)^2} \quad \text{--- (10.23)}$$

$\vec{E}_a$, is directed along PQ, which is the direction of the dipole moment $\vec{p}$ i.e. from the negative to the positive charge, parallel to the axis. If $r \gg l$, l^2 can be neglected compared to r^2 ,

$$|\vec{E}_a| = \frac{1}{4\pi\epsilon_0} \frac{2p}{r^3} \quad \text{--- (10.24)}$$

The field will be along the direction of the dipole moment $\vec{p}$.

Case 2: At a point on the equatorial line. As shown in Fig. 10.21 (a)

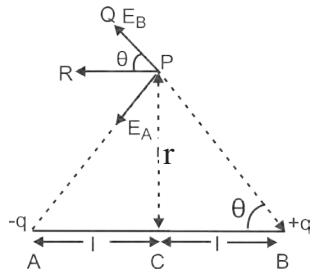


Fig. 10.21 (a): Electric field of a dipole at a point on the equatorial line.

Electric field at point P due to charge -q at A is:

$$\vec{E}_A = \frac{1}{4\pi\epsilon_0} \frac{(-q)}{(AP)^2} \overrightarrow{PA}$$

where $\overrightarrow{PA}$ is the unit vector direction along $\overrightarrow{PA}$.

Similarly, Electric field at P due to charge +q at B is:

$$\vec{E}_B = \frac{1}{4\pi\epsilon_0} \frac{(+q)}{(BP)^2} \overrightarrow{PQ}$$

where $\overrightarrow{PQ}$ is the unit vector directed along $\overrightarrow{PQ}$ or $\overrightarrow{BP}$

Electric field at P is the, sum of $\vec{E}_A$ and $\vec{E}_B$

$$\therefore \vec{E}_{eq} = \vec{E}_A + \vec{E}_B$$

Consider ΔACP

$$(AP)^2 = (PC)^2 + (AC)^2 = r^2 + l^2 = (BP)^2$$

$$\therefore |\vec{E}_A| = \frac{1}{4\pi\epsilon_0} \frac{q}{(r^2 + l^2)} \quad \text{--- (10.25)}$$

$$|\vec{E}_B| = \frac{1}{4\pi\epsilon_0} \frac{q}{(r^2 + l^2)} \quad \text{--- (10.26)}$$

$$|\vec{E}_A| = |\vec{E}_B|$$

The resultant of fields $\vec{E}_A$ and $\vec{E}_B$ acting at point P can be calculated by resolving these vectors $\vec{E}_A$ and $\vec{E}_B$ along the equatorial line and along a direction perpendicular to it.

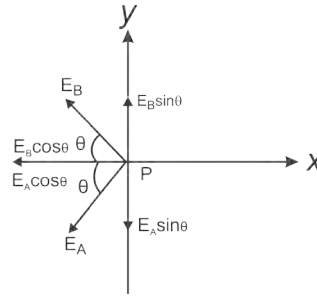


Fig. 10.21 (b): Components of the field at point P.

Consider Fig. 10.21 (b). Let the y-axis coincide with the equator of the dipole x-axis will be parallel to dipole axis, as shown. The origin is at point P.

The y-components of E_A and E_B are $E_A \sin \theta$ and $E_B \sin \theta$ respectively. They are equal in magnitude but opposite in direction and cancel each other. There is no contribution from them towards the resultant.

The x-components of E_A and E_B are $E_A \cos \theta$ and $E_B \cos \theta$ respectively. They are of equal magnitude and are in the same direction

$$\therefore |\vec{E}_{eq}| = E_A \cos \theta + E_B \cos \theta \quad \text{--- (10.27)}$$

By using Eq. 10.25 and 10.26

$$\begin{aligned} |\vec{E}_{eq}| &= 2E_A \cos \theta \\ &= 2 \left(\frac{q}{4\pi\epsilon_0(r^2 + l^2)} \right) \frac{l}{\sqrt{r^2 + l^2}} \\ &= \frac{2ql}{4\pi\epsilon_0(r^2 + l^2)^{3/2}} \end{aligned}$$

If $r \gg l$ then l^2 is very small compared to r^2

$$|\vec{E}_{eq}| = \frac{1}{4\pi\epsilon_0} \frac{p}{(r^2)^{3/2}} = \frac{1}{4\pi\epsilon_0} \frac{p}{r^3} \quad \text{--- (10.28)}$$

The direction of this field is along $-\vec{p}$ (anti-parallel to $\vec{p}$) as shown in Fig. 10.21 (c).

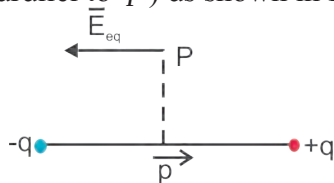


Fig. 10.21 (c): Electric field at point P is anti-parallel to $\vec{p}$.

Comparing Eq. 10.28 and 10.24 we find that the electric intensity at an axial point is twice that at a point on the equatorial position, lying at the same distance from the centre of the dipole.

10.10 Continuous Charge Distribution:

A system of charges can be considered as a continuous charge distribution, if the charges are located very close together, compared to their distances from the point where the intensity of electric field is to be found out.

The charge distribution is continuous in the sense that, a system of closely spaced charges is equivalent to a total charge which is continuously distributed along a line or a surface or a volume. To find the electric field due to continuous charge distribution, we define following terms for different types of charge distribution.

(a) Linear charge density (λ).

As shown in Fig. 10.22 charge q is uniformly distributed along a linear conductor of length l . The linear charge density λ is defined as,

$$\lambda = \frac{q}{l} \quad \text{--- (10.29)}$$

SI unit of λ is (C/m).

For example, charge distributed uniformly on a straight thin rod or a thin nylon thread. If the charge is not distributed uniformly over the length of thin conductor then charge dq on small element of length dl can be written as $dq = \lambda dl$



Fig. 10.22: Linear charge.

(b) Surface charge density (σ)

Suppose a charge q is uniformly distributed over a surface of area A . As shown in Fig. 10.23, then the surface charge density σ is defined as

$$\sigma = \frac{q}{A} \quad \text{--- (10.30)}$$

SI unit of σ is (C/m²)

For example, charge distributed uniformly on a thin disc or a synthetic cloth. If the charge is not distributed uniformly over the surface of a conductor, then charge dq on small area element dA can be written as $dq = \sigma dA$.

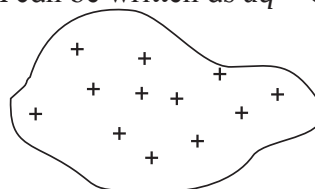


Fig. 10.23: Surface charge.

(c) Volume charge density (ρ)

Suppose a charge q is uniformly distributed throughout a volume V , then the volume charge density ρ is defined as the charge per unit volume.

$$\rho = \frac{q}{V} \quad \text{--- (10.31)}$$

S.I. unit of ρ is (C/m³)

For example, charge on a solid plastic sphere or a solid plastic cube.

If the charge is not distributed uniformly over the volume of a material, then charge dq over small volume element dV can be written as $dq = \rho dV$.

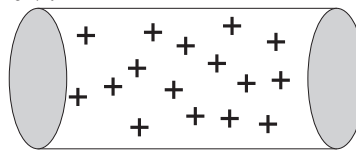


Fig. 10.24: Volume charge.

Electric field due to a continuous charge distribution can be calculated by adding electric fields due to all these small charges.



Can you tell?

The surface charge density of Earth is $\sigma = -1.33 \text{ nC/m}^2$. That is about 8.3×10^9 electrons per square meter. If that is the case why don't we feel it?



Do you know ?

Static charge can be useful

Static charges can be created whenever there is a friction between an insulator and other object. For example, when an insulator like rubber or ebonite is rubbed against a cloth, the friction between them causes electrons to be transferred from one to the other. This property of insulators is used in many applications such as Photocopier, Inkjet printer, Painting metal panels, Electrostatic precipitation/separators etc.

Static charge can be harmful

- When charge transferred from one body to other is very large sparking can take place. For example lightning in sky.
- Sparking can be dangerous while refuelling your vehicle.

- One can get static shock if charge transferred is large.
- Dust or dirt particles gathered on computer or TV screens can catch static charges and can be troublesome.

Precautions against static charge

- Home appliances should be grounded.
- Avoid using rubber soled footwear.
- Keep your surroundings humid. (dry air can retain static charges).



Internet my friend

- [https://www.physicsclassroom.com>class](https://www.physicsclassroom.com/class)
- [https://courses.lumenlearning.com>elect](https://courses.lumenlearning.com/elect)
- [https://www.khanacademy.org>science](https://www.khanacademy.org/science).
- [https://www.topper.com>guides>physics](https://www.topper.com/guides/physics)



Exercises

1. Choose the correct option.

- A positively charged glass rod is brought close to a metallic rod isolated from ground. The charge on the side of the metallic rod away from the glass rod will be
 - same as that on the glass rod and equal in quantity
 - opposite to that on the glass of and equal in quantity
 - same as that on the glass rod but lesser in quantity
 - same as that on the glass rod but more in quantity
- An electron is placed between two parallel plates connected to a battery. If the battery is switched on, the electron will
 - be attracted to the +ve plate
 - be attracted to the -ve plate
 - remain stationary
 - will move parallel to the plates
- A charge of $+7 \mu\text{C}$ is placed at the centre of two concentric spheres with radius 2.0 cm and 4.0 cm respectively. The ratio of the flux through them will be
 - 1:4
 - 1:2
 - 1:1
 - 1:16
- Two charges of 1.0 C each are placed one meter apart in free space. The force between them will be
 - 1.0 N
 - $9 \times 10^9 \text{ N}$
 - $9 \times 10^{-9} \text{ N}$
 - 10 N
- Two point charges of $+5 \mu\text{C}$ are so placed that they experience a force of $80 \times 10^{-3} \text{ N}$. They are then moved apart, so that the force is now $2.0 \times 10^{-3} \text{ N}$. The distance between them is now
 - 1/4 the previous distance
 - double the previous distance
 - four times the previous distance
 - half the previous distance
- A metallic sphere A isolated from ground is charged to $+50 \mu\text{C}$. This sphere is brought in contact with other isolated metallic sphere B of half the radius of sphere A. The charge on the two sphere will be now in the ratio
 - 1:2
 - 2:1
 - 4:1
 - 1:1
- Which of the following produces uniform electric field?

- (A) point charge
(B) linear charge
(C) two parallel plates
(D) charge distributed on a circular arc
- viii. Two point charges of $A = +5.0 \mu\text{C}$ and $B = -5.0 \mu\text{C}$ are separated by 5.0 cm. A point charge $C = 1.0 \mu\text{C}$ is placed at 3.0 cm away from the centre on the perpendicular bisector of the line joining the two point charges. The charge at C will experience a force directed towards
(A) point A
(B) point B
(C) a direction parallel to line AB
(D) a direction along the perpendicular bisector
- iii. Four charges of $+6 \times 10^{-8} \text{ C}$ each are placed at the corners of a square whose sides are 3 cm each. Calculate the resultant force on each charge and show its direction on a diagram drawn to scale.
[Ans: $6.89 \times 10^{-2} \text{ N}$]
- iv. The electric field in a region is given by $\vec{E} = 5.0 \text{ kN/C}$. Calculate the electric flux through a square of side 10.0 cm in the following cases
(a) the square is along the XY plane
[Ans: $5.0 \times 10^{-2} \text{ Vm}$]
(b) The square is along XZ plane
[Ans: Zero]
(c) The normal to the square makes an angle of 45° with the Z axis.
[Ans: $3.5 \times 10^{-2} \text{ Vm}$]

2. Answer the following questions.

- What is the magnitude of charge on an electron?
- State the law of conservation of charge.
- Define a unit charge.
- Two parallel plates have a potential difference of 10V between them. If the plates are 0.5 mm apart, what will be the strength of electric charge.
- What is uniform electric field?
- If two lines of force intersect at one point. What does it mean?
- State the units of linear charge density.
- What is the unit of dipole moment?
- What is relative permittivity?
- Three equal charges of $10 \times 10^{-8} \text{ C}$ respectively, each located at the corners of a right triangle whose sides are 15 cm, 20 cm and 25 cm respectively. Find the force exerted on the charge located at the 90° angle.
[Ans: $4.59 \times 10^{-3} \text{ N}$]
- A potential difference of 5000 volt is applied between two parallel plates 5 cm apart. A small oil drop having a charge of $9.6 \times 10^{-19} \text{ C}$ falls between the plates. Find (a) electric field intensity between the plates and (b) the force on the oil drop.
[Ans: (a) $1.0 \times 10^5 \text{ N/C}$
(b) $9.6 \times 10^{-14} \text{ N}$]

3. Solve numerical examples.

- Two small spheres 18 cm apart have equal negative charges and repel each other with the force of $6 \times 10^{-3} \text{ N}$. Find the total charge on both spheres.
[Ans: $q = 2.938 \times 10^{-7} \text{ C}$]
- A charge $+q$ exerts a force of magnitude -0.2 N on another charge $-2q$. If they are separated by 25.0 cm, determine the value of q .
[Ans: $q = 0.8333 \mu\text{C}$]
- Calculate the electric field due to a charge of $-8.0 \times 10^{-8} \text{ C}$ at a distance of 5.0 cm from it.
[Ans: $-2.88 \times 10^{-2} \text{ N/C}$]



Can you recall?

1. Do you recall that the flow of charged particles in a conductor constitutes a current?
2. An electric current in a metallic conductor such as a wire is due to flow of electrons, the negatively charged particles in the wire.
3. What is the role of the valence electrons which are the outermost electrons of an atom?

11.1 Introduction:

The valence electrons become de-localized when a large number of atoms come together in a metal. These are the conduction electrons or free electrons constituting an electric current when a potential difference is applied across the conductor.

11.2 Electric current :

Consider an imaginary gas of both negatively and positively charged particles. Fig. 11.1 shows the negatively and positively charged particles flowing randomly in various directions across a plane P. In a time interval t , let the amount of positive charge flowing in the forward direction be q^+ and the amount of negative charge flowing in the forward direction be q^- .

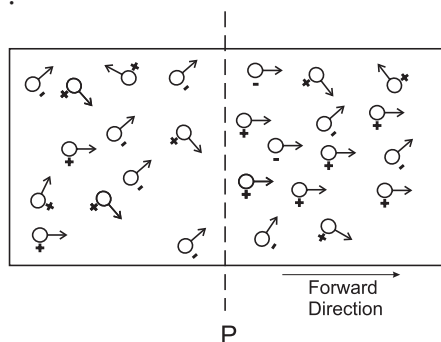


Fig. 11.1: Flow of charged particles.

Thus the net charge flowing in the forward direction is $q = q^+ - q^-$. For a steady flow, this quantity is proportional to the time t . The ratio $\frac{q}{t}$ is defined as the current I .

$$I = \frac{q}{t} \quad \text{--- (11.1)}$$

SI unit of the current is ampere (A), that of the charge and time is coulomb (C) and second (s) respectively.

Let I be the current varying with time. Let Δq be the amount of net charge flowing across

the plane P from time t to $t + \Delta t$, i.e. during the time interval Δt . Then the current is given by

$$I(I) = \lim_{\Delta t \rightarrow 0} \frac{\Delta q}{\Delta t} \quad \text{--- (11.2)}$$

Here, the current is expressed as the limit of the ratio $\Delta q / \Delta t$ as Δt tends to zero.

The current during lightening could be as high as 10,000 A, while the current in the house hold circuit could be of the order of a few amperes. Currents of the order of a milliamper (mA), a microampere (μA) or a nanoampere (nA) are common in semiconductor devices.

11.3 Flow of current through a conductor :

A current can be generated by positively or negatively charged particles. In an electrolyte, both positively and negatively charged particles take part in the conduction. In a metal, the free electrons are responsible for conduction. These electrons flow and generate a net current under the action of an applied electric field. As long as a steady field exists, the electrons continue to flow in the form of a steady current. Such steady electric fields are generated by cells and batteries.



Do you know ?

Sign convention : The direction of the current in a circuit is drawn in the direction in which positively charged particles would move, even if the current is constituted by the negatively charged particles, electrons, which move in the direction opposite to that the electric field. We use this as a convention.

11.4 Drift speed :

Imagine a copper rod with no current flowing through it. Fig 11.2 shows the schematic of a conductor with the free electrons

in random motion. There is no net motion of these electrons in any direction. If electric field is applied along the length of the copper rod, and a current is set up in the rod, these electrons still move randomly, but tend to 'drift' in a particular direction. Their direction is opposite to that of the applied electric field.

Direction of electric field : Direction of an electric field at a point is the direction of the force on the test charge placed at that point.

The electrons under the action of the applied electric field drift with a drift speed V_d . The drift speed in a copper conductor is of the order of 10^{-4} m/s- 10^{-5} m/s, whereas the electron random speed is of the order of 10^6 m/s.

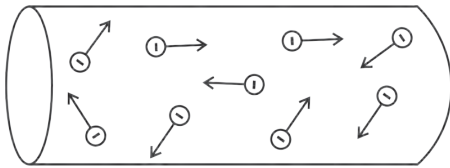


Fig. 11.2: Free electrons in random motion inside the conductor.

How is the current through a conductor related to the drift speed of electrons? Figure 11.3 shows a part of conducting wire with its free electrons having the drift speed V_d in the direction opposite to the electric field $\vec{E}$.

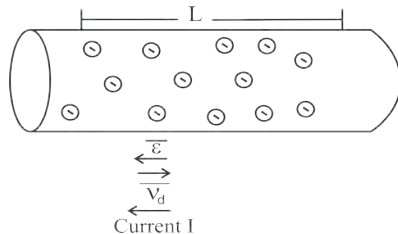


Fig. 11.3: Conducting wire with the applied electric field.

It is assumed that all the electron move with the same drift speed V_d and that, the current I is the same throughout the cross section (A) of the wire. Consider the length L of the wire. Let n be the number of free electrons per unit volume of the wire. Then the total number of electrons in the length L of the conducting wire is nAL . The total charge in the length L is,

$$q = n A L e \quad \text{--- (11.3)}$$

where e is the electron charge.

This is total charge that moves through any cross section of the wire in a certain time interval t ,

$$t = \frac{L}{V_d} \quad \text{--- (11.4)}$$

From the Eq. (11.1), and Eq. (11.3), the current

$$I = \frac{q}{t} = \frac{n A L e}{L/V_d} = n A V_d e \quad \text{--- (11.5)}$$

Hence

$$V_d = \frac{I}{n A e} = \frac{J}{n e} \quad \text{--- (11.6)}$$

where $J = I/A$ is current density. J is uniform over the cross sectional area A of the wire. Its unit is A/m²

$$\text{Here, } J = \frac{I}{A} \quad \text{--- (11.7)}$$

From Eq. (11.6),

$$\vec{J} = (n e) \vec{V}_d \quad \text{--- (11.8)}$$

For electrons, $n e$ is negative and $\vec{J}$ and $\vec{V}_d$ have opposite directions, $\vec{V}_d$ is the drift velocity.

Example 11.1: A metallic wire of diameter 0.02m contains 10^{28} free electrons per cubic meter. Find the drift velocity for free electrons, having an electric current of 100 amperes flowing through the wire.

(Given : charge on electron = 1.6×10^{-19} C)

Solution: Given

$$e = 1.6 \times 10^{-19} \text{ C}$$

$$n = 10^{28} \text{ electrons/m}^3$$

$$D = 0.02\text{m} \quad r = D/2 = 0.01\text{m}$$

$$I = 100 \text{ A}$$

$$V_d = \frac{J}{n e} = \frac{I}{n A e}$$

where A is the cross sectional area of the wire.

$$A = \pi r^2 = 3.142 \times (0.01)^2$$

$$= 3.142 \times 10^{-4} \text{ m}^2$$

$$V_d = \frac{100}{3.142 \times 10^{-4} \times 10^{28} \times 1.6 \times 10^{-19}}$$

$$= \frac{10^{2+4-9}}{5.027}$$

$$V_d = 10^{-3} \times 0.1989 = 1.9 \times 10^{-4} \text{ m/s}$$

Example 11.2: A copper wire of radius 0.6 mm carries a current of 1A. Assuming the current to be uniformly distributed over a cross sectional area, find the magnitude of current density.

Solution: Given

$$r = 0.6 \text{ mm} = 0.6 \times 10^{-3} \text{ m}$$

$$I = 1 \text{ A}$$

$$J = ?$$

$$\text{Area of copper wire} = \pi r^2$$

$$= 3.142 \times (0.6)^2 \times 10^{-6}$$

$$= 3.142 \times 0.36 \times 10^{-6}$$

$$= 1.1311 \times 10^{-6} \text{ m}^2$$

$$J = \frac{I}{A} = \frac{1}{1.1311 \times 10^{-6}}$$

$$J = 0.884 \times 10^6 \text{ A/m}^2$$

11.5 Ohm's law :

The relationship between the current through a conductor and applied potential difference was first discovered by German scientist George Simon Ohm in 1828 AD. This relationship is known as Ohm's law.

It states that "The current I through a conductor is directly proportional to the potential difference V applied across its two ends provided the physical state of the conductor is unchanged".

The graph of current versus potential difference across the conductor is a straight line as shown in Fig. 11.4

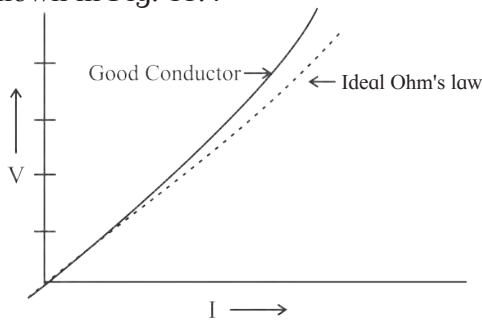


Fig. 11.4: I-V curve for a conductor.

In general, $I \propto V$

$$\text{or } V = IR \quad \text{or } R = \frac{V}{I}, \quad \text{--- (11.9)}$$

where R is a proportionality constant and is called the resistance of the conductor. The unit of resistance is ohm (Ω),

$$1\Omega = \frac{1 \text{ volt}}{1 \text{ ampere}}$$

If potential difference of 1 volt across a conductor produces a current of 1 ampere through it, then the resistance of the conductor is 1Ω .

Reciprocal of resistance is called conductance.

$$C = \frac{1}{R} \quad \text{--- (11.10)}$$

The unit of conductance is siemens or $(\Omega)^{-1}$

Example 11.3: A Flashlight uses two 1.5V batteries to provide a steady current of 0.5A in the filament. Determine the resistance of the glowing filament.

Solution:

$$R = \frac{V}{I} = \frac{3}{0.5} = 6.0\Omega$$

$\therefore$ Resistance of the glowing filament is 6.0Ω .

Physical origin of Ohm's law :

We know that electrical conduction in a conductor is due to mobile charge carriers, the electrons. It is assumed that these conduction electrons are free to move inside the volume of the conductor. During their random motion, electrons collide with the ion cores within the conductor. It is assumed that electrons do not collide with each other. These random motions average to zero. On the application of an electric field $\vec{E}$, the motion of the electrons is a combination of the random motion of electrons due to collisions and that due to the electric field $\vec{E}$. The electrons drift under the action of the field $\vec{E}$ and move in a direction opposite to the direction of the field $\vec{E}$.

Consider an electron of mass m subjected to an electric field $\vec{E}$. The force experienced by the electron will be $\vec{F} = e\vec{E}$. The acceleration experienced by the electron will then be

$$\vec{a} = \frac{e\vec{E}}{m} \quad \text{--- (11.11)}$$

The type of collision the conduction electrons undergo is such that the drift velocity attained before the collision has nothing to do with the drift velocity after the collision. After the collision, the electron will move in random direction, but will still drift in the direction opposite to $\vec{E}$.

Let τ be the average time between two successive collisions. Thus on an average, the

electrons will acquire a drift speed $V_d = a\tau$, where a is the acceleration given by Eq (11.11). Also, at any given instant of time, the average drift speed of the electron will also be $V_d = a\tau$. From Eq. (11.11),

$$V_d = a\tau = \frac{eE\tau}{m} \quad \text{--- (11.12)}$$

From the Eq. (11.6) and Eq. (11.12),

$$V_d = \frac{J}{ne} = \frac{eE\tau}{m} \quad \text{--- (11.13)}$$

which gives

$$E = \left(\frac{m}{e^2 n \tau} \right) J \quad \text{--- (11.14)}$$

or, $E = \rho J$, where ρ is the resistivity of the material and

$$\rho = \frac{m}{ne^2 \tau} \quad \text{--- (11.15)}$$

For a given material, m , n , e^2 and τ will be constant and ρ will also be constant, ρ is independent of $\vec{E}$, the externally applied electric field.

11.6 Limitations of the Ohm's law:

Ohm's law is obeyed by various materials and devices. The devices for which potential difference (V) versus current (I) curve is a straight line passing through origin, inclined to V -axis, are called linear devices or ohmic devices (Fig. 11.4). Resistance of these devices is constant. Several conductors obey the Ohm's law. They follow the linear I - V characteristic.

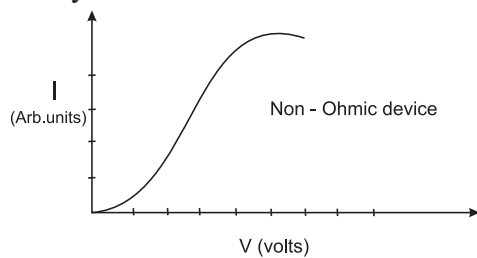


Fig. 11.5: I - V curve for non-Ohmic devices.

The devices for which the I - V curve is not a straight line as shown in Fig. 11.5 are called non-ohmic devices. They do not obey the Ohm's law and the resistance of these devices is a function of V or I ; e.g. liquid electrolytes,

vacuum tubes, junction diodes, thermistors etc. Resistance R for such non-linear devices at a particular value of the potential difference V is given by,

$$R = \lim_{\Delta I \rightarrow 0} \frac{\Delta V}{\Delta I} = \frac{dV}{dI} \quad \text{--- (11.16)}$$

where ΔV is the potential difference between the two values of potential

$$V - \frac{\Delta V}{2} \quad \text{to} \quad V + \frac{\Delta V}{2},$$

and ΔI is the corresponding change in the current.

11.7 Electrical Energy and Power:

Consider a resistor AB connected to a cell in a circuit shown in Fig. 11.6 with current flowing from A to B . The cell maintains a potential difference V between the two terminals of the resistor, higher potential at A and lower at B . Let Q be the charge flowing in time Δt through the resistor from A to B . The potential difference V between the two points A and B , is equal to the amount of work W , done to carry a unit positive charge from A to B . It is given by

$$V = \frac{W}{Q}, \quad W = VQ \quad \text{--- (11.17)}$$

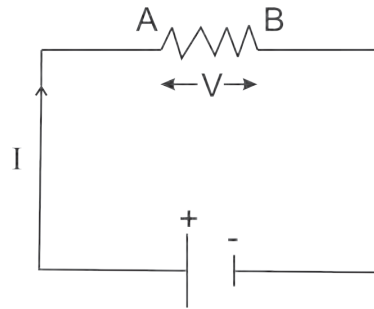


Fig. 11.6: A simple circuit with a cell and a resistor.

The cell provides this energy through the charge Q , to the resistor AB where the work is performed. When the charge Q flows from the higher potential point A to the lower potential point B , i.e. through a decrease in potential of value V , its potential energy decreases by an amount

$$\Delta U = QV = I \Delta t V \quad \text{--- (11.18)}$$

where I is current due to the charge Q flowing in time Δt . Where will this energy go? By the principle of conservation of energy, it is

converted into some other form of energy.

In the limit as $\Delta t \rightarrow 0$,

$$\frac{dU}{dt} = I.V \quad \text{--- (11.19)}$$

Here, $\frac{dU}{dt}$ is power, the time rate of transfer of energy and is given by,

$$P = \frac{dU}{dt} = I.V \quad \text{--- (11.20)}$$

We can also say that this power is transferred by the cell to the resistor or any other device in place of the resistor, such as a motor, a rechargeable battery etc.

Because of the presence of an electric field, the free electrons move across a resistor and there would be an increase in their kinetic energy as they move. When the electrons collide with the ion cores the energy gained by them is shared among the ion cores. Consequently, vibrations of the ions increase, resulting in heating up of the resistor. Thus, some amount of energy is dissipated in the form of heat in a resistor. The energy dissipated in time interval Δt is given by Eq. (11.18). The energy dissipated per unit time is actually the power dissipated and is given by Eq. (11.20).

Using Eq. (11.20), and using Ohm's law, $V=IR$,

$$\therefore P = \frac{V^2}{R} = I^2 R \quad \text{--- (11.21)}$$

It is the power dissipation across a resistor which is responsible for heating it up. For example, the filament of an electric bulb heats up to incandescence, radiating out heat and light.

Example 11.4 : An electric heater takes 6A current from a 230V supply line, calculate the power of the heater and electric energy consumed by it in 5 hours.

Solution : Given

$$I = 6A, V = 230V$$

We know that,

$$P = I \times V = (6A)(230V) = 1380 \text{ W}$$

$$P = 1.38 \text{ kW}$$

$$\text{Energy consumed} = \text{Power} \times \text{time}$$

$$= (1.38 \text{ kW}) \times (5 \text{ h})$$

$$= 6.90 \text{ kWh} (1.0 \text{ Kwh} = 1 \text{ unit of power})$$

$$= 6.9 \text{ units of electrical energy.}$$

11.8 Resistors:

Resistors are used to limit the current following through a particular path of a circuit. Commercially available resistors are mainly of two types :

Carbon resistors and Wire wound resistors. High value resistors are mostly carbon resistors. They are small and inexpensive. The values of these resistors are colour coded to mark their values in ohms. The colour coding is standardized by Electronic Industries Association (EIA). One such resistor is shown in Fig. 11.7.

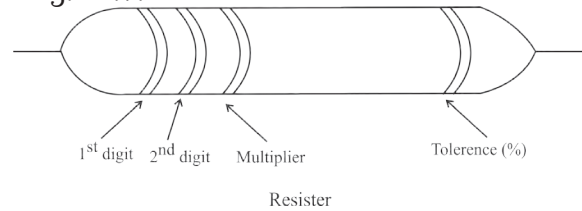


Fig. 11.7: Carbon composition resistor.

Colour code:

Colours	1st digit	2nd digit	Multiplier	Tolerance
Black	0	0	$\times 10^0$	
Brown	1	1	$\times 10^1$	$\pm 1\%$
Red	2	2	$\times 10^2$	$\pm 2\%$
Orange	3	3	$\times 10^3$	
Yellow	4	4	$\times 10^4$	
Green	5	5	$\times 10^5$	
Blue	6	6	$\times 10^6$	
Violet	7	7	$\times 10^7$	
Gray	8	8	$\times 10^8$	
White	9	9	$\times 10^9$	
For Gold			$\times 10^{-1}$	$\pm 5\%$
For Silver			$\times 10^{-2}$	$\pm 10\%$
No colour			-	$\pm 20\%$

Easy Bytes:

Finding it difficult to memorize the colour code sequence? No need to worry, we have a one liner which will help you out "B. B. Roy in Great Britain has Very Good Wife"

B B R O Y G B V G W

This funny one liner makes it easy to recall the sequence of digits and multipliers.

In the four band resistor colour code illustrated in the above table, the first three bands (closest together) indicate the value in ohms. The first two bands indicate two numbers and third band often called decimal multiplier. The fourth band separated by a space from the three value bands, (so that you know which end to start reading from), indicates tolerance of the resistor.

Example

i. Colour code of resistor is

	Yellow	Violet	Orange	Gold
Value :	4	7	10^3	$\pm 5\%$

i.e. $47 \times 10^3 = 47000\Omega = 47k\Omega \pm 5\%$

The value of the resistor is $47k\Omega \pm 5\%$

ii. From given values of resistor; find the colour bands of this resistor

$330\Omega = 33 \times 10$

3 3 10^1

Orange Orange Brown tolerance band

11.8.1 Rheostat:

A rheostat shown in Fig. 11.8 is an adjustable resistor used in applications that require adjustment of current or resistance in an electric circuit. The rheostat can be used to adjust potential difference between two points in a circuit, change the intensity of lights and control the speed of motors, etc. Its resistive element can be a metal wire or a ribbon, carbon films or a conducting liquid, depending upon the application. In hi-fi equipment, rheostats are used for volume control.

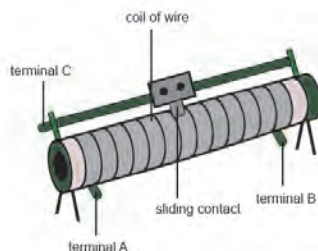


Fig. 11.8: Rheostat.

11.8.2 Combination of Resistors:

I. Series combination of Resistors:

In series combination of resistors, these are connected in single electrical path as shown in Fig 11.9. Hence the same electric current flows through each resistor in a series combination.

Because of series combination, the supply voltage between two resistors R_1 and R_2 is V_1 and V_2 , respectively and the same current I flows through the resistor R_1 and the resistor R_2 . i.e. in series combination, supply voltage is divided and the current remains the same in all the resistors.

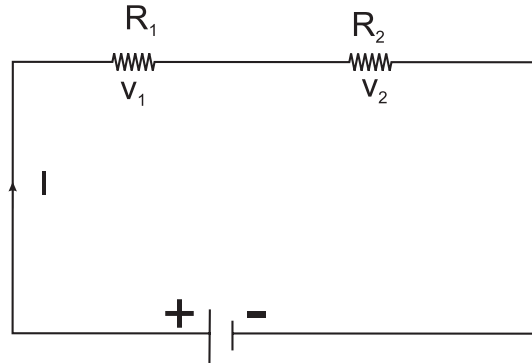


Fig. 11.9: Series combination of two resistors R_1 and R_2 .

According to Ohm's law,

$$R_1 = \frac{V_1}{I}, \quad R_2 = \frac{V_2}{I} \quad \text{--- (11.22)}$$

$$\text{Total voltage } V = V_1 + V_2 \quad \text{--- (11.23)}$$

From equation.... (11.22) and (11.23)

we write

$$V = I(R_1 + R_2) \quad \text{--- (11.24)}$$

$$\therefore V = I R_s \quad \text{--- (11.25)}$$

Thus the equivalent resistance of the series circuit $R_s = R_1 + R_2$

When a number of resistors are connected in series, the equivalent resistance is equal to the sum of individual resistances.

For n number of resistors,

$$R_s = R_1 + R_2 + R_3 + \dots + R_n = \sum_{i=1}^{i=n} R_i \quad \text{--- (11.26)}$$

II. Parallel Combination of Resistors:

In the parallel combination, the resistors are connected in such a way that the same voltage is applied across each resistor.

A number of resistors are said to be connected in parallel if all of them are connected between the same two electrical points each having individual path as shown in Fig. 11.10.

In parallel combination the total current I is divided into I_1 and I_2 as shown in the circuit

diagram Fig.11.10, whereas voltage V across them remains the same,

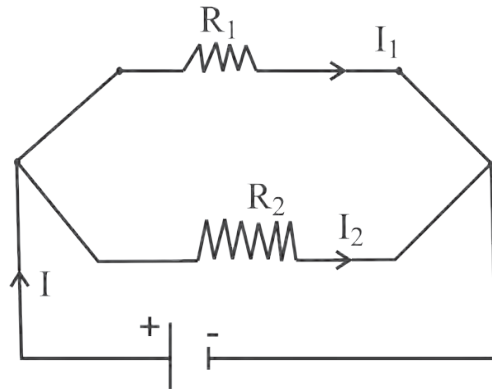


Fig. 11.10 : Two resistors in parallel combination.

$$I = I_1 + I_2 \quad \text{--- (11.27)}$$

where I_1 is current flowing through R_1 and I_2 is current flowing through R_2 .

When Ohm's law is applied to R_1

$$V = I_1 R_1 \quad \text{i.e. } I_1 = \frac{V}{R_1} \quad \text{--- (11.28a)}$$

Ohm's law applied to R_2

$$V = I_2 R_2 \quad \text{i.e. } I_2 = \frac{V}{R_2} \quad \text{--- (11.28b)}$$

From Eq. (11.27) and Eq. (11.28),

$$\therefore I = \frac{V}{R_1} + \frac{V}{R_2},$$

$$\text{If, } I = \frac{V}{R_p},$$

$$\frac{V}{R_p} = \frac{V}{R_1} + \frac{V}{R_2},$$

$$\therefore \frac{1}{R_p} = \frac{1}{R_1} + \frac{1}{R_2}, \quad \text{--- (11.29)}$$

where R_p is the equivalent resistance in parallel combination.

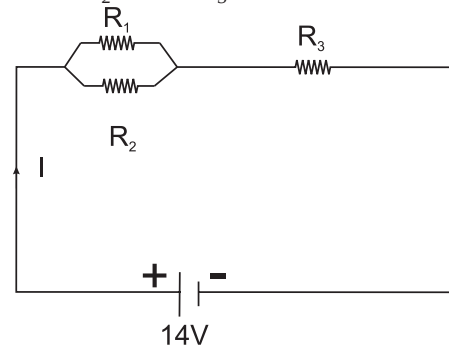
If n resistors $R_1, R_2, R_3, \dots, R_n$ are connected in parallel, the equivalent resistance of the combination is given by

$$\frac{1}{R_p} = \frac{1}{R_1} + \frac{1}{R_2} + \frac{1}{R_3} + \dots + \frac{1}{R_n} = \sum_{i=1}^n \frac{1}{R_i} \quad \text{--- (11.30)}$$

Thus when a number of resistors are connected in parallel, the reciprocal of the equivalent resistance is equal to the sum of the reciprocals of individual resistances.

Example 11.5: Calculate i) total resistance and ii) total current in the following circuit.

$$R_1 = 3\Omega, R_2 = 6\Omega, R_3 = 5\Omega, V = 14V$$



Circuit diagram

Solution:

i) Total resistance $= R_T = R_p + R_3$

$$R_p = \frac{R_1 R_2}{R_1 + R_2} = \frac{3 \times 6}{9} = 2\Omega$$

$$R_T = 2 + 5 = 7\Omega$$

$$\text{Total Resistance} = 7\Omega$$

ii) Total current :

$$I = \frac{V}{R_T} = \frac{14V}{7\Omega}$$

$$I = 2A$$

11.9 Specific Resistance (Resistivity):

At a particular temperature, the resistance of a given conductor is observed to depend on the nature of material of conductor, the area of its cross-section, and its length.

It is found that resistance R of a conductor of uniform cross section is

i. directly proportional to its length l ,

$$\text{i.e. } R \propto l$$

ii. inversely proportional to its area of cross section A ,

$$\text{i.e. } R \propto \frac{l}{A}$$

From i and ii

$$R = \rho \frac{l}{A} \quad \text{--- (11.31)}$$

where ρ is a constant of proportionality and it is called specific resistance or resistivity of the material of the conductor at a given temperature.

From Eq. (11.31), we write

$$\rho = \frac{RA}{l} \quad \text{--- (11.32)}$$

SI unit of resistivity is ohm-meter. Resistivity of a conductor is numerically the resistance per unit length, and per unit area of cross-section of material of the conductor.

i.e. when, $R = 1\Omega$, $A = 1\text{m}^2$ and $l = 1\text{m}$,

then, $\rho = 1\Omega\text{m}$

Conductivity : Reciprocal of resistivity is called conductivity of a material, $\sigma = (\rho)^{-1}$.

SI unit of σ is: $(\Omega\text{m})^{-1}$ i.e. siemens/meter (Sm^{-1})

Table 11.1 : Resistivity of various materials

Material	Resistivity ρ ($\Omega\cdot\text{m}$)	Material	Resistivity ρ ($\Omega\cdot\text{m}$)
Conductors		Semiconductors	
Silver	1.59×10^{-8}	Carbon	3.5×10^{-5}
Copper	1.72×10^{-8}	Germanium	0.5
Gold	2.44×10^{-8}	Silicon	3×10^4
Aluminium	2.82×10^{-8}	Insulators	
Tungsten	5.6×10^{-8}	Glass	$10^{11}\text{-}10^{13}$
Iron	9.7×10^{-8}	Mica	$10^{11}\text{-}10^{15}$
Mercury	95.8×10^{-8}	Rubber (hard)	$10^{13}\text{-}10^{16}$
Nichrome (alloy)	100×10^{-8}	Teflon	10^{16}
		Wood (maple)	3×10^8

Example 11.6: Calculate the resistance per metre, at room temperature, of a constantan (alloy) wire of diameter 1.25mm. The resistivity of constantan at room temperature is $5.0 \times 10^{-7} \Omega\text{m}$.

Solution: $\rho = 5.0 \times 10^{-7} \Omega\text{m}$

$$d = 1.25 \times 10^{-3} \text{ m}$$

$$r = .625 \times 10^{-3} \text{ m}$$

$$\text{Cross-sectional Area} = \pi r^2$$

$$\text{Resistivity } \rho = \frac{RA}{l}$$

$$\text{Resistance per meter} = \frac{R}{l}$$

$$\text{i.e. } \frac{R}{l} = \frac{\rho}{A} = \frac{5 \times 10^{-7}}{(0.625 \times 10^{-3})^2 \times 3.142}$$

$$\frac{R}{l} = 0.41 \Omega\text{m}^{-1}$$

$\therefore$ Resistance per metre = $0.41 \Omega\text{m}^{-1}$

Resistivity ρ is a property of a material, while the resistance R refers to a particular object. Similarly, the electric field $\vec{E}$ at a point is specified in a material with the potential difference across the resistance, and the current density $\vec{J}$ in a material instead of the current I in the resistor. Then for an isotropic material,

$$\rho = \frac{E}{J} \quad \text{or} \quad \vec{E} = \rho \vec{J} \quad \text{---- (11.33)}$$

Again, the SI unit of ρ is

$$\frac{\text{unit}(E)}{\text{unit}(J)} = \frac{\text{V/m}}{\text{A/m}^2} = \frac{\text{V}}{\text{A}} \text{m} = \Omega\cdot\text{m}$$

In terms of conductivity σ of a material, from (11.33),

$$\vec{J} = \frac{1}{\rho} \vec{E} = \sigma \vec{E} \quad \text{--- (11.34)}$$

For a particular resistor, we had (Eq. 11.9) the resistance R given by

$$R = \frac{V}{I}$$

Compare this with the above Eq (11.33).

11.10 Variation of Resistance with Temperature:

Resistivity of a material varies with temperature. It is a property of material. Fig. 11.11 shows the variation of resistivity of copper as a function of temperature (K). It can be seen that the variation is linear over a certain range of temperatures. Such a linear relation can be expressed as,

$$\rho = \rho_0 [1 + \alpha(T - T_0)], \quad \text{--- (11.35)}$$

where T_0 is the chosen reference temperature and ρ_0 in the resistivity at the chosen temperature, for example, T_0 can be 0°C .

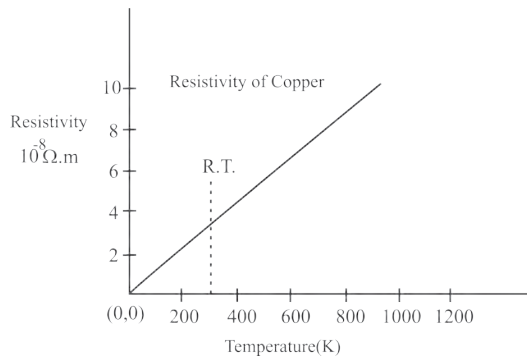


Fig. 11.11: Resistivity as a function of temperature (K).

In the above Eq. (11.35),

$$\alpha = \frac{\rho - \rho_0}{\rho_0(T - T_0)} = \frac{R - R_0}{R_0(T - T_0)} \quad \text{--- (11.36)}$$

Here, α is called the temperature coefficient of resistivity. Table (11.1) shows the resistivity of some of the metals. The temperature coefficient of resistance is defined as the increase in resistance per unit original resistance at the chosen reference temperature, per degree rise in temperature. The unit of α is $^{\circ}\text{C}^{-1}$ or $^{\circ}\text{K}^{-1}$ (per degree celcius or per degree kelvin).

From Eq. (11.36)

$$R = R_0 [1 + \alpha (T - T_0)] \quad \text{--- (11.37)}$$

For small difference in temperatures,

$$\alpha = \frac{1}{R_0} \cdot \frac{dR}{dT} \quad \text{--- (11.38)}$$



Do you know ?

Here, the temperature difference is more important than the temperature alone. Therefore, as the sizes of degrees on the Celsius scale and the Absolute scale are identical, any scale can be used.

Example 11.7: A piece of platinum wire has resistance of $2.5 \, \Omega$ at 0°C . If its temperature coefficient of resistance is $4 \times 10^{-3}/^{\circ}\text{C}$. Find the resistance of the wire at 80°C .

Solution:

$$R_0 = 2.5 \, \Omega$$

$$\alpha = 0.004/^{\circ}\text{C}$$

$$T - 0 = T = 80^{\circ}\text{C}$$

$$R_T = R_0(1 + \alpha T)$$

$$R_T = 2.5 (1 + 0.004 \times 80) = 2.5 (1 + 0.32)$$

$$R_T = 2.5 \times 1.32$$

$$R_T = 3.3 \, \Omega$$

Superconductivity :

We know that the resistivity of a metal decreases as the temperature decreases. In case of some metals and metal alloys, the resistivity suddenly drops to zero at a particular temperature (T_c). This temperature is called critical temperature, for example, mercury loses its resistance completely to zero at 4.2K .

Superconductivity can be harnessed so as to be useful for mankind. It is already in use in obtaining very high magnetic field (a few Tesla) in superconducting magnet. These magnets are used in research quality NMR spectrometers. For its operation, the current carrying coils are required to be kept at a temperature less than the critical temperature of the coil material.

11.11 Electromotive Force (emf):

When charges flow through a conductor, a potential difference has to be established between the two ends of the conductor. For a steady flow of charges, this potential difference is required to be maintained across the two ends of the conductor, the terminals. There is a device that does so by doing work on the charges, thereby maintaining the potential difference. Such a device is called an emf device and it provides the emf ε . The charges move in the conductor owing to the energy provided by the emf device. The device supplies this energy through the work it does.

You must have used some of these emf devices. Power cells, batteries, Solar cells, fuel cells, and even generators, are some examples of emf devices familiar to you.

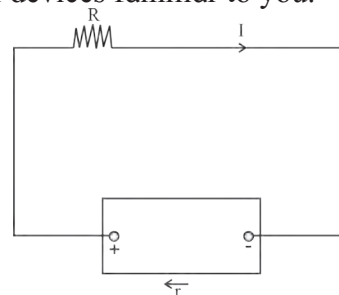


Fig. 11.12: Circuit with emf device.

Fig. 11.12 shows a circuit with an emf device and a resistor R . Here, the emf device keeps the positive terminal (+) at a higher electric potential than the negative terminal (-).

The emf is represented by an arrow from the negative terminal to the positive terminal of a device such as a Voltaic cell. When the circuit is open, there is no net flow of charge carriers within the device. When connected in a circuit, there is a flow of carriers from one terminal to the other terminal inside the emf device. The positive charge carriers move towards the positive terminal which acts as cathode inside the emf device. Thus the positive charge carriers move from the region of lower potential energy, to the region of higher potential energy which is cathode inside the emf device. Here, the energy source is chemical in nature. In a Solar cell, it is the photon energy in the Solar radiation.

Now suppose that a charge dq flows through the cross section of the circuit (Fig. 11.12), in time dt .

It is clear that the same amount of charge dq flows throughout the circuit, including the emf device. It enters the negative terminal (low potential terminal) and leaves the positive terminal (higher potential terminal). Hence, the device must do work dw on the charge dq , so that it moves in the above manner. Thus we define the emf of the emf device.

$$\varepsilon = \frac{dw}{dq} \quad \text{--- (11.39)}$$

The SI unit of emf is joule/coulomb (J/C).

In an ideal device, there is no internal resistance to the motion of charge carriers. The emf of the device is then equal to the potential difference across the two terminals of the device. In a real emf device, there is an internal resistance to the motion of charge carriers. If such a device is not connected in a circuit, there is no current through it. In that case the emf is equal to the potential difference across the two terminals of the emf device connected in a circuit, there is no current through it. If a current (I) flows through an emf device, there is an internal resistance (r) and the emf (ε) differs

from the potential difference across its two terminals (V).

$$V = \varepsilon - (I)(r) \quad \text{--- (11.40)}$$

The negative sign is due to the fact that the current I flows through the emf device from the negative terminal to the positive terminal.

By the application of Ohm's law Eq. (11.9),

$$V = IR$$

$$\text{Hence } IR = \varepsilon - Ir \quad \text{--- (11.41)}$$

Or

$$I = \frac{\varepsilon}{R + r} \quad \text{--- (11.42)}$$

Thus, the maximum current that can be drawn from the emf device is when $R = 0$, i.e.

$$I_{\max} = \frac{\varepsilon}{r} \quad \text{--- (11.43)}$$

This is the maximum allowed current from an emf device (or a cell). This decides the maximum current rating of a cell or a battery.

11.12 Cells in Series:

In a series combination, cells are connected in single electrical path, such that the positive terminal of one cell is connected to the negative terminal of the next cell, and so on. The terminal voltage of battery/cell is equal to the sum of voltages of individual cells in series, as shown in Fig 11.13 a.

Figure shows two 1.5V cells in series. This combination provides total voltage of 3.0V (1.5×2).

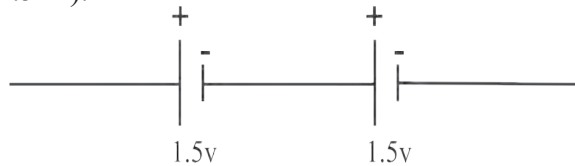


Fig. 11.13 (a): Cells in parallel.

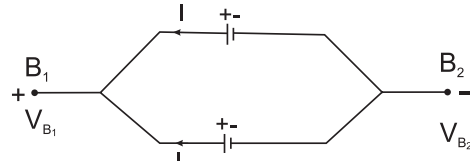


Fig. 11.13 (b): Cells in parallel.

The equivalent emf of n number of cells in series combination is the algebraic sum of their individual emf. The equivalent internal resistance of n cells in a series combination is the sum of their individual internal resistance.

$$V = \sum_i \varepsilon_i - I \cdot \sum_i r_i \quad \text{--- (11.44)}$$

• **Advantages of cells in series.**

- (i) The cells connected in series produce a larger resultant voltage.
- (ii) Cells which are damaged can be easily identified, hence can be easily replaced.

11.13 Cells in parallel:

Consider two cells which are connected in parallel. Here, positive terminals of all the cells are connected together and the negative terminals of all the cells are connected together. In parallel connection, the current is divided among the branches i.e. I_1 and I_2 as shown in Fig. 11.13b. Consider points B_1 and B_2 having potentials V_{B_1} and V_{B_2} , respectively.

For the first cell the potential difference across its terminals is,

$$V = V_{B_1} - V_{B_2} = \varepsilon_1 - I_1 r_1 \quad \text{--- (11.45)}$$

$$\therefore I_1 = \frac{\varepsilon_1 - V}{r_1} \quad \text{--- (11.46)}$$

Point B_1 and B_2 are connected exactly similarly to the second cell.

Hence, considering the second cell we write,

$$V = V_{B_1} - V_{B_2} = \varepsilon_2 - I_2 r_2; I_2 = \frac{\varepsilon_2 - V}{r_2} \quad \text{--- (11.47)}$$

We know that $I = I_1 + I_2$

Combining the last three equations,

$$\therefore I = \frac{\varepsilon_1}{r_1} - \frac{V}{r_1} + \frac{\varepsilon_2}{r_2} - \frac{V}{r_2} \quad \text{--- (11.48)}$$

$$= \left(\frac{\varepsilon_1}{r_1} + \frac{\varepsilon_2}{r_2} \right) - V \left(\frac{1}{r_1} + \frac{1}{r_2} \right)$$

$$\text{Thus, } V \left(\frac{1}{r_1} + \frac{1}{r_2} \right) = \left(\frac{\varepsilon_1}{r_1} + \frac{\varepsilon_2}{r_2} \right) - I$$

$$\therefore V \left(\frac{r_1 + r_2}{r_1 r_2} \right) = \frac{\varepsilon_1 r_2 + \varepsilon_2 r_1}{r_1 r_2} - I$$

$$V = \frac{\varepsilon_1 r_2 + \varepsilon_2 r_1}{r_1 + r_2} - I \frac{r_1 r_2}{r_1 + r_2} \quad \text{--- (11.49)}$$

If we replace the cells by a single cell connected between points B_1 and B_2 with the emf ε_{eq} and the internal resistance r_{eq} as in Fig. (11.13b),

then,

$$V = \varepsilon_{eq} - I r_{eq} \quad \text{--- (11.50)}$$

Considering Eq. (11.49) and Eq. (11.50) we can write,

$$\varepsilon_{eq} = \frac{\varepsilon_1 r_2 + \varepsilon_2 r_1}{r_1 + r_2}$$

$$r_{eq} = \frac{r_1 r_2}{r_1 + r_2}$$

$$\text{i.e. } \frac{1}{r_{eq}} = \frac{1}{r_1} + \frac{1}{r_2}$$

$$\frac{\varepsilon_{eq}}{r_{eq}} = \frac{\varepsilon_1}{r_1} + \frac{\varepsilon_2}{r_2} \quad \text{--- (11.51)}$$

For n number of cells connected in parallel with emf $\varepsilon_1, \varepsilon_2, \varepsilon_3, \dots, \varepsilon_n$ and internal resistance $r_1, r_2, r_3, \dots, r_n$

$$\frac{1}{r_{eq}} = \frac{1}{r_1} + \frac{1}{r_2} + \frac{1}{r_3} + \dots + \frac{1}{r_n} \quad \text{--- (11.52)}$$

$$\text{and } \frac{\varepsilon_{eq}}{r_{eq}} = \frac{\varepsilon_1}{r_1} + \frac{\varepsilon_2}{r_2} + \dots + \frac{\varepsilon_n}{r_n}$$

$$\text{--- (11.53)}$$

Substitution of emfs should be done algebraically by considering proper $\pm$ signs according to polarity.

- **Advantages of cells in parallel :** For cells connected in parallel in a circuit, the circuit will not break open even if a cell gets damaged or open.
- **Disadvantages of cells in parallel :** The voltage developed by the cells in parallel connection cannot be increased by increasing number of cells present in circuit.

11.14 Types of Cells:

Electrical cells can be divided into several categories like primary cell, secondary cell, fuel cell, etc.

A primary cell cannot be charged again. It can be used only once. Dry cells, alkaline cells are different examples of primary cells. Primary cells are low cost and can be used easily. But these are not suitable for heavy loads. Secondary cells are used for such applications. The secondary cell are rechargeable and can be reused. The chemical reaction in a secondary cells is reversible. Lead acid cell, and fuel cell

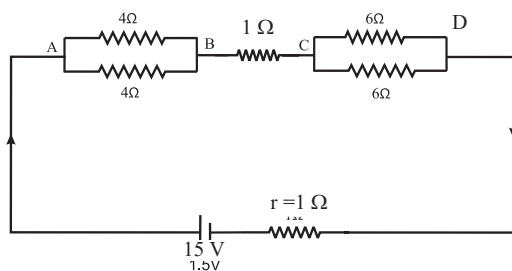
are some examples of secondary cells. Lead acid battery is used widely in vehicles and other applications which require high load currents. Solar cells are secondary cells that convert Solar energy into electrical energy.

Fuel cells vehicles (FCVs) are electric vehicles that use fuel cells instead of lead acid batteries to power the vehicles. Hydrogen is used as a fuel in fuel cells. The by-product after its burning is water. This is important in terms of reducing emission of greenhouse gases produced by traditional gasoline fueled vehicles. The hydrogen fuel cell vehicles are thus more environment friendly.

Example 11.8: A network of resistors is connected to a 15 V battery with internal resistance 1 Ω as shown in the circuit diagram.

Calculate

- The equivalent resistance,
- Current in each resistor,
- Voltage drops V_{AB} , V_{BC} and V_{DC} .



Solution :

- i) Equivalent Resistance (R_{eq}) = $R_{AB} + R_{BC} + R_{DC}$

$$R_{AB} = \frac{4 \times 4}{4 + 4} = 2\Omega, \quad R_{CD} = \frac{6 \times 6}{6 + 6} = 3\Omega$$

$$R_{BC} = 1\Omega$$

$$R_T = R_{eq} = 2 + 1 + 3 = 6\Omega$$

$$\therefore \text{Equivalent Resistance is } 6\Omega$$

- ii. Current in each resistor :

Total current I in the circuit is,

$$I = \frac{\epsilon}{R_T + r} = \frac{15}{6 + 1} = 2.1\text{A}$$

Consider resistors between A and B.

Let I_1 be the current through one of the 4 Ω resistors and I_2 be the current in the other resistor

$$I_1 \times 4 = I_2 \times 4$$

that is, $I_1 = I_2$ from symmetry of the two arms.

$$\text{But } I_1 + I_2 = I = 2.1\text{A}$$

$$\therefore I_1 = I_2 = 1.05\text{A}$$

that is, the current in each 4 Ω resistor is 1.05A, the current in 1 Ω resistor between B and C would be 2.1A.

Now, consider the resistances between C and D

Let I_3 be the current through one of the 6 Ω resistors and I_4 be the current in the other resistor.

$$I_3 \times 6 = I_4 \times 6$$

$$\therefore I_3 = I_4 = 1.05\text{A}$$

That is, current in each 6 Ω resistor is 1.05A

- iii. Voltage drop across BC is V_{BC}

$$V_{BC} = I \times 1 = 2.1 \times 2.1 = 2\text{ V}$$

Voltage drop across CD is V_{CD}

$$V_{CD} = I \times R_{CD} = 2.1 \times 3 = 6.3\text{V}$$

[Note : Total voltage drop across AD is (4.2 V + 2.1 V + 6.3 V) = 12.6 V, while its emf is 15 V. The loss of the voltage is 2.4 V].



Internet my friend

<https://www.britannica.com/science/superconductivityphysics>



Exercises

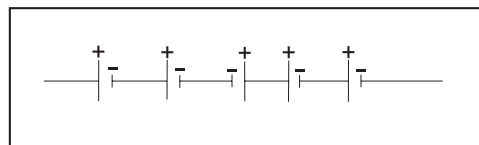
1. Choose correct alternative

- i) You are given four bulbs of 25 W, 40 W, 60 W and 100 W of power, all operating at 230 V. Which of them has the lowest resistance?
(A) 25 W (C) 40 W
(C) 60 W (D) 100 W
- ii) Which of the following is an ohmic conductor?
(A) transistor (B) vacuum tube
(C) electrolyte (D) nichrome wire
- iii) A rheostat is used
(A) to bring on a known change of resistance in the circuit to alter the current
(B) to continuously change the resistance in any arbitrary manner and thereby alter the current
(C) to make and break the circuit at any instant
(D) neither to alter the resistance nor the current
- iv) The wire of length L and resistance R is stretched so that its radius of cross-section is halved. What is its new resistance?
(A) $5R$ (B) $8R$
(C) $4R$ (D) $16R$
- v) Masses of three pieces of wires made of the same metal are in the ratio 1:3:5 and their lengths are in the ratio 5:3:1. The ratios of their resistances are
(A) 1:3:5 (B) 5:3:1
(C) 1:15:125 (D) 125:15:1
- vi) The internal resistance of a cell of emf 2V is 0.1Ω it is connected to a resistance of 0.9Ω . The voltage across the cell will be
(A) 0.5 V (B) 1.8 V
(C) 1.95 V (D) 3V
- vii) 100 cells each of emf 5V and internal resistance 1Ω are to be arranged so as to produce maximum current in a 25Ω resistance. Each row contains equal

number of cells. The number of rows should be

- (A) 2 (B) 4
(C) 5 (D) 100

- viii) Five dry cells each of voltage 1.5 V are connected as shown in diagram

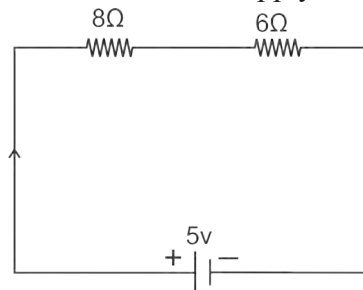


What is the overall voltage with this arrangement?

- (A) 0V (B) 4.5V
(C) 6.0V (D) 7.5V

2. Give reasons / short answers

- i) In given circuit diagram two resistors are connected to a 5V supply.



a] Calculate potential difference across the 8Ω resistor.

b] A third resistor is now connected in parallel with 6Ω resistor. Will the potential difference across the 8Ω resistor be larger, smaller or the same as before? Explain the reason for your answer.

- ii) Prove that the current density of a metallic conductor is directly proportional to the drift speed of electrons.

3. Answer the following questions.

- i) Distinguish between Ohmic and non-ohmic substances; explain with the help of example.
- ii) DC current flows in a metal piece of non-uniform cross-section. Which of these quantities remains constant along the conductor: current, current density or drift speed?

4. Solve the following problems.

- i) What is the resistance of one of the rails of a railway track 20 km long at 20°C ? The cross section area of rail is 25 cm^2 and the rail is made of steel having resistivity at 20°C as $6\times 10^{-8}\ \Omega\text{ m}$.
[Ans: $0.48\ \Omega$]
- ii) A battery after a long use has an emf 24 V and an internal resistance $380\ \Omega$. Calculate the maximum current drawn from the battery? Can this battery drive starting motor of car?
[Ans: 0.063 A]
- iii) A battery of emf 12 V and internal resistance $3\ \Omega$ is connected to a resistor. If the current in the circuit is 0.5 A ,
a] Calculate resistance of resistor.
b] Calculate terminal voltage of the battery when the circuit is closed.
[Ans: a) $21\ \Omega$, b) 10.5 V]
- iv) The magnitude of current density in a copper wire is 500 A/cm^2 . If the number of free electrons per cm^3 of copper is 8.47×10^{22} calculate the drift velocity of the electrons through the copper wire (charge on an $e = 1.6\times 10^{-19}\text{ C}$)
[Ans: $3.69\times 10^{-4}\text{ m/s}$]
- v) Three resistors $10\ \Omega$, $20\ \Omega$ and $30\ \Omega$ are connected in series combination.
i] Find equivalent resistance of series combination.
ii] When this series combination is connected to 12V supply, by neglecting the value of internal resistance, obtain potential difference across each resistor.
[Ans: i) $60\ \Omega$, ii) 2 V , 4 V , 6 V]
- vi) Two resistors $1\text{ k}\Omega$ and $2\text{ k}\Omega$ are connected in parallel combination.
i] Find equivalent resistance of parallel combination
ii] When this parallel combination is connected to 9 V supply, by neglecting internal resistance calculate current through each resistor.
[Ans: i) $0.66\text{ k}\Omega$, ii) 9 mA , 4.5 mA]
- vii) A silver wire has a resistance of $4.2\ \Omega$ at 27°C and resistance $5.4\ \Omega$ at 100°C . Determine the temperature coefficient of resistance.
[Ans: $3.91\times 10^{-3}/^{\circ}\text{C}$]
- viii) A 6m long wire has diameter 0.5 mm. Its resistance is $50\ \Omega$. Find the resistivity and conductivity.
[Ans: $1.636\times 10^{-6}\ \Omega/\text{m}$, $6.112\times 10^5\text{ m}/\Omega$]
- ix) Find the value of resistances for the following colour code.
1. Blue Green Red Gold
[Ans: $6.5\text{ k}\Omega \pm 5\%$]
2. Brown Black Red Silver
[Ans: $1.0\text{ k}\Omega \pm 10\%$]
3. Red Red Orange Gold
[Ans: $2.2\text{ k}\Omega \pm 5\%$]
4. Orange White Red Gold
[Ans: $3.9\text{ k}\Omega \pm 5\%$]
5. Yellow Violet Brown Silver
[Ans: $4.70\text{ k}\Omega \pm 10\%$]
- x) Find the colour code for the following value of resistor having tolerance $\pm 10\%$
a) 330Ω b) 100Ω c) $47\text{ k}\Omega$
d) 160Ω e) $1\text{ k}\Omega$
- xi) A current 4A flows through an automobile headlight. How many electrons flow through the headlight in a time 2hrs.
[Ans : 1.8×10^{23}]
- xii) The heating element connected to 230V draws a current of 5A. Determine the amount of heat dissipated in 1 hour ($J = 4.2\text{ J/cal.}$).
[Ans : 985.7 kcal]

**Can you recall?**

1. What is a bar magnet?
2. What are the magnetic lines of force?
3. What are the rules concerning the lines of force?
4. If you freely hang a bar magnetic horizontally, in which direction will it become stable?

12.1 Introduction:

The history of magnetism dates back to earlier than 600 B.C., but it is only in the twentieth century that scientists began to understand it and developed technologies based on this understanding. William Gilbert (1544-1603) was the first to systematically investigate the phenomenon of magnetism using scientific method. He also discovered that Earth is a weak magnet. Danish physicist Hans Oersted (1777-1851) suggested a link between electricity and magnetism. James Clerk Maxwell (1831-1879) proved that electricity and magnetism represent different aspects of the same fundamental force field.

In electrostatics you have learnt about the relationship between the electric field and force due to electric charges and electric dipoles. Analogous concepts exist in magnetism except that magnetic poles do not exist in isolation, and we always have a magnetic dipole or a

quadrupole. In this Chapter the main focus will be on elementary aspects of magnetism and terrestrial magnetism.

12.2 Magnetic Lines of Force and Magnetic Field:

You have studied properties of electric lines of force earlier in the Chapter on electrostatics. In a similar manner, magnetic lines of force originate from the north pole and end at the south pole of a bar magnet. The magnetic lines of force of a magnet have the following properties:

- i) The magnetic lines of force of a magnet or a solenoid form closed loops. This is in contrast to the case of an electric dipole, where the electric lines of force originate from the positive charge and end on the negative charge, without forming a complete loop (see Fig. 12.4).
- ii) The direction of the net magnetic field $\vec{B}$ at a point is given by the tangent to the magnetic line of force at that point in the direction of line of force.
- iii) The number of lines of force crossing per unit area decides the magnitude of the magnetic field $\vec{B}$.
- iv) The magnetic lines of force do not intersect. This is because had they intersected, the direction of magnetic field would not be unique at that point.

**Do you know ?****Some commonly known facts about magnetism.**

- (i) Every magnet regardless of its size and shape has two poles called north pole and south pole.
- (ii) If a magnet is broken into two or more pieces then each piece behaves like an independent magnet with somewhat weaker magnetic field.
Thus isolated magnetic monopoles do not exist. The search for magnetic monopoles is still going on.
- (iii) Like magnetic poles repel each other, whereas unlike poles attract each other.
- (iv) When a bar magnet/ magnetic needle is suspended freely or is pivoted, it aligns itself in geographically North-South direction.

**Try this**

You can take a bar magnet and a small compass needle. Place the bar magnet at a fixed position on a paper and place the needle at various positions. Noting the orientation of the needle, the magnetic field direction at various locations can be traced.

Density of lines of force i.e., the number of lines of force per unit area normal to the surface

around a particular point determines the strength of the magnetic field at that point. The number of lines of force is called magnetic flux (ϕ). SI unit of magnetic flux (ϕ) is weber (Wb). For a specific case of uniform magnetic field which is normal to the finite area A , the magnitude of magnetic field strength B at a point in the area A is given by

$$\text{Magnetic Field} = \frac{\text{magnetic flux}}{\text{area}}$$

i.e. $B = \frac{\phi}{A}$ --- (12.1)

SI unit of magnetic field (B) is expressed as weber/m² or Tesla.

$$1 \text{ Tesla} = 10^4 \text{ Gauss.}$$

However, magnetic lines are only a crude way of representing magnetic field. It is a pictorial representation of the strength of the magnetic field (B). It is better defined in terms of Lorentz force law which you will learn in std XII.

12.3 The Bar magnet:

A bar magnet is said to have magnetic pole strength $+q_m$ and $-q_m$ at the north and south poles, respectively. The separation of magnetic poles inside the magnet is $2l$. As the bar magnet has two poles, with equal and opposite pole strength, it is called a magnetic dipole. This is analogous to an electric dipole. The magnetic dipole moment, therefore, becomes $\vec{m} = q_m \cdot 2\vec{l}$ ($2\vec{l}$ is a vector from south pole to north pole) in analogy with the electric dipole moment.

SI unit of pole strength (q_m) is A m.

SI unit of magnetic dipole moment m is A m².

Axis:- It is the line passing through both the poles of a bar magnet. Obviously, there is only one axis for a given bar magnet.

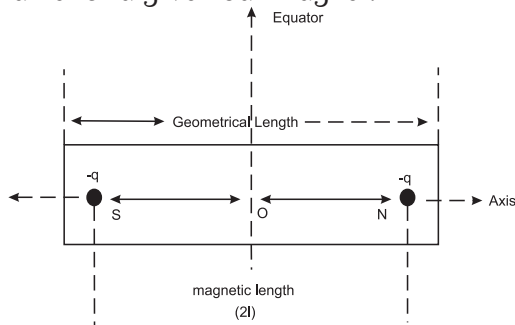


Fig. 12.1: Bar magnet

Equator:- A line passing through the centre of a magnet and perpendicular to its axis is called magnetic equator. The plane containing all equators is called the equatorial plane. The locus of points, on the equatorial plane, which are equidistant from the centre of the magnet is called the equatorial circle. The popularly known 'equator' in Geography is actually an 'equatorial circle'. Such a circle with any diameter is an equator.

Magnetic length (2l):- It is the distance between the two poles of a magnet.

$$\text{Magnetic length (2l)} = \frac{5}{6} \times \text{Geometric length}$$

--- (12.2)

12.3.1 Magnetic field due to a bar magnet at a point along its axis and at a point along its equator:

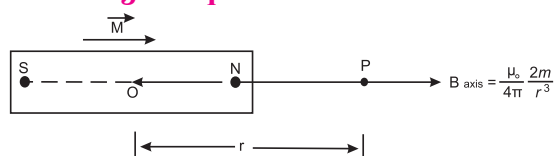


Fig. 12.2 (a): Magnetic field at a point along the axis of the magnet.

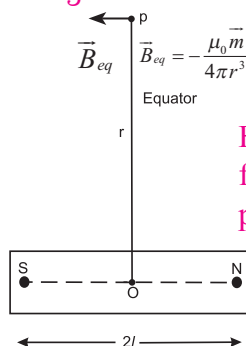


Fig. 12.2 (b): Magnetic field along the equatorial point.

Consider a bar magnet of dipole length $2l$ and magnetic dipole moment $\vec{m}$ as shown in Fig. 12.2 (a). We will now find magnetic field at a point P along the axis of the bar magnet.

Let r be the distance of point P from the centre O of the magnetic dipole.

$$OS = ON = l$$

$$\therefore NP = SP = \sqrt{r^2 + l^2}$$

We now use the electrostatic analogy to obtain the magnetic field due to a bar magnet at a large distance $r \gg l$. Consider the electric field due to an electric dipole with a dipole moment p .

The Electrostatic Analogue:

As suggested by Maxwell, electricity and magnetism could be studied analogously. The pole strength (q_m) in magnetism can be

compared with charge q in electrostatics. Accordingly, we can write the equivalent physical quantities in electrostatics and magnetism as shown in table 12.1.

Table 12.1: The Electrostatic Analogue

Quantity	Electrostatics	Magnetism
Basic physical quantity	Electrostatic charge	Magnetic pole
Field	Electric Field $\vec{E}$	Magnetic Field $\vec{B}$
Constant	$\frac{1}{4\pi\epsilon_0}$	$\frac{\mu_0}{4\pi}$
Dipole moment	$\vec{p} = q (2\vec{l})$ along (-ve) $\rightarrow$ (+ve) charge	$\vec{m} = q_m (2\vec{l})$ (bar magnet) along S $\rightarrow$ N pole
Force	$\vec{F} = q \vec{E}$	$\vec{F} = q_m \vec{B}$
Energy (In external field) of a dipole	$U = - \vec{p} \cdot \vec{E}$	$U = - \vec{m} \cdot \vec{B}$
Coulomb's law	$F = \left(\frac{1}{4\pi\epsilon_0} \right) \frac{q_1 q_2}{r^2}$	No analogous law as magnetic monopoles do not exist
Axial field for a short dipole	$\frac{2p}{4\pi\epsilon_0 r^3}$ along $\vec{p}$	$\frac{\mu_0 2\vec{m}}{4\pi r^3}$
Equatorial field for a short dipole	$\frac{p}{4\pi\epsilon_0 r^3}$ opposite to $\vec{p}$	$\frac{-\mu_0 \vec{m}}{4\pi r^3}$

You have studied the electric field due to an electric dipole of length $2l$ ($p=2ql$) at a distance r along the dipolar axis (Eq. 10.24) which is given by,

$$|\vec{E}_a| = \frac{1}{4\pi\epsilon_0} \frac{2p}{r^3}, \quad r \gg l$$

The electric field on the equator (Eq. 10.28) is antiparallel to $\vec{p}$ and is given by

$$|\vec{E}_{eq}| = \frac{1}{4\pi\epsilon_0} \frac{p}{r^3}, \quad r \gg l$$

Using the analogy given in Table 12.1, we can thus write the axial magnetic field of a bar magnet at a distance r , $r \gg l$, $2l$ being the length of bar magnet,

$$\vec{B}_a = + \frac{\mu_0}{4\pi} \frac{2\vec{m}}{r^3} \quad \text{--- (12.3)}$$

Similarly, the equatorial magnetic field

$$\vec{B}_{eq} = - \frac{\mu_0 \vec{m}}{4\pi r^3} \quad \text{--- (12.4)}$$

Negative sign shows that the direction of $\vec{B}_{eq}$ is opposite to $\vec{m}$.

For the same distance from centre O of a bar magnet,

$$B_{axis} = 2B_{eq} \quad \text{--- (12.5)}$$

12.3.2 Magnetic field due to a bar magnet at an arbitrary point:

Fig. 12.3 Shows a bar magnet of magnetic moment $\vec{m}$ with centre at O. P is any point in its magnetic field. Magnetic moment $\vec{m}$ is resolved (*about the centre of the magnet*) into components along $\vec{r}$ and perpendicular to $\vec{r}$. For the component $m \cos \theta$ along $\vec{r}$, the point P is an axial point.

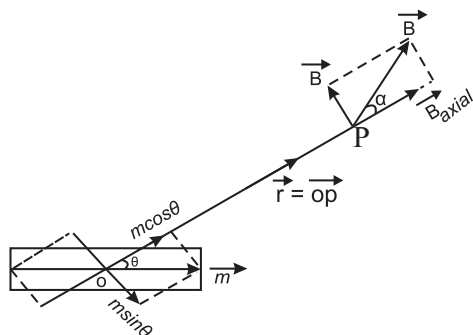


Fig. 12.3: Magnetic field at an arbitrary point.

Also, for the component $m \sin \theta$ perpendicular to $\vec{r}$, the point P is an equatorial point at the same distance $\vec{r}$. Using the results of axial and equatorial fields, we get

$$B_a = \frac{\mu_0}{4\pi} \frac{2m \cos \theta}{r^3} \quad \text{--- (12.6)}$$

directed along $m \cos \theta$ and

$$B_{eq} = \frac{\mu_0}{4\pi} \frac{m \sin \theta}{r^3} \quad \text{--- (12.7)}$$

directed opposite to $m \sin \theta$

Thus, the magnitude of the resultant magnetic field B , at point P is given by

$$B = \sqrt{B_a^2 + B_{eq}^2}$$

$$\therefore B = \frac{\mu_0}{4\pi} \frac{m}{r^3} \sqrt{[2 \cos \theta]^2 + [\sin \theta]^2}$$

$$\therefore B = \frac{\mu_0}{4\pi} \frac{m}{r^3} \sqrt{3 \cos^2 \theta + 1} \quad \text{--- (12.8)}$$

Let α be the angle made by the direction of $\vec{B}$ with $\vec{r}$. Then, by using eq (12.6) and eq (12.7),

$$\tan \alpha = \frac{B_{eq}}{B_a} = \frac{1}{2} (\tan \theta) \quad \text{--- (12.9)}$$

The angle between directions of $\vec{B}$ and $\vec{m}$ is then $(\theta + \alpha)$.

Example 12.1: A short magnetic dipole has magnetic moment 0.5 A m^2 . Calculate its magnetic field at a distance of 20 cm from the centre of magnetic dipole on (i) the axis (ii) the equatorial line (Given $\mu_0 = 4\pi \times 10^{-7} \text{ SI units}$)

Solution :

$$m = 0.5 \text{ A m}^2, \quad r = 20 \text{ cm} = 0.2 \text{ m}$$

$$\begin{aligned} B_a &= \frac{\mu_0}{4\pi} \frac{2m}{r^3} = \frac{10^{-7} \times 2 \times 0.5}{(0.2)^3} = \frac{1 \times 10^{-7}}{8 \times 10^{-3}} \\ &= \frac{1}{8} \times 10^{-4} = 1.25 \times 10^{-5} \text{ Wb / m}^2 \end{aligned}$$

$$\begin{aligned} B_{eq} &= \frac{\mu_0}{4\pi} \frac{m}{r^3} \\ &= \frac{10^{-7} \times 0.5}{(0.2)^3} = \frac{5 \times 10^{-8}}{8 \times 10^{-3}} = 0.625 \times 10^{-5} \text{ Wb / m}^2 \end{aligned}$$

12.4 Gauss' Law of Magnetism:

The Gauss' law for electric field is known to you. It states that the net electric flux through a closed Gaussian surface is proportional to the net electric charge enclosed by the surface (Eq. (10.18)). The Gauss' law for magnetic fields states that the net magnetic flux Φ_B through a closed Gaussian surface is zero, i.e.,

$$\Phi_B = \int \vec{B} \cdot d\vec{S} = 0 \quad (\text{Gauss' law for magnetic fields})$$

The magnetic force lines of (a) bar magnet, (b) current carrying finite solenoid, and (c) electric dipole are shown in Fig.12.4(a), 12.4(b) and 12.4(c), respectively. The curves labelled (i) and (ii) are cross sections of three dimensional closed Gaussian surfaces.

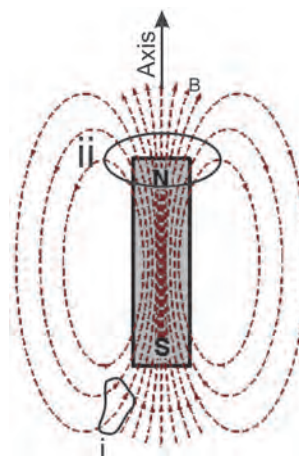


Fig. 12.4 (a): Bar magnet.

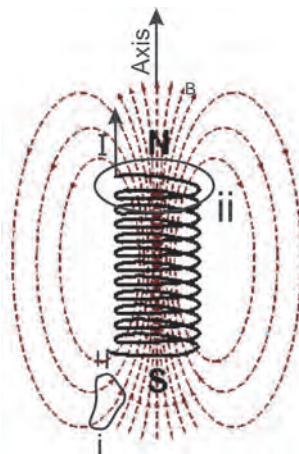


Fig. 12.4 (b): Current (I) carrying solenoid.

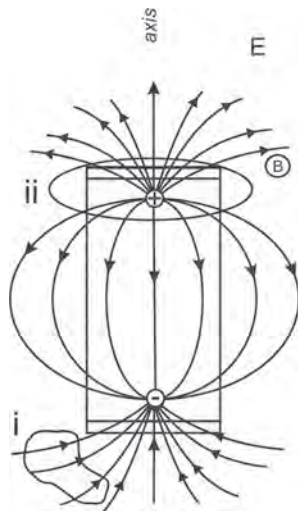


Fig. 12.4 (c): Electric dipole.

If we compare the number of lines of force entering in and leaving out of the surface (i), it is clearly seen that they are equal. The Gaussian surface does not include poles. It means that the flux associated with any closed surface is equal to zero. When we consider surface (ii), in Fig. 12.4 (b), we are enclosing the North pole. As even a thin slice of a bar magnet will have North and South poles associated with it, the closed Gaussian surface will also include a South pole. However in Fig. 12.4(c), for an electric dipole, the field lines begin from positive charge and end on negative charge. For a closed surface (ii), there is a net outward flux since it does include a net (positive) charge. According to the Gauss' law of electrostatics as studied earlier, $\Phi_E = \int \vec{E} \cdot d\vec{S} = \frac{q}{\epsilon_0}$, where q is the positive charge enclosed. Thus, situation is entirely different from magnetic lines of force, which are shown in Fig. 12.4(a) and Fig. 12.4(b). Thus, Gauss' law of magnetism can be written as $\Phi_B = \int \vec{B} \cdot d\vec{S} = 0$.

From the above we conclude that for electrostatics, an isolated electric charge exists but an isolated magnetic pole does not exist. In short, only dipoles exist in case of magnetism.

12.5 Earth's Magnetism:

It is common experience that a bar magnet or a magnetic needle suspended freely in air always aligns itself along geographic N-S direction. If it has a freedom to rotate about horizontal axis, it inclines with some angle with the horizontal in the vertical N-S plane.

This fact clearly indicates that there is some magnetic field present everywhere on the Earth. This is called Terrestrial Magnetism. It is extremely useful during navigation.

Magnetic parameters of the Earth are described below. The magnetic lines of force enter the Earth's surface at the north pole and emerge from the south pole.

Unless and otherwise stated, the directions mentioned (South, North, etc.) are always, Geographic.

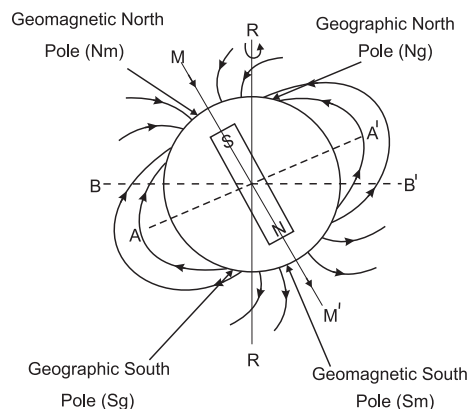


Fig. 12.5: Earth's magnetism.

Magnetic Axis :- The Earth is considered to be a huge magnet. Magnetic north pole (N) of the Earth is located below Antarctica while the south pole (S) is below north Canada. The straight line NS joining these two poles is called the magnetic axis, MM'.

Magnetic equator :- A great circle in the plane perpendicular to magnetic axis is magnetic equatorial circle, AA'. It happens to pass through India near Thiruvananthapuram.

Geographic Meridian:- A plane perpendicular to the surface of the Earth (vertical plane) perpendicular to geographic axis is geographic meridian. (Fig.12.6)

Magnetic Meridian:- A plane perpendicular to surface of the Earth (Vertical plane) and passing through the magnetic axis is magnetic meridian. Direction of resultant magnetic field of the Earth is always along or parallel to magnetic meridian. (Fig.12.6)

Magnetic declination:- Angle between the geographic and the magnetic meridian at a place is called 'magnetic declination' (α). The declination is small in India. It is $0^\circ 58'$ west at Mumbai and $0^\circ 41'$ east at Delhi. Thus, at both these places, magnetic needle shows true North

accurately (Fig.12.6).

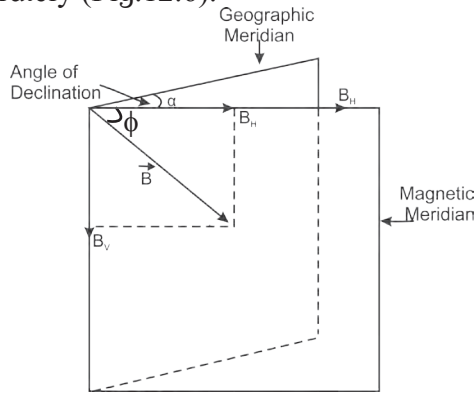


Fig. 12.6: Magnetic declination.

Magnetic inclination or angle of dip (ϕ):- Angle made by the direction of resultant magnetic field with the horizontal at a place is inclination or angle of dip at the place (Fig. 12.7).

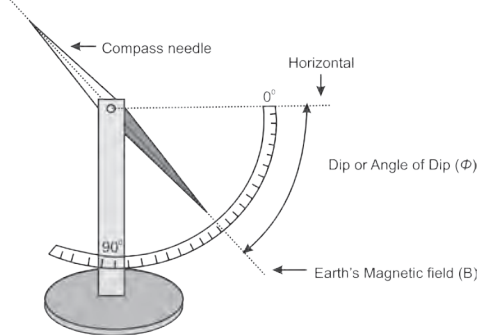


Fig. 12.7: Magnetic inclination.

Earth's magnetic field:- Magnetic force experienced per unit pole strength is magnetic field $\vec{B}$ at that place. It can be resolved in components along the horizontal, $\vec{B}_H$ and along vertical, $\vec{B}_V$. The vertical component can be conveniently determined. The two components can be related with the angle of dip (ϕ) as,

$$B_H = B \cos \phi, B_V = B \sin \phi$$

$$\frac{B_V}{B_H} = \tan \phi \quad \text{--- (12.10)}$$

$$B^2 = B_V^2 + B_H^2$$

$$\therefore B = \sqrt{B_V^2 + B_H^2} \quad \text{--- (12.11)}$$

Special cases

- 1) At the magnetic North pole, $\vec{B} = \vec{B}_V$, directed upward, $\vec{B}_H = 0$ and $\phi = 90^\circ$.
- 2) At the magnetic south pole, $\vec{B} = \vec{B}_V$, directed downward, $\vec{B}_H = 0$ and $\phi = 270^\circ$.
- 3) Anywhere on the magnetic great circle

(magnetic equator) $B = B_H$ along South to North, $B_V = 0$ and $\phi = 0$

Magnetic maps of the Earth:-

Magnetic elements of the Earth (B_H , α and ϕ) vary from place to place and also with time. The maps providing these values at different locations are called magnetic maps. These are extremely useful for navigation. Magnetic maps drawn by joining places with the same value of a particular element are called Iso-magnetic charts.

Lines joining the places of equal horizontal components (B_H) are known as 'Isodynamic lines'

Lines joining the places of equal declination (α) are called Isogonic lines.

Lines joining the places of equal inclination or dip (ϕ) are called Aclinic lines.

Example 12.2: Earth's magnetic field at the equator is approximately 4×10^{-5} T. Calculate Earth's dipole moment. (Radius of Earth = 6.4×10^6 m, $\mu_0 = 4\pi \times 10^{-7}$ SI units)

Solution: Given

$$B_{eq} = 4 \times 10^{-5} \text{ T}$$

$$r = 6.4 \times 10^6 \text{ m}$$

Assume that Earth is a bar magnet with N and S poles being the geographical South and North poles, respectively. The equatorial magnetic field due to Earth's dipole can be written as

$$B_{eq} = \frac{\mu_0 m}{4\pi r^3}$$

$$m = 4\pi B_{eq} \times r^3 / \mu_0$$

$$= 4 \times 10^{-5} \times (6.4 \times 10^6)^3 \times 10^7$$

$$= 1.05 \times 10^{20} \text{ A m}^2$$

Example 12.3: At a given place on the Earth, a bar magnet of magnetic moment $\vec{m}$ is kept horizontal in the East-West direction. P and Q are the two neutral points due to magnetic field of this magnet and $\vec{B}_H$ is the horizontal component of the Earth's magnetic field.

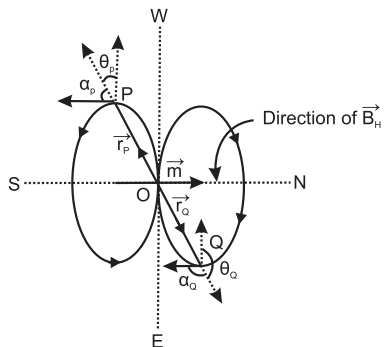
(A) Calculate the angles between position vectors of P and Q with the direction of $\vec{m}$.

(B) Points P and Q are 1 m from the centre of the bar magnet and $B_H = 3.5 \times 10^{-5}$ T. Calculate

magnetic dipole moment of the bar magnet.

Neutral point is that point where the resultant magnetic field is zero.

Solution: (A) As seen from the figure, the direction of magnetic field $\vec{B}$ due to the bar magnet is opposite to $\vec{B}_H$ at the points P and Q. Also, $(\theta + \alpha) = 90^\circ$ at P and it is 270° at Q.



$$\tan \alpha = \frac{1}{2} \tan \theta$$

$$\therefore \tan \theta = 2 \tan \alpha$$

$$= 2 \tan(90^\circ - \theta) \text{ and } 2 \tan(270^\circ - \theta)$$

$$\therefore \tan \theta = \pm 2 \cot \theta$$

$$\therefore \tan^2 \theta = 2$$

$$\therefore \tan \theta = \pm \sqrt{2}$$

$$\therefore \theta = \tan^{-1}(\pm \sqrt{2})$$

$$\therefore \theta = 54^\circ 44' \text{ and } 180^\circ - 54^\circ 44' = 116^\circ 4'$$

(B)

$$\tan^2 \theta = 2$$

$$\therefore \sec^2 \theta = 1 + \tan^2 \theta = 1 + 2 = 3$$

$$\therefore \cos^2 \theta = \frac{1}{3}$$

$$r = 1 \text{ m and } B = B_H = 3.5 \times 10^{-5} \text{ T} \quad (\text{Given})$$

we have,

$$B = \frac{\mu_0}{4\pi} \frac{m}{r^3} \sqrt{3 \cos^2 \theta + 1}$$

$$\therefore m = \frac{B_H \times r^3}{\left(\frac{\mu_0}{4\pi}\right) \sqrt{3 \cos^2 \theta + 1}}$$

$$= \frac{3.5 \times 10^{-5} \times 1^3}{10^{-7} \times \sqrt{3 \frac{1}{3} + 1}}$$

$$\therefore m = \frac{350}{\sqrt{2}} = 247.5 \text{ A m}^2$$

Always remember:

In this Chapter we have used B as a symbol for magnetic field. Calling it magnetic induction is unreasonable. We have used the words **magnetic field** which are used in spoken language.



Exercises

1. Choose the correct option.

- Let r be the distance of a point on the axis of a bar magnet from its center. The magnetic field at r is always proportional to
(A) $1/r^2$ (B) $1/r^3$ (C) $1/r$
(D) not necessarily $1/r^3$ at all points
- Magnetic meridian is the plane
(A) perpendicular to the magnetic axis of Earth
(B) perpendicular to geographic axis of Earth
(C) passing through the magnetic axis of Earth
(D) passing through the geographic axis of Earth

- The horizontal and vertical component of magnetic field of Earth are same at some place on the surface of Earth. The magnetic dip angle at this place will be
(A) 30° (B) 45°
(C) 0° (D) 90°
- Inside a bar magnet, the magnetic field lines
(A) are not present
(B) are parallel to the cross sectional area of the magnet
(C) are in the direction from N pole to S pole
(D) are in the direction from S pole to N pole

- v) A place where the vertical components of Earth's magnetic field is zero has the angle of dip equal to
(A) 0° (B) 45°
(C) 60° (D) 90°
- vi) A place where the horizontal component of Earth's magnetic field is zero lies at
(A) geographic equator
(B) geomagnetic equator
(C) one of the geographic poles
(D) one of the geomagnetic poles
- vii) A magnetic needle kept nonparallel to the magnetic field in a nonuniform magnetic field experiences
(A) a force but not a torque
(B) a torque but not a force
(C) both a force and a torque
(D) neither force nor a torque
- ii) A magnet makes an angle of 45° with the horizontal in a plane making an angle of 30° with the magnetic meridian. Find the true value of the dip angle at the place.
[Ans: $\tan^{-1}(0.866)$]
- iii) Two small and similar bar magnets have magnetic dipole moment of 1.0 Am^2 each. They are kept in a plane in such a way that their axes are perpendicular to each other. A line drawn through the axis of one magnet passes through the center of other magnet. If the distance between their centers is 2 m, find the magnitude of magnetic field at the mid point of the line joining their centers.
[Ans: $\sqrt{5} \times 10^{-7} \text{ T}$]
- iv) A circular magnet is made with its north pole at the centre, separated from the surrounding circular south pole by an air gap. Draw the magnetic field lines in the gap. [The magnet is hypothetical magnet]. Draw a diagram to illustrate the magnetic lines of force between the south poles of two such magnets.
- v) Two bar magnets are placed on a straight line with their north poles facing each other on a horizontal surface. Draw magnetic lines around them. Mark the position of any neutral points (points where there is no resultant magnetic field) on your diagram.

2. Answer the following questions in brief.

- i) What happens if a bar magnet is cut into two pieces transverse to its length/ along its length?
- ii) What could be the equation for Gauss' law of magnetism, if a monopole of pole strength p is enclosed by a surface?

3. Answer the following questions in detail.

- i) Explain the Gauss' law for magnetic fields.
- ii) What is a geographic meridian. How does the declination vary with latitude? Where is it minimum?
- iii) Define the Angle of Dip. What happens to angle of dip as we move towards magnetic pole from magnetic equator?

4. Solve the following Problems.

- i) A magnetic pole of bar magnet with pole strength of 100 A m is 20 cm away from the centre of a bar magnet. Bar magnet has pole strength of 200 A m and has a length 5 cm. If the magnetic pole is on the axis of the bar magnet, find the force on the magnetic pole.

[Ans: $2.5 \times 10^{-2} \text{ N}$]



Internet my friend

<https://www.ngdc.noaa.gov>

13. Electromagnetic Waves and Communication System



Can you recall?

1. What is a wave?
2. What is the difference between longitudinal and transverse waves?
3. What are electric and magnetic fields and what are their sources?
4. What are Lenz's law, Ampere's law and Faraday's law?
5. By which mechanism heat is lost by hot bodies?

13.1 Introduction :

The information age in which we live is based almost entirely on the physics of electromagnetic (EM) waves. We are now globally connected by TV, cellphone and internet. All these gadgets use EM waves as carriers for transmission of signals. Energy from the Sun, an essential requirement for life on Earth, reaches us by means of EM waves that travel through nearly 150 million km of empty space. There are EM waves from light bulbs, heated engine blocks of automobiles, x-ray machines, lightning flashes, and some radioactive materials. Stars, other objects in our milky way galaxy and other galaxies are known to emit EM waves. Hence, it is important for us to make a careful study of the properties of EM waves.

13.2 EM wave:

There are four basic laws which describe the behaviour of electric and magnetic fields, the relation between them and their generation by charges and currents. These laws are as follows.

- (1) Gauss' law for electrostatics, which is essentially the Coulomb's law, describes the relationship between static electric charges and the electric field produced by them.
- (2) Gauss' law for magnetism, which is similar to the Gauss' law for electrostatics mentioned above, states that "magnetic monopoles which are thought to be magnetic charges equivalent to the electric charges, do not exist". Magnetic poles always occur in pairs.
- (3) Faraday's law which gives the relation between electromotive force (emf) induced in a circuit when the magnetic flux linked

with the circuit changes.

- (4) Ampere's law gives the relation between the induced magnetic field associated with a loop and the current flowing through the loop. Maxwell (1831-1879) noticed a major flaw in the Ampere's law for time dependant fields. He noticed that the magnetic field can be generated not only by electric current but also by changing electric field. Therefore in the year 1861, he added one more term to the equation describing this law. This term is called the displacement current. This term is extremely important and the EM waves which are an outcome of these equations would not have been possible in absence of this term.

As a result, the set of four equations describing the above four laws is called Maxwell's equations.

In 1888, H. Hertz (1857-1894) succeeded in producing and detecting the existence of EM waves. He also demonstrated their properties namely reflection, refraction and interference.

In 1895, an Indian physicist Sir Jagdish Chandra Bose (1858-1937) produced EM waves ranging in wavelengths from 5 mm to 25 nm. His work, however, remained confined to laboratory only.

In 1896, an Italian physicist G. Marconi (1874-1937) became pioneer in establishing wireless communication. He was awarded the Nobel prize in physics in 1909 for his work in developing wireless telegraphy, telephony and broadcasting.

13.2.1 Sources of EM waves:

According to Maxwell's theory, "accelerated charges radiate EM waves". Consider a charge oscillating with some frequency. This produces

an oscillating electric field in space, which produces an oscillating magnetic field which in turn is a source of oscillating electric field. Thus varying electric and magnetic fields regenerate each other.

Waves that are caused by the acceleration of charged particles and consist of electric and magnetic fields vibrating sinusoidally at right angles to each other and to the direction of propagation are called EM waves or EM radiation. Figure 13.1 shows an EM wave propagating along z-axis. The time varying electric field is along the x-axis and time varying magnetic field is along the y-axis.

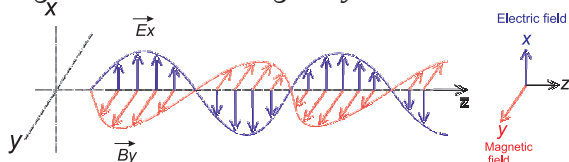


Fig. 13.1: EM wave propagating along z-axis.



Do you know ?

In 1865, Maxwell proposed that an oscillating electric charge radiates energy in the form of EM wave. EM waves are periodic changes in electric and magnetic fields, which propagate through space. Thus, energy can be transported in the form of EM waves.

Maxwell's Equations for Charges and Currents in Vacuum

$$1) \int \vec{E} \cdot d\vec{S} = \frac{Q_{in}}{\epsilon_0} \text{ (Gauss' law)}$$

Here $\vec{E}$ is the electric field and ϵ_0 is the permittivity of vacuum. The integral is over a closed surface S . The law states that electric flux through any closed surface S is equal to the total electric charge Q_{in} enclosed by the surface divided by ϵ_0 . Gauss' law describes the relation between an electric charge and electric field it produces.

$$2) \int \vec{B} \cdot d\vec{S} = 0 \text{ (Gauss' law for magnetism).}$$

Here $\vec{B}$ is the magnetic field. The integral is over a closed surface S . The law states that magnetic flux through a closed surface is always zero, i.e., the magnetic field lines are continuous closed curves, having neither beginning nor end.

$$3) \int \vec{E} \cdot d\vec{l} = -\frac{d\phi_m}{dt}$$

(Faraday's law with Lenz's law)

Here ϕ_m is the magnetic flux and the integral is over a closed loop. Time varying magnetic field induces an electromotive force (emf) and hence, an electric field. The direction of the induced emf is such that the change is opposed.

$$4) \int \vec{B} \cdot d\vec{l} = \mu_0 I + \epsilon_0 \mu_0 \frac{d\phi_E}{dt}$$

(Ampere-Maxwell law)

Here μ_0 is the permeability of vacuum and the integral is over a closed loop, I is the current flowing through the loop. ϕ_E is the electric flux linked with the circuit. Magnetic field is generated by moving charges and also by varying electric fields.

13.2.2 Characteristics of EM waves:

- 1) The electric and magnetic fields, $\vec{E}$ and $\vec{B}$ are always perpendicular to each other and also to the direction of propagation of the EM wave. Thus the EM waves are transverse waves.
- 2) The cross product $\vec{E} \times \vec{B}$ gives the direction in which the EM wave travels. $\vec{E} \times \vec{B}$ also gives the energy carried by EM wave.
- 3) The $\vec{E}$ and $\vec{B}$ fields vary sinusoidally and are in phase.
- 4) EM waves are produced by accelerated electric charges.
- 5) EM waves can travel through free space as well as through solids, liquids and gases.
- 6) In free space, EM waves travel with velocity c , equal to that of light in free space.

$$c = \frac{1}{\sqrt{\mu_0 \epsilon_0}} = 3 \times 10^8 \text{ m/s},$$

where μ_0 ($4\pi \times 10^{-7} \text{ Tm/A}$) is permeability and ϵ_0 ($8.85 \times 10^{-12} \text{ C}^2/\text{Nm}^2$) is permittivity of free space.

- 7) In a given material medium, the velocity (v_m) of EM waves is given by $v_m = \frac{1}{\sqrt{\mu\epsilon}}$ where μ is the permeability and ϵ is the

permittivity of the given medium.

- 8) The EM waves obey the principle of superposition.
- 9) The ratio of the amplitudes of electric and magnetic fields is constant at any point and is equal to the velocity of the EM wave.

$$|\vec{E}_0| = c |\vec{B}_0| \text{ or } \frac{|\vec{E}_0|}{|\vec{B}_0|} = \frac{1}{\sqrt{\mu_0 \epsilon_0}} \text{ --- (13.1)}$$

$|\vec{E}_0|$ and $|\vec{B}_0|$ are the amplitudes of $\vec{E}$ and $\vec{B}$ respectively.

- 10) As the electric field vector ($\vec{E}_0$) is more prominent than the magnetic field vector ($\vec{B}_0$), it is responsible for optical effects due to EM waves. For this reason, electric vector is called light vector.
- 11) The intensity of a wave is proportional to the square of its amplitude and is given by the equations

$$I_E = \frac{1}{2} \epsilon_0 E_0^2, I_B = \frac{1}{2} \frac{B_0^2}{\mu_0} \text{ --- (13.2)}$$

- 12) The energy of EM waves is distributed equally between the electric and magnetic fields. $I_E = I_B$

Example 13.1: Calculate the velocity of EM waves in vacuum.

Solution: The velocity of EM wave in free space is given by

$$c = \frac{1}{\sqrt{\mu_0 \epsilon_0}} = \frac{1}{\sqrt{(8.85 \times 10^{-12} \frac{C^2}{Nm^2})(4\pi \times 10^{-7} \frac{T.m}{A})}}$$

$$c = 3.00 \times 10^8 \text{ m/s}$$

Example 13.2: In free space, an EM wave of frequency 28 MHz travels along the x-direction. The amplitude of the electric field is $E = 9.6$ V/m and its direction is along the y-axis. What is amplitude and direction of magnetic field B ?

Solution : We have,

$$|B| = \frac{|E|}{c} = \frac{9.6 \text{ V/m}}{3 \times 10^8 \text{ m/s}}$$

$$B = 3.2 \times 10^{-8} \text{ T}$$

It is given that E is along y-direction and the wave propagates along x-axis. The magnetic field B should be in a direction perpendicular to



Do you know ?

According to quantum theory, an electron, while orbiting around the nucleus in a stable orbit does not emit EM radiation even though it undergoes acceleration. It will emit an EM radiation only when it falls from an orbit of higher energy to one of lower energy.

EM waves (such as X-rays) are produced when fast moving electrons hit a target of high atomic number (such as molybdenum, copper, etc.).

An electric charge at rest has an electric field in the region around it but has no magnetic field. When the charge moves, it produces both electric and magnetic fields. If the charge moves with a constant velocity, the magnetic field will not change with time and as such it cannot produce an EM wave. But if the charge is accelerated, both the magnetic and electric fields change with space and time and an EM wave is produced. Thus an oscillating charge emits an EM wave which has the same frequency as that of the oscillation of the charge.

both x- and y-axes. As per property (2) of EM waves, $\vec{E} \times \vec{B}$ should be along the direction of propagation which is along the x- axis

Since $(+\hat{j}) \times (+\hat{k}) = \hat{i}$, B is along the $\hat{k}$, i.e., along the z-direction.

Thus, the amplitude of $B = 3.2 \times 10^{-8} \text{ T}$ and its direction is along the z-axis.

Example 13.3: A beam of red light has an amplitude 2.5 times the amplitude of second beam of the same colour. Calculate the ratio of the intensities of the two waves.

Solution: Intensity $\propto$ (Amplitude)²
 $I_2 \propto (a)^2$ and $I_1 \propto (2.5a)^2$

$$\therefore \frac{I_1}{I_2} = \frac{(2.5a)^2}{a^2} = (2.5)^2 = 6.25$$

In an EM wave, the magnetic field and electric field both vary sinusoidally with x . For a wave travelling along x-axis having $\vec{E}$ along y-axis and $\vec{B}$ along the z axis, with reference to Chapter 8, we can write E_y and B_z as

$$E_y = E_0 \sin(kx - \omega t) \quad \text{--- (13.2)}$$

$$\text{and } B_z = B_0 \sin(kx - \omega t), \quad \text{--- (13.3)}$$

where E_0 is the amplitude of the electric field E_y and B_0 is the amplitude of the magnetic field B_z . $k = \frac{2\pi}{\lambda}$ is the propagation constant and λ

is the wavelength of the wave. $\omega = 2\pi\nu$ is the angular frequency of oscillations, ν being the frequency of the wave.

Both the electric and magnetic fields attain their maximum (and minimum) values at the same time and at the same point in space, i.e., $\vec{E}$ and $\vec{B}$ oscillate in phase with the same frequency.

Example 13.4: An EM wave of frequency 50 MHz travels in vacuum along the positive x-axis and $\vec{E}$ at a particular point, x and at a particular instant of time t is $9.6 \hat{j}$ V/m. Find the magnitude and direction of $\vec{B}$ at this point x and at time t.

Solution : $B = \frac{E}{c} = \frac{9.6}{3 \times 10^8} = 3.2 \times 10^{-8} \text{ T}$

As the wave propagates along +x axis and E is along +y axis, direction of B will be along +z-axis i.e. $B = 3.2 \times 10^{-8} \hat{k} \text{ T}$.

Example 13.5: For an EM wave propagating along x direction, the magnetic field oscillates along the z-direction at a frequency of 3×10^{10} Hz and has amplitude of 10^{-9} T.

- What is the wavelength of the wave?
- Write the expression representing the corresponding electric field.

Solution :

$$\text{a) } \lambda = \frac{c}{\nu} = \frac{3 \times 10^8 \text{ m/s}}{3 \times 10^{10} / \text{s}} = 10^{-2} \text{ m}$$

b) $E_0 = cB_0 = (3 \times 10^8 \text{ m/s}) \times (10^{-9} \text{ T}) = 0.3 \text{ V/m}$. Since B acts along z-axis, E acts along y-axis. Expression representing the oscillating electric field is

$$E_y = E_0 \sin(kx - \omega t)$$

$$E_y = E_0 \sin \left[\left(\frac{2\pi}{\lambda} \right) x - (2\pi\nu)t \right]$$

$$E_y = E_0 \sin 2\pi \left[\frac{x}{\lambda} - \nu t \right]$$

$$E_y = E_0 \sin 2\pi \left[\frac{x}{10^{-2}} - 3 \times 10^{10} t \right]$$

$$E_y = E_0 \sin 2\pi [100x - 3 \times 10^{10} t] \text{ V/m}$$

Example 13.6: The magnetic field of an EM wave travelling along x-axis is $\vec{B} = \hat{k} 4 \times 10^{-4} \sin(\omega t - kx)$. Here B is in tesla, t is in second and x is in m. Calculate the peak value of electric force acting on a particle of charge $5 \mu\text{C}$ travelling with a velocity of $5 \times 10^5 \text{ m/s}$ along the y-axis.

Solution :

$$B_0 = 4 \times 10^{-4} \text{ T}, q = 5 \mu\text{C} = 5 \times 10^{-6} \text{ C}$$

$$\nu = 5 \times 10^5 \text{ m/s}$$

$$E_0 = cB_0 = (3 \times 10^8) \times (4 \times 10^{-4}) = 12 \times 10^4 \text{ N/C}$$

$$\begin{aligned} \text{Maximum electric force} &= qE_0 \\ &= (5 \times 10^{-6}) (12 \times 10^4) \\ &= 60 \times 10^{-2} \\ &= 0.6 \text{ N} \end{aligned}$$

13.3 Electromagnetic Spectrum:

The orderly distribution (sequential arrangement) of EM waves according to their wavelengths (or frequencies) in the form of distinct groups having different properties is called the EM spectrum (Fig. 13.2). The properties of different types of EM waves are given in Table 13.1.

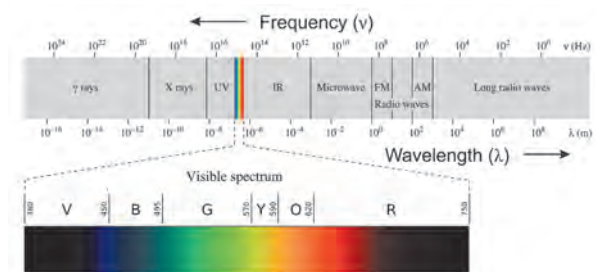


Fig. 13.2: Electromagnetic spectrum.

We briefly describe different types of EM waves in the order of decreasing wavelength (or increasing frequency).

13.3.1 Radio waves :

Radio waves are produced by accelerated motion of charges in a conducting wire. The

frequency of waves produced by the circuit depends upon the magnitudes of the inductance and the capacitance (This will be discussed in XIIth standard). Thus, by choosing suitable values of the inductance and the capacitance, radio waves of desired frequency can be produced.

Properties :

- 1) They have very long wavelengths ranging from a few centimetres to a few hundreds of kilometres.
- 2) The frequency range of AM band is 530 kHz to 1710 kHz. Frequency of the waves used for TV-transmission range from 54 MHz to 890 MHz, while those for FM radio band range from 88 MHz to 108MHz.

Notation used for high frequencies

1 kHz = one kilo Hertz = 1000 Hz = 10^3 Hz

1 MHz = one mega Hertz = 10^6 Hz

1 GHz = one giga Hertz = 10^9 Hz

Notation used for small wavelengths

1 μm = one micrometer = 10^{-6} m

1 $\text{\AA}$ = one angstrom = 10^{-10} m = 10^{-8} cm

1 nm = one nanometer = 10^{-9} m

Uses :

- 1) Radio waves are used for wireless communication purpose.
- 2) They are used for radio broadcasting and transmission of TV signals.
- 3) Cellular phones use radio waves to transmit voice communication in the ultra high frequency (UHF) band.

Table 13.1: Properties of different types of EM waves

Name	Wavelength range in m	Frequency range in Hz	Generated By
Gamma rays	6×10^{-13} to 1×10^{-10}	5×10^{20} to 3×10^{18}	a) Transitions of nuclear energy levels b) Radioactive substances
X-rays	1×10^{-11} to 3×10^{-8}	3×10^{19} to 1×10^{16}	a) Bombardment of high energy electrons (keV) on a high atomic number target (Cu, Mg, Co) b) Energy level transitions of innermost orbital electrons
Ultraviolet (UV waves)	3×10^{-8} to 4×10^{-7}	1×10^{16} to 8×10^{14}	Rearrangement of orbital electrons of atom between energy levels. As in high voltage gas discharge tube, the Sun and mercury vapour lamp, etc.
Visible light	4×10^{-7} to 8×10^{-7}	8×10^{14} to 4×10^{14}	Rearrangement of outer orbital electrons in atoms and molecules e.g., gas discharge tube
Infrared (IR) radiations	8×10^{-7} to 3×10^{-4}	4×10^{14} to 1×10^{12}	Hot objects
Microwaves	3×10^{-4} to 6×10^{-2}	1×10^{12} to 5×10^9	Special electronic devices such as klystron tube
Radio waves	6×10^{-4} to 1×10^5	5×10^{11} to 8×10^{10}	Acceleration of electrons in circuits

13.3.2 Microwaves :

These waves were discovered of by H. Hertz in 1888. Microwaves are produced by oscillator electric circuits containing a capacitor and an inductor. They can be produced by special vacuum tubes.

Properties

- 1) They heat certain substances on which

they are incident.

- 2) They can be detected by crystal detectors.

Uses

- 1) Used for the transmission of TV signals.
- 2) Used for long distance telephone communication.
- 3) Microwave ovens are used for cooking.
- 4) Used in radar systems for the location of

distant objects like ships, aeroplanes etc,

- 5) They are used in the study of atomic and molecular structure.

13.3.3 Infrared waves

These waves were discovered by William Herschel (1737-1822) in 1800. All hot bodies are sources of infrared rays. About 60% of the solar radiations are infrared in nature. Thermocouples, thermopile and bolometers are used to detect infrared rays.

Properties

- 1) When infrared rays are incident on any object, the object gets heated.
- 2) These rays are strongly absorbed by glass.
- 3) They can penetrate through thick columns of fog, mist and cloud cover.

Uses

- 1) Used in remote sensing.
- 2) Used in diagnosis of superficial tumours and varicose veins.
- 3) Used to cure infantile paralysis and to treat sprains, dislocations and fractures.
- 4) They are used in Solar water heaters and cookers.
- 5) Special infrared photographs of the body called thermograms, can reveal diseased organs because these parts radiate less heat than the healthy organs.
- 6) Infrared binoculars and thermal imaging cameras are used in military applications for night vision.
- 7) Used to keep green house warm.
- 8) Used in remote controls of TV, VCR, etc

13.3.4 Visible light :

It is the most familiar form of EM waves. These waves are detected by human eye. Therefore this wavelength range is called the visible light. The visible light is emitted due to atomic excitations.

Properties :

- 1) Different wavelengths give rise to different colours. These are given in Table 13.2.
- 2) Visible light emitted or reflected from objects around us provides us information about those objects and hence about the surroundings.



Do you know ?

Stars and galaxies emit different types of waves. Radio waves and visible light can pass through the Earth's atmosphere and reach the ground without getting absorbed significantly. Thus the radio telescopes and optical telescopes can be placed on the ground. All other type of waves get absorbed by the atmospheric gases and dust particles. Hence, the γ -ray, X-ray, ultraviolet, infrared, and microwave telescopes are kept aboard artificial satellites and are operated remotely from the Earth. Even though the visible radiation reaches the surface of the Earth, its intensity decreases to some extent due to absorption and scattering by atmospheric gases and dust particles. Optical telescopes are therefore located at higher altitudes.

The Indian Giant Metrewave Radio Telescope (GMRT) near Pune is an important milestone in the field of Radio-astronomy. Also, Indian Astronomical Observatory houses the Himalayan Chandra Telescope (HCT), the 2 m optical-IR Telescope, which is situated at Hanle, Ladakh, at an altitude of 4500 m.

Table 13.2: Wavelengths of colours in visible light

Colour	Wavelength
violet	380-450 nm
blue	450-495 nm
green	495-570 nm
yellow	570-590 nm
orange	590-620 nm
red	620-750 nm

13.3.5 Ultraviolet rays :

Ultraviolet rays were discovered by J. Ritter (1776-1810) in 1801. They can be produced by the mercury vapour lamp, electric spark and carbon arc lamp. They can also be obtained by striking electrical discharge in hydrogen and xenon gas tubes. The Sun is the most important natural source of ultraviolet rays, most of which are absorbed by the ozone layer in the Earth's atmosphere.

Properties :

- 1) They produce fluorescence in certain materials, such as 'phosphors'.
- 2) They cause photoelectric effect.
- 3) They cannot pass through glass but pass through quartz, fluorite, rock salt etc.
- 4) They possess the property of synthesizing vitamin D, when skin is exposed to them.

Uses :

- 1) Ultraviolet rays destroy germs and bacteria and hence they are used for sterilizing surgical instruments and for purification of water.
- 2) Used in burglar alarms and security systems.
- 3) Used to distinguish real and fake gems.
- 4) Used in analysis of chemical compounds.
- 5) Used to detect forgery.



Do you know ?

1. A fluorescent light bulb is coated from within with a powder inside and contains a gas; electricity causes the gas to emit ultraviolet radiation, which then stimulates the tube coating to emit light.

2. The pixels of a television or computer screen fluoresce when electrons from an electron gun strike them.

3. What we call 'visible light' is just the part of the EM spectrum that human eyes see. Many other animals would define 'visible' somewhat differently. For instance, many animals including insects and birds, see in the UV region. Natural world is full of signals that animals see and humans cannot. Many birds including bluebirds, budgies, parrots and even peacocks have ultraviolet patterns that make them even more vivid to each other than they are to us.

13.3.6 X-rays:

German physicist W. C. Rontgen (1845-1923) discovered X-rays in 1895 while studying cathode rays (which is a stream of electrons emitted by the cathode in a vacuum tube). X-rays are also called Rontgen rays. X-rays are produced when cathode rays are suddenly stopped by an obstacle.

Properties

- 1) They are high energy EM waves.
- 2) They are not deflected by electric and magnetic fields.
- 3) X-rays ionize the gases through which they pass.
- 4) They have high penetrating power.
- 5) Their over dose can kill living plant and animal overdose tissues and hence are harmful.

Uses

- 1) Useful in the study of the structure of crystals.
- 2) X-ray photographs are useful to detect bone fracture. X-rays have many other medical uses such as CT scan.
- 3) X-rays are used to detect flaws or cracks in metals.
- 4) These are used for detection of explosives, opium etc.

13.3.7 Gamma Rays (γ -rays)

Discovered by P. Villard (1860-1934) in 1900. Gamma rays are emitted from the nuclei of some radioactive elements such as uranium, radium etc.

Properties

- 1) They are highest energy EM waves. (energy range keV - GeV)
- 2) They are highly penetrating.
- 3) They have a small ionising power.
- 4) They kill living cells.

Uses

- 1) Used as insecticide disinfection for wheat and flour.
- 2) Used for food preservation.
- 3) Used in radiotherapy for the treatment of cancer and tumour.
- 4) They are used to produce nuclear reactions.

13.4 Propagation of EM Waves:

You must have seen a TV antenna used to receive the TV signals from the transmitting tower or from a satellite. In communication using radio waves, an antenna in the transmitter radiates the EM waves, which travel through space and reach the receiving antenna at the other end. As the EM wave travels away from the transmitter; the strength of the wave keeps

on decreasing. Several factors influence the propagation of EM waves and the path they follow. It is also important to understand the composition of the Earth's atmosphere as it plays a vital role in the propagation of EM waves. Different layers of Earth's atmosphere are shown in Fig. 13.3.



Do you know ?

Ionizing radiations :

Ultraviolet, X-ray and gamma rays have sufficient energy to cause ionization i.e. they strip electrons from atoms and molecules lying along their path. The atoms lose their electrons and are then known as ions. Ionization is harmful to human beings because it can kill or damage living cells, or make them grow abnormally as cancers. Fluorescent lamps are based on ionization of gas. Ionizing radiation is also used in various equipments in laboratory and industry.

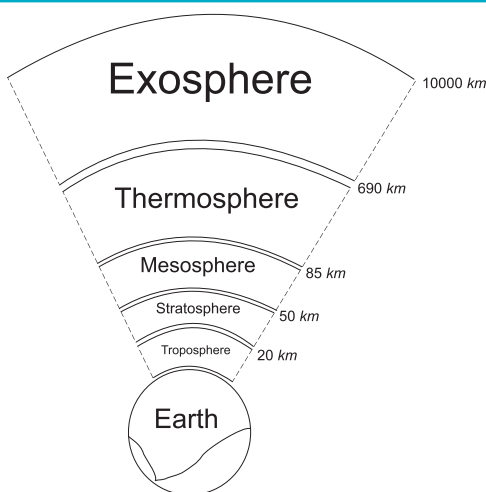


Fig 13.3: Earth and atmospheric layers.

Different modes of propagation of EM waves are described below and are shown in Fig. 13.4.



Do you know ?

X-rays have many practical applications in medicine and industry. Because X-ray photons are of such high energy, they can penetrate several centimetres of solid matter and can be used to visualize the interiors of materials that are opaque to ordinary light.

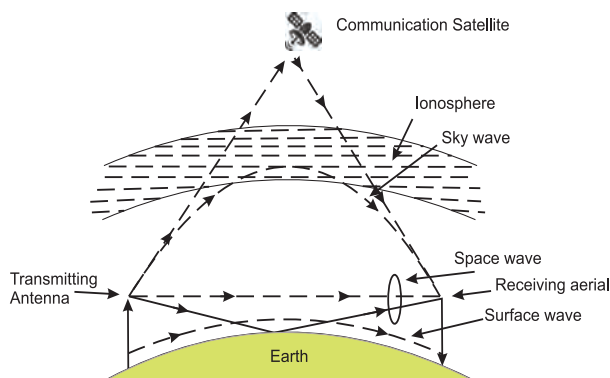


Fig. 13.4: Propagation of EM waves.

13.4.1 Ground (surface) wave:

When a radio wave from a transmitting antenna propagates near surface of the Earth so as to reach the receiving antenna, the wave propagation is called ground wave or surface wave propagation.

In this mode, radio waves travel close to the surface of the Earth and move along its curved surface from transmitter to receiver.

The radio waves induce currents in the ground and lose their energy by absorption. Therefore, the signal cannot be transmitted over large distances. Radio waves having frequency less than 2 MHz (in the medium frequency band) are transmitted by ground wave propagation. This is suitable for local broadcasting only. For TV or FM signals (very high frequency), ground wave propagation cannot be used.

13.4.2 Space wave:

When the radio waves from the transmitting antenna reach the receiving antenna either directly along a straight line (line of sight) or after reflection from the ground or satellite or after reflection from troposphere, the wave propagation is called space wave propagation. The radio waves reflected from troposphere are called tropospheric waves. Radio waves with frequency greater than 30 MHz can pass through the ionosphere (60 km - 1000 km) after suffering a small deviation. Hence, these waves cannot be transmitted by space wave propagation except by using a satellite. Also, for TV signals which have high frequency, transmission over long distance is not possible by means of space wave propagation.

The maximum distance over which a signal can reach is called its range. For larger TV

coverage, the height of the transmitting antenna should be as large as possible. This is the reason why the transmitting and receiving antennas are mounted on top of high rise buildings.

Range is the straight line distance from the point of transmission (the top of the antenna) to the point on Earth where the wave will hit while travelling along a straight line. Range is shown by d in Fig. 13.5. Let the height of the transmitting antenna (AA') situated at A be h . B represents the point on the surface of the Earth at which the space wave hits the Earth. The triangle OA'B is a right angled triangle. From $\Delta OA'B$ we can write

$$OA'^2 = A'B^2 + OB^2$$

$$(R+h)^2 = d^2 + R^2$$

$$\text{or } R^2 + h^2 + 2Rh = d^2 + R^2$$

As $h \ll R$, we can ignore h^2 and write

$$d \cong \sqrt{2Rh}$$

The range can be increased by mounting the receiver at a height h' say at a point C on the surface of the Earth. The range increases to $d + d'$ where d' is $\sqrt{2Rh'}$. Thus

$$\text{Total range} = d + d' = \sqrt{2Rh} + \sqrt{2Rh'}$$

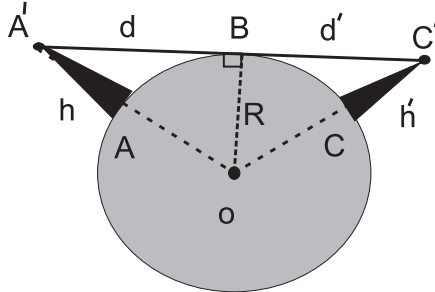


Fig. 13.5: Range of the signal (not to scale).

Example 13.7: A radar has a power of 10 kW and is operating at a frequency of 20 GHz. It is located on the top of a hill of height 500 m. Calculate the maximum distance upto which it can detect object located on the surface of the Earth. (Radius of Earth = 6.4×10^6 m)

Solution:

Maximum distance (range) =

$$d = \sqrt{2Rh}$$

$$= \sqrt{2 \times (6.4 \times 10^6) \times 500} \text{ m}$$

$$= 8 \times 10^4 = 80 \text{ km,}$$

where R is radius of the Earth and h is the height of the radar above Earth's surface.

Example 13.8: If the height of a TV transmitting antenna is 128 m, how much square area can be covered by the transmitted signal if the receiving antenna is at the ground level? (Radius of the Earth = 6400 km)

Solution:

$$\begin{aligned} \text{Range} = d &= \sqrt{2Rh} \\ &= \sqrt{2(6400 \times 10^3)(128)} \\ &= \sqrt{16.384 \times 10^5 \times 10^3} \\ &= \sqrt{16.384 \times 10^8} \\ &= 4.047 \times 10^4 \\ &= 40.470 \text{ km} \end{aligned}$$

$$\begin{aligned} \text{Area covered} &= \pi d^2 = 3.14 \times (40)^2 \\ &= 5144.58 \text{ km}^2 \end{aligned}$$

Example 13.9: The height of a transmitting antenna is 68 m and the receiving antenna is at the top of a tower of height 34 m. Calculate the maximum distance between them for satisfactory transmission in line of sight mode. (radius of Earth = 6400 km)

$$h_t = 68 \text{ m, } h_r = 34 \text{ m, } R = 6400 \text{ km} = 6.4 \times 10^6 \text{ m}$$

Solution:

$$\begin{aligned} d_{\max} &= \sqrt{2Rh_t} + \sqrt{2Rh_r} \\ &= \sqrt{2 \times 6.4 \times 10^6 \times 68} + \sqrt{2 \times 6.4 \times 10^6 \times 34} \\ &= \sqrt{870 \times 10^3} + \sqrt{435 \times 10^3} \\ &= 29.5 \times 10^3 + 20.9 \times 10^3 \\ &= 50.4 \times 10^3 \text{ m} \\ &= 50.4 \text{ km} \end{aligned}$$

13.4.3 Sky wave propagation:

When radio waves from a transmitting antenna reach the receiving antenna after reflection in the ionosphere, the wave propagation is called sky wave propagation.

The sky waves include waves of frequency between 3 MHz and 30 MHz. These waves can suffer multiple reflections between the ionosphere and the Earth. Therefore, they can be transmitted over large distances.

Critical frequency : It is the maximum value of the frequency of radio wave which can be reflected back to the Earth from the ionosphere when the waves are directed normally to ionosphere.

Skip distance (zone) : It is the shortest distance from a transmitter measured along the surface of the Earth at which a sky wave of fixed frequency (if greater than critical frequency) will be returned to the Earth so that no sky waves can be received within the skip distance.

13.5 Introduction to Communication System:

Communication is exchange of information. Since ancient times it is practiced in various ways e.g., through speaking, writing, singing, using body language etc. After the discovery of electricity in the late 19th century, human communication systems changed dramatically. Modern communication is based upon the discoveries and inventions by a number of scientists like J. C. Bose (1858-1937), S. F. B. Morse (1791-1872), G. Marconi (1874-1937) and Alexander Graham Bell (1847-1922) in the 19th and 20th centuries.

In the 20th century we could send messages over large distances using analogue signals, cables and radio waves. With the advancements of digitization technologies, we can now communicate with the entire world almost in real time.

The ability to communicate is an important feature of modern life. We can speak directly to others all around the world and generate vast amount of information every day.

Here we will briefly discuss how communication systems work. A communication system is a device or set up used in transmission and reception of information from one place to another.

13.5.1 Elements of a communication system:

There are three basic (essential) elements of every communication system: a) Transmitter, b) Communication channel and c) Receiver.

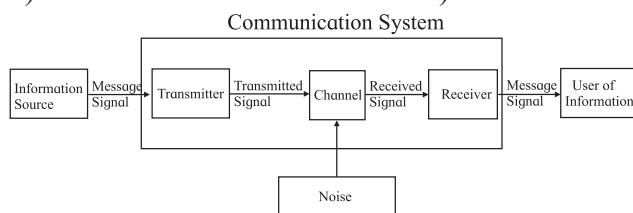


Fig. 13.6: Block diagram of the basic elements of a communication system.

In a communication system, as shown in Fig. 13.6, the transmitter is located at one

place and the receiver at another place. The communication channel is a passage through which signals transfer in between a transmitter and a receiver. This channel may be in the form of wires or cables, or may also be wireless, depending on the types of communication system.

There are two basic modes of communication: (i) point to point communication and (ii) broadcast.

In point to point communication mode, communication takes place over a link between a single transmitter and a receiver e.g. Telephony. In the broadcast mode there are large number of receivers corresponding to the single transmitter e.g., Radio and Television transmission.

13.5.2 Commonly used terms in electronic communication system:

Following terms are useful to understand any communication system:

1) Signal :- The information converted into electrical form that is suitable for transmission is called a signal. In a radio station, music and speech are converted into electrical form by a microphone for transmission into space. This electrical form of sound is the signal. A signal can be analog or digital as shown in Fig. 13.7.

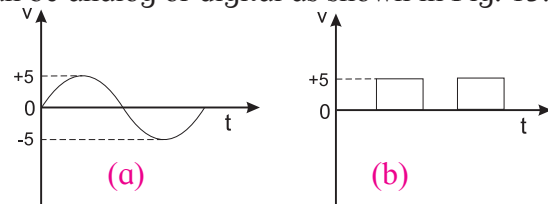


Fig 13.7: (a) Analog signal. (b) Digital signal.

(i) Analog signal: A continuously varying signal (voltage or current) is called an analog signal. Since a wave is a fundamental analog signal, sound and picture signals in TV are analog in nature (Fig 13.7 a)

(ii) Digital signal: A signal (voltage or current) that can have only two discrete values is called a digital signal. For example, a square wave is a digital signal. It has two values viz, +5 V and 0 V. (Fig- 13.7 b)

2) Transmitter :- A transmitter converts the signal produced by a source of information into a form suitable for transmission through a channel and subsequent reception.

3) Transducer :- A device that converts one form of energy into another form of energy is called a transducer. For example, a microphone converts sound energy into electrical energy. Therefore, a microphone is a transducer. Similarly, a loudspeaker is a transducer which converts electrical energy into sound energy.

4) Receiver :- The receiver receives the message signal at the channel output, reconstructs it in recognizable form of the original message for delivering it to the user of information.

5) Noise :- A random unwanted signal is called noise. The source generating the noise may be located inside or outside the system. Efforts should be made to minimise the noise level in a communication system.

6) Attenuation :- The loss of strength of the signal while propagating through the channel is known as attenuation. It occurs because the channel distorts, reflects and refracts the signals as it passes through it.

7) Amplification :- Amplification is the process of raising the strength of a signal, using an electronic circuit called amplifier.

8) Range :- The maximum (largest) distance between a source and a destination up to which the signal can be received with sufficient strength is termed as range.

9) Bandwidth :- The bandwidth of an electronic circuit is the range of frequencies over which it operates efficiently.

10) Modulation :- The signals in communication system (e.g. music, speech etc.) are low frequency signals and cannot be transmitted over large distances. In order to transmit the signal to large distances, it is superimposed on a high frequency wave (called carrier wave). This process is called modulation. Modulation is done at the transmitter and is an important part of a communication system.

11) Demodulation :- The process of regaining signal from a modulated wave is called demodulation. This is the reverse process of modulation.

12) Repeater :- It is a combination of a transmitter and a receiver. The receiver receives the signal from the transmitter, amplifies it and transmits it to the next repeater. Repeaters are

used to increase the range of a communication system. These are shown in Fig. 13.8.

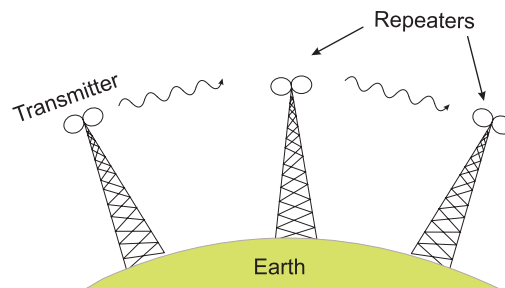


Fig.13.8: Use of repeater station to increase the range of communication.



Do you know ?

To transmit a signal we need an antenna or an aerial. For efficient transmission and reception, the transmitting and receiving antennas must have a length at least $\lambda/4$ where λ is the wavelength of the signal.

For an audio signal of 15kHz, the required length of the antenna is $\lambda/4$ which can be seen to be equal to 5 km.

The highest TV tower in Rameshwaram, Tamilnadu, is the tallest tower in India and is ranked 32nd in the world with pinnacle height of 323 metre. It is used for television broadcast by the Doordarshan.

13.6 Modulation:

As mentioned earlier, an audio signal has low frequency (< 20 KHz). Low frequency signals can not be transmitted over large distances. Because of this, a high frequency wave, called a carrier wave, is used. Some characteristic (e.g. amplitude, frequency or phase) of this wave is changed in accordance with the amplitude of the signal. This process is known as modulation. Modulation also helps avoid mixing up of signals from different transmitters as different carrier wave frequencies can be allotted to different transmitters. Without the use of these waves, the audio signals, if transmitted directly by different transmitters, would have got mixed up.

Modulation can be done by modifying the (i) amplitude (amplitude modulation) (ii) frequency (frequency modulation), and (iii) phase (phase modulation) of the carrier wave in proportion to the amplitude or intensity of the

signal wave keeping the other two properties same. Figure 13.9 (a) shows a carrier wave and (b) shows the signal. The carrier wave is a high frequency wave while the signal is a low frequency wave. Amplitude modulation, frequency modulation and phase modulation of carrier waves are shown in Fig. 13.9 (c), (d) and (e) respectively.

Amplitude modulation (AM) is simple to implement and has large range. It is also cheaper. Its disadvantages are that (i) it is not very efficient as far as power usage is concerned (ii) it is prone to noise and (iii) the reproduced signal may not exactly match the original signal. In spite of this, these are used for commercial broadcasting in the long, medium and short wave bands.

Frequency modulation (FM) is more complex as compared to amplitude modulation and, therefore is more difficult to implement. However, its main advantage is that it reproduces the original signal closely and is less

susceptible to noise. This modulation is used for high quality broadcast transmission.

Phase modulation (PM) is easier than frequency modulation. It is used in determining the velocity of a moving target which cannot be done using frequency modulation.

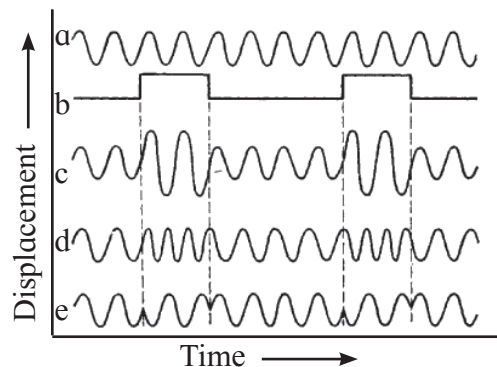


Fig. 13.9: (a) Carrier wave, (b) signal (c) AM (d) FM and (e) PM.



Internet my friend

<https://www.iiap.res.in/centers/iao>



Exercises

1. Choose the correct option.

- The EM wave emitted by the Sun and responsible for heating the Earth's atmosphere due to green house effect is
(A) Infra-red radiation (B) X ray
(C) Microwave (D) Visible light
- Earth's atmosphere is richest in
(A) UV (B) IR
(C) X-ray (D) Microwaves
- How does the frequency of a beam of ultraviolet light change when it travels from air into glass?
(A) No change (B) increases
(C) decreases (D) remains same
- The direction of EM wave is given by
(A) $\vec{E} \times \vec{B}$ (B) $\vec{E} \cdot \vec{B}$
(C) along $\vec{E}$ (D) along $\vec{B}$
- The maximum distance upto which TV transmission from a TV tower of height h can be received is proportional to
(A) $h^{1/2}$ (B) h
(C) $h^{3/2}$ (D) h^2

- The waves used by artificial satellites for communication purposes are
(A) Microwave
(B) AM radio waves
(C) FM radio waves
(D) X-rays
- If a TV telecast is to cover a radius of 640 km, what should be the height of transmitting antenna?
(A) 32000 m (B) 53000 m
(C) 42000 m (D) 55000 m

2. Answer briefly.

- State two characteristics of an EM wave.
- Why are microwaves used in radar?
- What are EM waves?
- How are EM waves produced?
- Can we produce a pure electric or magnetic wave in space? Why?
- Does an ordinary electric lamp emit EM waves?
- Why do light waves travel in vacuum whereas sound wave cannot?

- viii) What are ultraviolet rays? Give two uses.
- ix) What are radio waves? Give its two uses.
- x) Name the most harmful radiation entering the Earth's atmosphere from the outer space.
- xi) Give reasons for the following:
 (i) Long distance radio broadcast uses short wave bands.
 (ii) Satellites are used for long distance TV transmission.
- xii) Name the three basic units of any communication system.
- xiii) What is a carrier wave?
- xiv) Why high frequency carrier waves are used for transmission of audio signals?
- xv) What is modulation?
- xvi) What is meant by amplitude modulation?
- xvii) What is meant by noise?
- xviii) What is meant by bandwidth?
- xix) What is demodulation?
- xx) What type of modulation is required for television broadcast?
- xxi) How does the effective power radiated by an antenna vary with wavelength?
- xxii) Why should broadcasting programs use different frequencies?
- xxiii) Explain the necessity of a carrier wave in communication.
- xxiv) Why does amplitude modulation give noisy reception?
- xxv) Explain why is modulation needed.
- 2. Solve the numerical problem.**
- i) Calculate the frequency in MHz of a radio wave of wavelength 250 m. Remember that the speed of all EM waves in vacuum is 3.0×10^8 m/s.
 [Ans: 1.2 MHz]
- ii) Calculate the wavelength in nm of an X-ray wave of frequency 2.0×10^{18} Hz.
 [Ans: 0.15 nm]
- iii) The speed of light is 3×10^8 m/s. Calculate the frequency of red light of wavelength of 6.5×10^{-7} m.
 [Ans: $\nu = 4.6 \times 10^{14}$ Hz]
- iv) Calculate the wavelength of a microwave of frequency 8.0 GHz.
 [Ans: 3.75 cm]
- v) In a EM wave the electric field oscillates sinusoidally at a frequency of 2×10^{10} Hz. What is the wavelength of the wave?
 [Ans: 1.5×10^{-2} m]
- vi) The amplitude of the magnetic field part of a harmonic EM wave in vacuum is $B_0 = 5 \times 10^{-7}$ T. What is the amplitude of the electric field part of the wave?
 [Ans: 150V/m]
- vii) A TV tower has a height of 200 m. How much population is covered by TV transmission if the average population density around the tower is $1000/\text{km}^2$? (Radius of the Earth = 6.4×10^6 m)
 [Ans: 8×10^6]
- viii) Height of a TV tower is 600 m at a given place. Calculate its coverage range if the radius of the Earth is 6400 km. What should be the height to get the double coverage area?
 [Ans: 87.6 km, 1200 m]
- ix) A transmitting antenna at the top of a tower has a height 32 m and that of the receiving antenna is 50 m. What is the maximum distance between them for satisfactory communication in line of sight mode? Given radius of Earth is 6.4×10^6 m.
 [Ans: 45.537 km]



Can you recall?

1. Your mobile handset is very efficient gadget.
2. International Space Station works using solar energy.
3. A LED TV screen produces brighter and vivid colours.
4. Good and bad conductor of electricity.

14.1 Introduction:

Modern life is heavily dependent on many electronic gadgets. It could be a cell phone, a smart watch, a computer or even an LED lamp, they all have one common factor, semiconductor devices that make them work. Semiconductors have made our life very comfortable and easy.

Semiconductors are materials whose electrical properties can be tailored to suit our requirements. Before the discovery of semiconductors, electrical properties of materials could be of two types, conductors or insulators. Conductors such as metals have a very high electrical conductivity, for example, conductivity of silver is $6.25 \times 10^7 \text{ Sm}^{-1}$ whereas an insulator or a bad conductor like glass has a very low electrical conductivity of the order of 10^{-10} Sm^{-1} . Electrical conductivity of silicon, a semiconductor, for example is $1.56 \times 10^{-3} \text{ Sm}^{-1}$. It lies between that of a good conductor and a bad conductor. A semiconductor can be customised to have its electrical conductivity as per our requirement. Temperature dependence of electrical conductivity of a semiconductor can also be controlled. Table 14.1 gives electrical conductivity of some materials which are commonly used.

14.2 Electrical conduction in solids:

Electrical conduction in a solid takes place by transport of charge carriers. It depends on its temperature, the number of charge carriers, how easily these carries can move inside a solid (mobility), its crystal structure, types and the nature of defects present in a solid etc. There can be three types of electrical conductors. It could be a good conductor, a semiconductor or a bad conductor.

1. Conductors (Metals): The best example of a conductor is any metal. They have a large

number of free electrons available for electrical conduction. (A typical metal will have 10^{28} electrons per m^3). Metals are good conductors of electricity due to the large number of free electrons present in them.

2. Insulators: Glass, wood or rubber are some common examples of insulators. Insulators have very small number (10^{23} per m^3) of free electrons.

3. Semiconductors: Silicon, germanium, gallium arsenide, gallium nitride, cadmium sulphide are some of the commonly used semiconductors. The electrical conductivity of a semiconductor is between the conductivity of a metal and that of an insulator. The number of charge carriers in a semiconductor can be controlled as per our requirement. Their structure can also be designed to suit our requirement. Such materials are very useful in electronic industry and find applications in almost every gadget of daily use such as a cell phone, a solar cell or a complex system such as a satellite or the International Space Station.

Table 14.1: Electrical conductivities of some commonly used materials

Material	Electrical conductivity (S m^{-1})
Silver	6.30×10^7
Copper	5.96×10^7
Aluminium	3.5×10^7
Gold	4.10×10^7
Nichrome	9.09×10^5
Platinum	9.43×10^6
Germanium	2.17
Silicon	1.56×10^{-3}
Air	3×10^{-15} to 8×10^{-15}
Glass	10^{-11} to 10^{-15}
Teflon	10^{-25} to 10^{-23}
Wood	10^{-16} to 10^{-24}



Do you know ?

Electrical conductivity σ of a solid is given by $\sigma = nqu$, where,

n = charge carrier density
(number of carriers per unit volume)

q = charge on the carriers

μ = mobility of carriers

Mobility of a charge carrier is the measure of the ease with which a carrier can move in a material under the action of an external electric field. It depends upon many factors such as mass of the carrier, whether the material is crystalline or amorphous, the presence of structural defects in a material, the nature of impurities in a material and so on.

Figure 14.1 shows the temperature dependence of the electrical conductivity of a typical metal and a semiconductor. When the temperature of a semiconductor is increased, its electrical conductivity also increases. The electrical conductivity of a metal decreases with increase in its temperature.

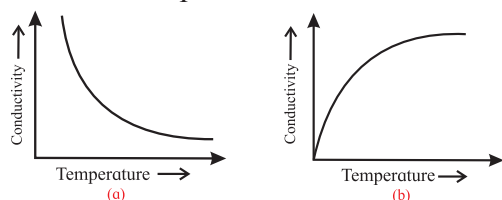


Fig. 14.1: Temperature dependence of electrical conductivity of (a) metals and (b) semiconductors.

Variation of electrical conductivity of semiconductors with change in its temperature is a very useful property and finds applications in a large number of electronic devices. A broad classification of semiconductors can be:

a. Elemental semiconductors: Silicon, germanium

b. Compound Semiconductors: Cadmium sulphide, zinc sulphide, etc.

c. Organic Semiconductors: Anthracene, doped phthalocyanines, polyaniline etc.

Elemental semiconductors and compound semiconductors are widely used in electronic industry. Discovery of organic semiconductors is relatively new and they find lesser applications.

Electrical properties of semiconductors are different from metals and insulators due to their unique conduction mechanism. The electronic configuration of the elemental semiconductors silicon and germanium plays a very important role in their electrical properties. They are from the fourth group of elements in the periodic table. They have a valence of four. Their atoms are bonded by covalent bonds. At absolute zero temperature, all the covalent bonds are completely satisfied in a single crystal of pure silicon or germanium.

The conduction mechanism in a semiconductor can be better understood with the help of the band theory of solids.

14.3 Band theory of solids, a brief introduction:

We begin with the way electron energies in an isolated atom are distributed. An isolated atom has its nucleus at the center which is surrounded by a number of revolving electrons. These electrons are arranged in different and discrete energy levels.

When a solid is formed, a large number of atoms are packed in it. The outermost electronic energy levels in a solid are occupied by electrons from all atoms in a solid. Sharing of the outermost energy levels and resulting formation of energy bands can be easily understood by considering formation of solid sodium.

The electronic configuration of sodium (atomic number 11) is $1s^2, 2s^2, 2p^6, 3s^1$. The outermost level $3s$ can take one more electron but it is half filled in sodium.

When solid sodium is formed, atoms interact with each other through the electrons in each atom. The energy levels are filled according to the Pauli's exclusion principle. According to this principle, no two electrons can have the same set of quantum numbers, or in simple words, no two electrons with similar spin can occupy the same energy level.

Any energy level can accommodate only two electrons (one with spin up state and the other with spin down state). According to this principle, there can be two states per energy level. Figure 14.2 (a) shows the allowed energy

levels of an isolated sodium atom by horizontal lines. The curved lines represent the potential energy of an electron near the nucleus due to Coulomb interaction.

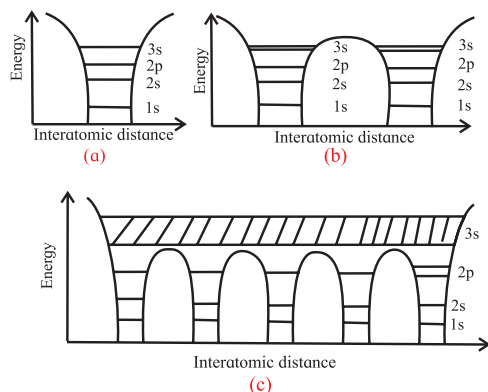


Fig. 14.2: Potential energy diagram, energy levels and bands (a) isolated atom, (b) two atoms, (c) sodium metal.

Consider two sodium atoms close enough so that outer 3s electrons are equally likely to be on any atom. The 3s electrons from both the sodium atoms need to be accommodated in the same level. This is made possible by splitting the 3s level into two sub-levels so that the Pauli's exclusion principle is not violated. Figure 14.2 (b) shows the splitting of the 3s level into two sub levels. When solid sodium is formed, the atoms come close to each other (distance between them $\sim 2 - 3\text{\AA}$). Therefore, the electrons from different atoms interact with each other and also with the neighbouring atomic cores. The interaction between the outer most electrons is more due to overlap while the inner core electrons remain mostly unaffected. Each of these energy levels is split into a large number of sub levels, of the order of Avogadro's number. This is because the number of atoms in solid sodium is of the order of this number. The separation between the sublevels is so small that the energy levels appear almost continuous. This continuum of energy levels is called an energy band. The bands are called 1s band, 2s band, 2p band and so on. Figure 14.2 c shows these bands in sodium metal. Broadening of valence and higher bands is more because of stronger interaction of these electrons.

For sodium atom, the topmost occupied energy level is the 3s level. This level is called the valence level. Corresponding energy band is called the **valence band**. Thus, the valence

band in solid sodium is the topmost occupied energy band. The valence band is half filled in sodium. Figure 14.3 shows the energy bands in sodium.

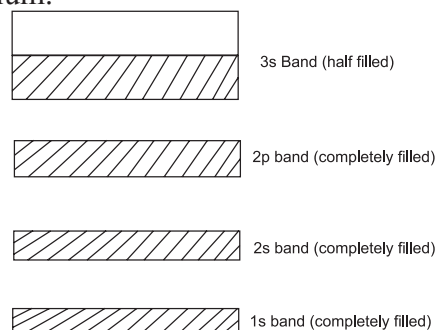


Fig. 14.3: Energy bands in sodium.

When sufficient energy is provided to electrons from the valence band they are raised to higher levels. The immediately next energy level that electrons from valence band can occupy is called conduction level. The band formed by conduction levels is called **conduction band**. In sodium valence and conduction bands overlap.

In a semiconductor or an insulator, there is a gap between the bottom of the conduction band and the top of the valence band. This is called the energy gap or the band gap.

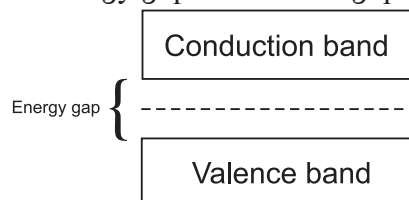


Fig. 14.4: Energy bands for a typical solid.

Figure 14.4 shows the conduction band, the energy gap and the valence band for a typical solid which is not a good conductor. **It is important to remember that this structure is related to the energy of electrons in a solid and it does not represent the physical structure of a solid in any way.**

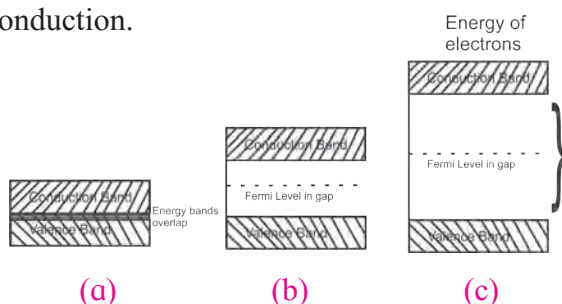
All the energy levels in a band, including the topmost band, in a semiconductor are completely occupied at absolute zero. At some finite temperature T , few electrons gain thermal energy of the order of kT , where k is the Boltzmann constant.

Electrons in the bands below the valence band cannot move to higher band since these are already occupied. Only electrons from the valence band can be excited to the empty

Formation of energy bands in a solid is a result of the small distances between atoms, the resulting interaction amongst electrons and the Pauli's exclusion principle.

conduction band, if the thermal energy gained by these electrons is greater than the band gap. In case of sodium, electrons from the 3s band can gain thermal energy and occupy a slightly higher energy level because the 3s band is only half filled.

Electrons can also gain energy when an external electric field is applied to a solid. Energy gained due to electric field is smaller, hence only electrons at the topmost energy level gain such energy and participate in electrical conduction.



(a) (b) (c)
Fig. 14.5: Band structure of a (a) metal, (b) semiconductor, and an insulator (c).

The difference in electrical conductivities of various solids can be explained on the basis of the band structure of solids. Band structure in a metal, semiconductor and an insulator is different. Figure 14.5 shows a schematic representation of band structure of a metal, a semiconductor and an insulator.

For metals, the valence band and the conduction band overlap and there is no band gap as shown in Fig.14.5 (a). Electrons, therefore, find it easy to gain electrical energy when some external electric field is applied. They are, therefore, easily available for conduction.

In case of semiconductors, the band gap is fairly small, of the order of one electron volt or less as shown in Fig.14.5 (b). When excited, electrons gain energy and occupy energy levels in conduction band easily and can take part in electric conduction.

Insulators, on the contrary, have a wide gap between valence band and conduction band as shown in Fig.14.5 (c). Diamond, for example, has a band gap of about 5.0 eV. In an insulator, therefore, electrons find it very difficult to gain

sufficient energy and occupy energy levels in the conduction band.

The magnitude of the band gap plays a very important role in electronic properties of a solid.

Table 14.2: Magnitude of energy gap in silicon, germanium and diamond.

Material	Energy gap (eV) At 300 K
Silicon	1.12
Germanium	0.66
Diamond	5.47

1 eV is the energy gained by an electron while it overcomes a potential difference of one volt. $1 \text{ eV} = 1.6 \times 10^{-19} \text{ J}$.

14.4 Intrinsic Semiconductor:

A pure semiconductor such as pure silicon or pure germanium is called an intrinsic semiconductor. Silicon (Si) has atomic number 14 and its electronic configuration is $1s^2 2s^2 2p^6 3s^2 3p^2$. Its valence is 4. Each atom of Si forms four covalent bonds with its neighbouring atoms. One Si atom is surrounded by four Si atoms at the corners of a regular tetrahedron Fig. 14.6.

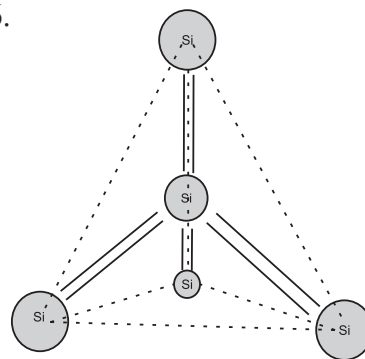


Fig. 14.6: Structure of silicon.

At absolute zero temperature, all valence electrons are tightly bound to respective atoms and the covalent bonds are complete. Electrons are not available to conduct electricity through the crystal because they cannot gain enough energy to get into higher energy levels. At room temperature, however, a few covalent bonds are broken due to thermal agitation and some valence electrons can gain energy. Thus we can say that a valence electron is moved to the conduction band. It creates a vacancy in the valence band as shown in Fig. 14.7.

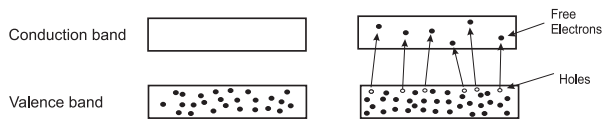


Fig. 14.7: Creation of vacancy in the valence band.

These vacancies of electrons in the valence band are called holes. The holes are thus absence of electrons in the valence band and they carry an *effective positive charge*.

For an intrinsic semiconductor, the number of holes per unit volume, (the number density, n_h) and the number of free electrons per unit volume, (the number density, n_e) is the same.

$$n_h = n_e$$

Electric conduction through an intrinsic semiconductor is quite interesting. There are two different types of charge carriers in a semiconductor. One is the electron and the other is the hole or absence of electron. Electrical conduction takes place by transportation of both carriers or any one of the two carriers in a semiconductor. When a semiconductor is connected in a circuit, electrons, being negatively charged, move towards positive terminal of the battery. Holes have an *effective positive charge*, and move towards negative terminal of the battery. Thus, the current through a semiconductor is carried by two types of charge carriers which move in opposite directions. This conduction mechanism makes semiconductors very useful in designing a large number of electronic devices. Figure 14.8 represents the current through a semiconductor.

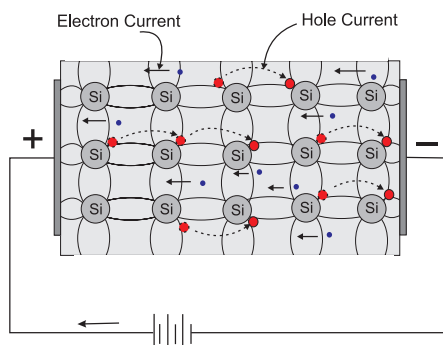


Fig. 14.8: Current through a semiconductor, transport of electrons and holes.

14.5 Extrinsic semiconductors:

The electric conductivity of an intrinsic semiconductor is very low at room temperature; hence no electronic devices can be fabricated

using them. Addition of a small amount of a suitable impurity to an intrinsic semiconductor increases its conductivity appreciably. **The process of adding impurities to an intrinsic semiconductor is called doping.** The semiconductor with impurity is called a doped semiconductor or an extrinsic semiconductor. The impurity is called the dopant. The parent atoms are called hosts. The dopant material is so selected that it does not disturb the crystal structure of the host. The size and the electronic configuration of the dopant should be compatible with that of the host. Silicon or germanium can be doped with a pentavalent impurity such as phosphorus (P) arsenic (As) or antimony (Sb). They can also be doped with a trivalent impurity such as boron (B) aluminium (Al) or indium (In).

Addition of pentavalent or trivalent impurities in intrinsic semiconductors gives rise to different conduction mechanisms. This is very useful in designing many electronic devices. Extrinsic semiconductors can be of two types a) n-type semiconductor or b) p-type semiconductor.

a) n-type semiconductor: When silicon or germanium crystal is doped with a pentavalent impurity such as phosphorus, arsenic, or antimony we get n-type semiconductor. Figure 14.9 shows the schematic electronic structure of antimony.

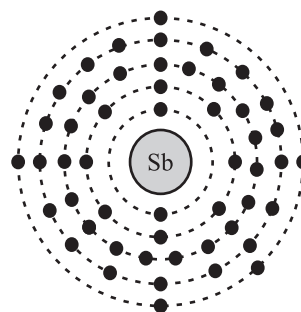


Fig. 14.9: Schematic electronic structure of antimony.

When a dopant atom of 5 valence electrons occupies the position of a Si atom in the crystal lattice, 4 electrons from the dopant form bonds with 4 neighbouring Si atoms and the fifth electron from the dopant remains very weakly bound to its parent atom. Figure 14.10 shows a pentavalent impurity in silicon lattice.

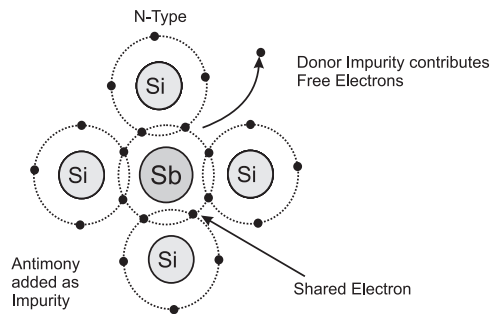


Fig. 14.10: Pentavalent impurity in silicon crystal.

To make this electron free even at room temperature, very small energy is required. It is 0.01 eV for Ge and 0.05 eV for Si.



Do you know ?

One cm^3 specimen of a metal or semiconductor has of the order of 10^{22} atoms. In a metal, every atom donates at least one free electron for conduction, thus 1 cm^3 of metal contains of the order of 10^{22} free electrons, whereas 1 cm^3 of pure germanium at 20°C contains about 4.2×10^{22} atoms, but only 2.5×10^{13} free electrons and 2.5×10^{13} holes. Addition of 0.001% of arsenic (an impurity) donates 10^{17} extra free electrons in the same volume and the electrical conductivity is increased by a factor of 10,000.

Since every pentavalent dopant atom donates one electron for conduction, it is called a donor impurity. As this semiconductor has large number of electrons in conduction band and its conductivity is due to negatively charged carriers, it is called n-type semiconductor. The n-type semiconductor also has a few electrons and holes produced due to the thermally broken bonds. The density of conduction electrons (n_e) in a doped semiconductor is the sum total of the electrons contributed by donors and the thermally generated electrons from the host. The density of holes (n_h) is only due to the thermal breakdown of some covalent bonds of the host Si atoms. Some electrons and holes recombine continuously because they carry opposite charges. The number of free electrons exceeds the number of holes. Thus, in a semiconductor doped with pentavalent impurity, electrons (negative charge) are the majority carriers

and holes are the minority carriers. Therefore, it is called n-type semiconductor. **For n-type semiconductor, $n_e \gg n_h$.**

The free electrons donated by the impurity atoms occupy energy levels which are in the band gap and are close to the conduction band. They can be easily available for conduction. Figure 14.11 shows the schematic band structure of an n-type semiconductor.

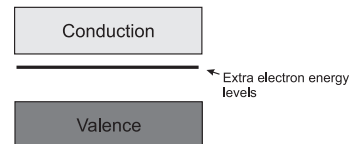


Fig.14.11: Schematic band structure of an n-type semiconductor.

Extrinsic semiconductors are thus far better conductors than intrinsic semiconductors. The conductivity of an extrinsic semiconductor can be controlled by controlling the amount of impurities added. The amount of impurities is expressed as part per million or ppm, that is, one impurity atom per one million atoms of the host.

Features of n-type semiconductors: These are materials doped with pentavalent impurity (donors) atoms. Electrical conduction in these materials is due to electrons as majority charge carriers.

1. The donor atom lose electrons and become positively charged ions.
2. Number of free electrons is very large compared to the number of holes, $n_e \gg n_h$. Electrons are majority charge carriers.
3. When energy is supplied externally, negatively charged free electrons (majority charges carries) and positively charged holes (minority charge carriers) are available for conduction.

b) p-type semiconductor: When silicon or germanium crystal is doped with a trivalent impurity such as boron, aluminium or indium, we get a p-type semiconductor. Figure 14.12 shows the schematic electronic structure of boron.

The dopant trivalent atom has one valence electron less than that of a silicon atom. Every trivalent dopant atom shares its three electrons with three neighbouring Si atoms to form

covalent bonds. But the fourth bond between silicon atom and its neighbour is not complete.

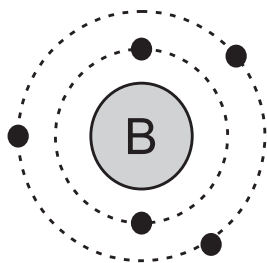


Fig. 14.12: Schematic electronic structure of boron.

Figure 14.13 shows a trivalent impurity in a silicon crystal. The incomplete bond can be completed by another electron in the neighbourhood from Si atom. Since each donor trivalent atom can accept an electron, it is called an acceptor impurity. The shared electron creates a vacancy in its place. This vacancy or the absence of electron is a hole.

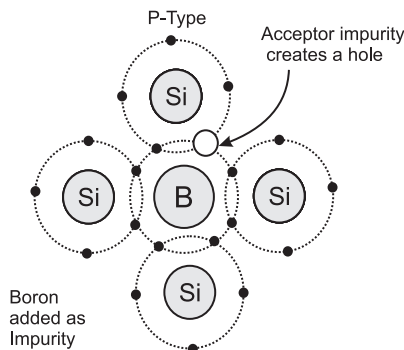


Fig. 14.13: A trivalent impurity in a silicon crystal.

Thus, a hole is available for conduction from each acceptor impurity atom. Holes are majority carriers and electrons are minority carriers in such materials. Acceptor atoms are negatively charged and majority carriers are holes (positively charged). Therefore, extrinsic semiconductor doped with trivalent impurity is called a p-type semiconductor. For a p-type semiconductor, $n_h \gg n_e$.

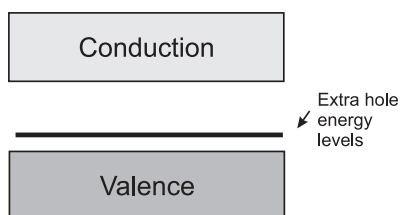


Fig. 14.14: Schematic band structure of a p-type semiconductor.

These vacancies of electrons are created in the valence band; therefore we can say that the holes are created in the valence band. The impurity levels are created just above the valence band in the band gap. Electrons from valence band can easily occupy these levels and conduct electricity. Figure 14.14 shows the schematic band structure of a p-type semiconductor.

Features of p-type semiconductor: These are materials doped with trivalent impurity atoms (acceptors). Electrical conduction in these materials is due to holes as majority charge carriers.

1. The acceptor atoms acquire electron and become negatively charged-ions.
2. Number of holes is very large compared to the number of free electrons. ($n_h \gg n_e$). Holes are majority charge carriers.
3. When energy is supplied externally, positively charged holes (majority charge carriers) and negatively charged free electrons (minority charge carriers) are available for conduction.

c) Charge neutrality of extrinsic semiconductors: The n-type semiconductor has excess of electrons but these extra electrons are supplied by the donor atoms which become positively charged. Since each atom of donor impurity is electrically neutral, the semiconductor as a whole is electrically neutral. Here, excess electron refers to an excess with reference to the number of electrons needed to complete the covalent bonds in a semiconductor crystal. These extra free electrons increase the conductivity of the semiconductor.

Similarly, a p-type semiconductor has holes or absence of electrons in some energy levels. When an electron from a host atom fills this level, the host atom is positively charged and the dopant atom is negatively charged but the semiconductor as a whole is electrically neutral. Thus, n-type as well as p-type semiconductors are electrically neutral.

Always remember, for a semiconductor,

$$n_e \cdot n_h = n_i^2$$

Example 14.1: A pure Si crystal has 4×10^{28} atoms m^{-3} . It is doped by 1ppm concentration of antimony. Calculate the number of electrons and holes. Given $n_i = 1.2 \times 10^{16}/\text{m}^3$.

Solution: 1 ppm = 1 part per million = $1/10^6$

$$\therefore \text{no. of Sb atoms} = \frac{4 \times 10^{28}}{10^6} = 4 \times 10^{22}$$

As one pentavalent impurity atom donates one free electron to the crystal,

Number of free electrons in the crystal
 $n_e = 4 \times 10^{22} \text{ m}^{-3}$

Number of holes,

$$n_h = \frac{(n_i)^2}{n_e} = \frac{(1.2 \times 10^{16})^2}{4 \times 10^{22}}$$

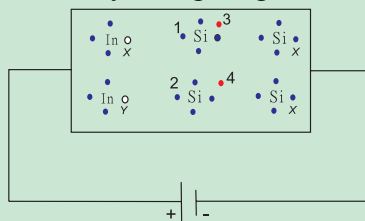
$$n_h = 3.6 \times 10^9 \text{ m}^{-3}$$



Do you know ?

Transportation of holes

Consider a p-type semiconductor connected to terminals of a battery as shown. When the circuit is switched on, electrons at 1 and 2 are attracted to the positive terminal of the battery and occupy nearby holes at x and y. This generates holes at the positions 1 and 2 previously occupied by electrons. Next, electrons at 3 and 4 move towards the positive terminal and create holes in the positions they occupied previously.



Finally, the hole is captured at the negative terminal by the electron supplied by the battery at that end. This keeps the density of holes constant and maintains the current so long as the battery is working.

Thus, physical transportation is of the electrons only. However, we feel that the holes are moving towards the negative terminal of the battery. Positive charge is attracted towards negative terminal. Thus holes, which are not actual charges, behave like a positive charge. In this case, there is an indirect movement of electrons and their drift speed is less than that in the n-type semiconductors. The mobility of holes is, therefore, less than that of the electrons.

Example 14.2: A pure silicon crystal at temperature of 300 K has electron and hole concentration $1.5 \times 10^{16} \text{ m}^{-3}$ each. ($n_e = n_h$). Doping by indium increases n_h to $4.5 \times 10^{22} \text{ m}^{-3}$. Calculate n_e for the doped silicon crystal.

Solution: We know,
 $n_e n_h = n_i^2$ and $n_e = \frac{(n_i)^2}{n_h}$

Given

$$n_i = 1.5 \times 10^{16} \text{ m}^{-3} \text{ and } n_h = 4.5 \times 10^{22} \text{ m}^{-3}$$

$$n_e = \frac{(1.5 \times 10^{16})^2}{4.5 \times 10^{22}} = 5 \times 10^9 \text{ m}^{-3}$$

14.6 p-n junction:

When n-type and p-type semiconductor materials are fused together, a p-n junction is formed. A p-n junction shows many interesting properties and it is the basis of almost all modern electronic devices. Figure 14.15 shows a schematic structure of a p-n junction.

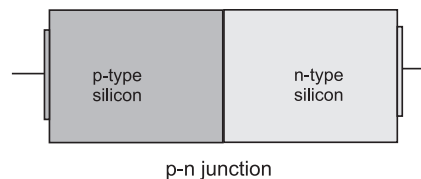


Fig. 14.15: Schematic structure of a p-n junction.

Diffusion: When n-type and p-type semiconductor materials are fused together, initially, the number of electrons in the n-side of the junction is very large compared to the number of electrons on the p-side. The same is true for the number of holes on the p-side and on the n-side. Thus, the density of carriers on both sides is different and a large density gradient exists on both sides of the p-n junction. This density gradient causes migration of electrons from the n-side to the p-side of the junction. They fill up the holes in the p-type material and produce negative ions.

When the electrons from the n-side of a junction migrate to the p-side, they leave behind positively charged donor ions on the n-side. Effectively, holes from the p-side migrate into the n-region.

As a result, in the p-type region near the junction there are negatively charged acceptor ions, and in the n-type region near the junction there are positively charged donor ions. The transfer of electrons and holes across the p-n

junction is called diffusion. The extent up to which the electrons and the holes can diffuse across the junction depends on the density of the donor and the acceptor ions on the n-side and the p-side respectively, of the junction. Figure 14.16 shows the diffusion of charge carriers across the junction.

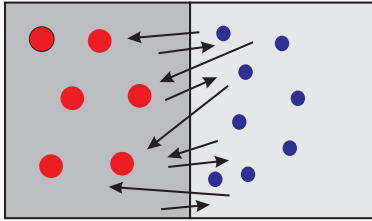


Fig. 14.16: Diffusion of charge carriers across a junction.

Depletion region: The diffusion of carriers across the junction and resultant accumulation of positive and negative charges across the junction builds a potential difference across the junction. This potential difference is called the potential barrier. The magnitude of the potential barrier for silicon is about 0.6 – 0.7 volt and for germanium, it is about 0.3 – 0.35 volt. This potential barrier always exists even if the device is not connected to any external power source. It prevents continuous diffusion of carriers across the junction. A state of electrostatic equilibrium is thus reached across the junction.

Free charge carriers cannot be present in a region where there is a potential barrier. The regions on either side of a junction, therefore, becomes completely devoid of any charge carriers. This region across the p-n junction where there are no charges is called the depletion layer or the depletion region. Figure 14.17 shows the potential barrier and the depletion layer.

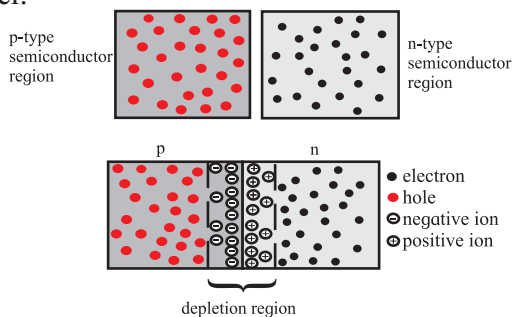


Fig. 14.17: Potential barrier and the depletion layer.

The potential across a junction and width of the potential barrier can be controlled. This is very interesting and useful property of a p-n junction.

The n-side near the boundary of a p-n junction becomes positive with respect to the p-side because it has lost electrons and the p-side has lost holes. Thus the presence of impurity ions on both sides of the junction establishes an electric field across this region such that the n-side is at a positive voltage relative to the p-side. Figure 14.18 shows the electric field thus produced.

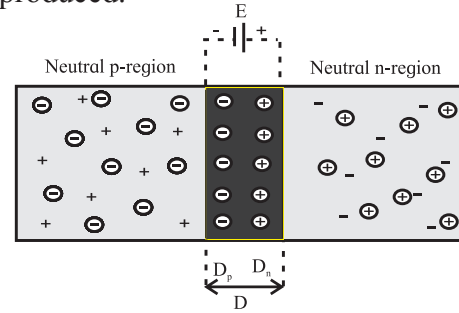


Fig. 14.18: Electric field across a junction.

Biasing a p-n junction: As a result of potential barrier across depletion region, charge carriers require some extra energy to overcome the barrier. A suitable voltage needs to be applied to the junction externally, so that these charge carriers can overcome the potential barrier and move across the junction. Figure 14.19 shows two possibilities of applying this external voltage across the junction.

Figure 14.19 (a) shows a p-n junction connected in an electric circuit where the p-region is connected to the positive terminal and the n-region is connected to the negative terminal of an external voltage source. This external voltage effectively opposes the built-in potential of the junction. The width of potential barrier is thus reduced. Also, negative charge carriers (electrons) from the n-region are pushed towards the junction. A similar effect is experienced by positive charge carriers (holes) in the p-region and they are pushed towards the junction. Both the charge carriers thus find it easy to cross over the barrier and contribute towards the electric current. Such arrangement of a p-n junction in an electric circuit is called **forward bias**.

Figure 14.19 (b) shows the other possibility, where, the p-region is connected to the negative terminal and the n-region is connected to the positive terminal of the external voltage source. This external voltage effectively adds to the

built-in potential of the junction. The width of potential barrier is thus increased. Also, the negative charge carriers (electrons) from the n-region are pulled away from the junction. Similar effect is experienced by the positive charge carriers (holes) in the p-region and they are pulled away from the junction. Both the charge carriers thus find it very difficult to cross over the barrier and thus do not contribute towards the electric current. Such arrangement of a p-n junction in an electric circuit is called **reverse bias**.

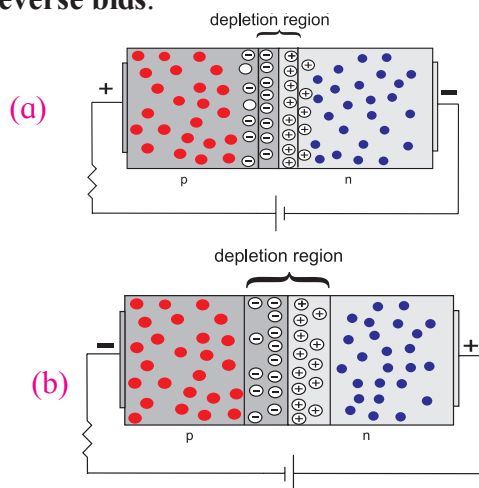


Fig. 14.19: Forward biased (a) and reverse biased (b) junction.

Therefore, when used in forward bias mode, a p-n junction allows a large current to flow across. This current is normally of the order of a few milliamperes, (10^{-3} A). A reverse biased p-n junction on the other hand, carries a very small current that is normally a few microamperes (10^{-6} A).

A p-n junction can be thus used as a one way switch or a gate in an electric circuit. It conducts easily in forward bias and acts as an open switch in reverse bias.

Features of the depletion region:

1. It is formed by diffusion of electrons from n-region to the p-region. This leaves positively charged ions in the n-region.
2. The p-region accumulates electrons (negative charges) and the n-region accumulates the holes (positive charges).
3. The accumulation of charges on either sides of the junction results in forming a potential barrier and prevents flow of charges across it.

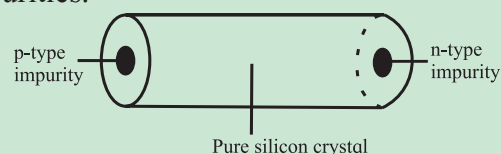
4. There are no charges in this region.
5. The depletion region has higher potential on the n-side and lower potential on the p-side of the junction.



Do you know ?

Fabrication of p-n junction diode:

It was mentioned previously, for easy understanding, that a p-n junction is formed by fusing a p-type and a n-type material together. However, in practice, a p-n junction is formed from a crystalline structure of silicon or germanium by adding carefully controlled amounts of donor and acceptor impurities.



The impurities grow on either side of the crystal after heating in a furnace. Electrons and holes combine at the center and the depletion region develops. A junction is thus formed. Electrodes are inserted after cutting transverse sections and hundreds of diodes are prepared. All semiconductor devices, including ICs, are fabricated by 'growing' junctions at the required locations.

Mobility of a hole is less than that of an electron and the hole current is lesser. This imbalance between the two currents is removed by increasing the doping percentage in the p-region. This ensures that the same current flows through the p-region and the n-region of the junction.

14.7 A p-n junction diode:

A p-n junction, when provided with metallic connectors on each side is called a junction diode or simply, a diode. (Diode is a device with two electrodes or di-electrodes). Figure 14.20 shows the circuit symbol for a junction diode.

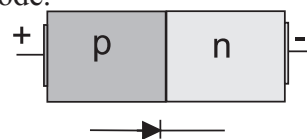


Fig. 14.20: Circuit symbol for a p-n junction diode.

The 'arrow' indicates the direction of the conventional current. The p-side is called the anode and the n-side is called the cathode of the diode. When a diode is connected across a battery, the carriers can gain additional energy to cross the barrier as per biasing.

A diode can be connected across a battery in two different ways, forward bias and reverse bias as shown in the (Fig. 14.21).

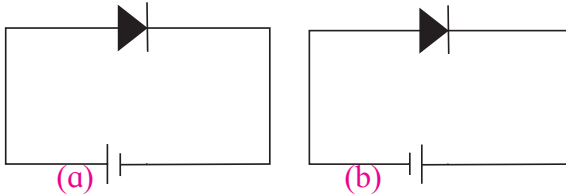


Fig. 14.21: (a) Forward bias, (b) Reverse bias.

The behavior of a diode in both cases is different. This is because the barrier potential is affected differently in the two cases. The barrier potential is reduced in forward biased mode and it is increased in reverse biased mode.

Carriers find it easy to cross the junction in forward bias and contribute towards current for two reasons; first the barrier width is reduced and second, they are pushed towards the junction and gain extra energy to cross the junction. The current through the diode in forward bias is, therefore, large. It is of the order of a few milliamperes (10^{-3} A) for a typical diode.

When connected in reverse bias, width of the potential barrier is increased and the carriers are pushed away from the junction so that very few thermally generated carriers can cross the junction and contribute towards current. This results in a very small current through a reverse biased diode. The current in reverse biased diode is of the order of a few microamperes (10^{-6} A).

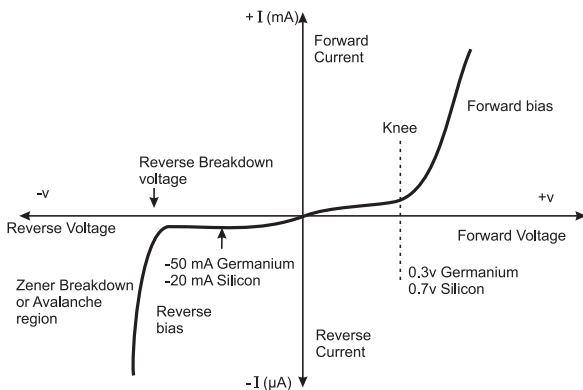


Fig. 14.22: Asymmetrical current flow through a diode.

The width of the depletion layer decreases with an increase in the application of a forward voltage. It increases when a reverse voltage is applied. We have discussed the reasons for this difference earlier. When the polarity of bias voltage is reversed, the width of the depletion layer changes. This results in asymmetrical current flow through a diode as shown in (Fig. 14.22).

A diode can be thus used as a one way switch in a circuit. It is forward biased when its anode is connected to be at a higher potential than that of the cathode. When the anode is at lower potential than that of the cathode, it is reverse biased. A diode can be zero biased if no external voltage is applied across it.

a) **Forward biased:** The positive terminal of the external voltage is connected to the anode (p-side) and negative terminal to the cathode (n-side) across the diode.

In case of forward bias, the width of the depletion region decreases and the p-n junction offers a low resistance path allowing a high current to flow across the junction (Fig. 14.23).

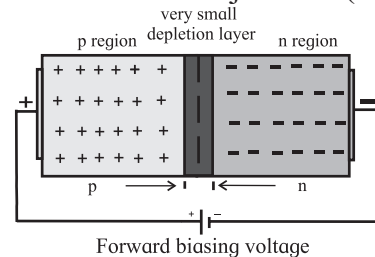


Fig. 14.23: Decrease in width of depletion region.

Figure 14.24 shows the I-V characteristic of a forward biased diode. Initially, the current is very low and then there is a sudden rise in the current. The point at which current rises sharply is shown as the 'knee' point on the I-V characteristic curve. The corresponding voltage is called the 'knee voltage'. It is about 0.7 V for silicon and 0.3 V for germanium.

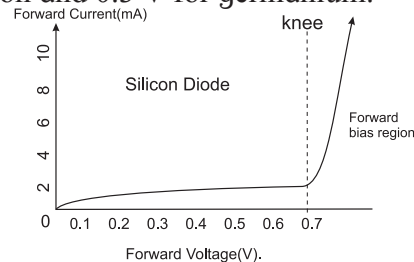


Fig. 14.24: I-V characteristic of a forward biased diode.

A diode effectively becomes a short circuit above this knee point and can conduct a very large current. Resistors are, therefore, used in series with diode to limit its current flow. If the current through a diode exceeds the specified value, it can heat up the diode due to the Joule heating and can result in its physical damage.

b) Reverse biased: The positive terminal of the external voltage is connected to the cathode (n-side) and negative terminal to the anode (p-side) across the diode. In case of reverse bias, the width of the depletion region increases and the p-n junction behaves like a high resistance (Fig. 14.25). Practically, no current flows through it with an increase in the reverse bias voltage. However, a very small **leakage current** does flow through the junction which is of the order of a few micro-amperes, (μA).

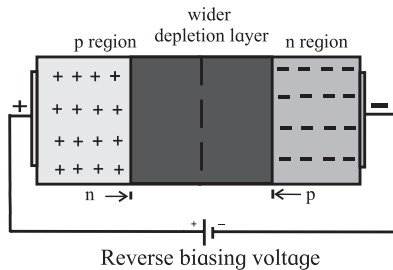


Fig. 14.25: Increase in width of depletion region.

When the reverse bias voltage applied to a diode is increased to sufficiently large value, it causes the p-n junction to overheat. The overheating of the junction results in a sudden rise in the current through the junction. This is because the covalent bonds break and a large number of carriers are available for conduction. The diode, thus, no longer behaves like a diode. This effect is called the avalanche breakdown. The reverse biased characteristic of a diode is shown in Fig 14.26.

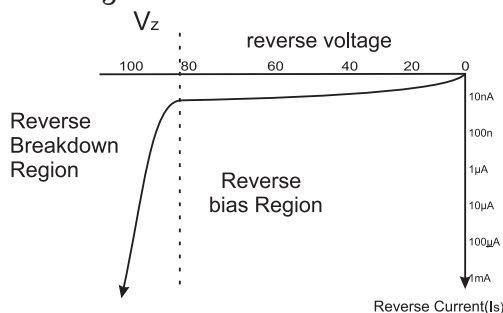
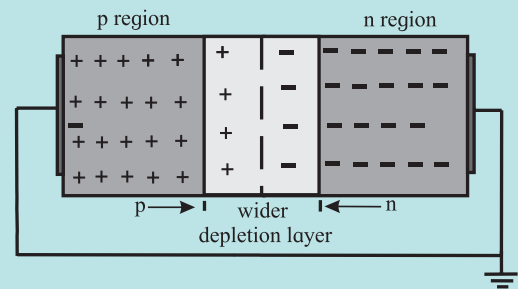


Fig. 14.26: Reverse biased characteristic of a diode.

Zero Biased Junction Diode.

When a diode is connected in a **zero bias** condition, no external potential energy is applied to the p-n junction. When the diode terminals are shorted together, some holes (majority carriers) in the p-side have enough thermal energy to overcome the potential barrier. Such carriers cross the barrier potential and contribute to current. This current is known as the **forward current**.

Similarly, some holes generated in the n-side (minority carriers), also move across the junction in the opposite direction and contribute to current. This current is known as the **reverse current**. This transfer of electrons and holes back and forth across the p-n junction is known as diffusion, as discussed previously.



Zero biased p-n junction diode

The potential barrier that exists in a junction prevents the diffusion of any more majority carriers across it. However, some minority carriers (few free electrons in the p-region and few holes in the n-region) do drift across the junction.

An equilibrium is established when the majority carriers are equal in number ($n_e = n_h$) and are moving in opposite directions. The net current flowing across the junction is zero. This is a state of '**dynamic equilibrium**'.

Minority carriers are continuously generated due to thermal energy. When the temperature of the p-n junction is raised, this state of equilibrium is changed. This results in generating more minority carriers and an increase in the leakage current. An electric current, however, cannot flow through the diode because it is not connected in any electric circuit.

c) Static and dynamic resistance of a diode:

One of the most important properties of a diode is its resistance in the forward biased mode and in the reverse biased mode. Figure 14.27 shows the I-V characteristics of an ideal diode.

An ideal diode offers zero resistance in forward biased mode and infinite resistance in reverse biased mode.

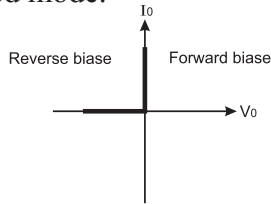


Fig. 14.27: I-V characteristics of an ideal diode.

The I-V characteristics of a forward biased diode (Fig. 14.24) is used to define two of its resistances i) the static (DC) resistance and ii) the dynamic (AC) resistance.

i) Static (DC) resistance: When a p-n junction diode is forward biased, it offers a definite resistance in the circuit. This resistance is called the static or DC resistance (R_g) of a diode. The DC resistance of a diode is the ratio of the DC voltage across the diode to the DC current flowing through it at a particular voltage.

$$R_g = \frac{V}{I}$$

ii) Dynamic (AC) resistance: The dynamic (AC) resistance of a diode, r_g , at a particular applied voltage, is defined as

$$r_g = \frac{\Delta V}{\Delta I}$$

The dynamic resistance of a diode depends on the operating voltage. It is the reciprocal of the slope of the characteristics at that point. Figure 14.28 shows how the DC and the AC resistance of a diode are found out.

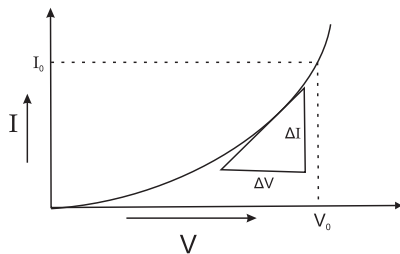
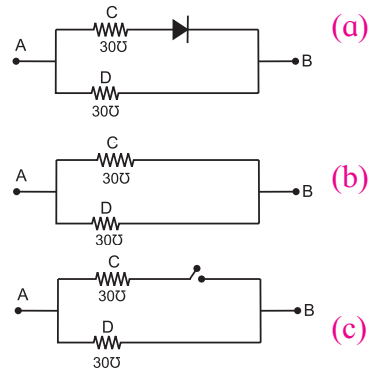


Fig. 14.28: DC and the AC resistance of a diode.

Example 14.3 Refer to the figure a shown below and find the resistance between point A and B when an ideal diode is (1) forward biased and (2) reverse biased.



Solution: We know that for an ideal diode, the resistance is zero when forward biased and infinite when reverse biased.

- i) Figure b shows the circuit when the diode is forward biased. An ideal diode behaves as a conductor and the circuit is similar to two resistances in parallel.

$$R_{AB} = (30 \times 30)/(30+30) = 900/60 = 15 \Omega$$

- ii) Figure c shows the circuit when the diode is reverse biased. It does not conduct and behaves as an open switch, path ACB. Therefore, $R_{AB} = 30 \Omega$, the only resistance in the circuit along the path ADB.

14.8 Semiconductor devices:

Semiconductor devices find applications in variety of fields. They have many advantages. They also have some disadvantages. Here we discuss some advantages and disadvantages.

14.8.1 Advantages:

1. Electronic properties of semiconductors can be controlled to suit our requirement.
2. They are smaller in size and light weight.
3. They can operate at smaller voltages (of the order of few mV) and require less current (of the order of μA or mA), therefore, consume lesser power.
4. Almost no heating effects occur, therefore these devices are thermally stable.
5. Faster speed of operation due to smaller size.
6. Fabrication of ICs is possible.

14.8.2 Disadvantages:

1. They are sensitive to electrostatic charges.
2. Not very useful for controlling high power.
3. They are sensitive to radiation.
4. They are sensitive to fluctuations in temperature.
5. They need controlled conditions for their manufacturing.
6. Very few materials are semiconductors.

14.9 Applications of semiconductors and p-n junction diode:

A p-n junction diode is the basic block of a number of semiconductor devices. A semiconductor device can have more than one junction. Properties of a device can be controlled by controlling the concentration of dopants.

1. **Solar cell:** Converts light energy into electric energy. Useful to produce electricity in remote areas and also for providing electricity for satellites, space probes and space stations.
2. **Photo resistor:** Changes its resistance when light is incident on it.
3. **Bi-polar junction transistor:** These are devices with two junctions and three terminals. A transistor can be a p-n-p or n-p-n transistor. Conduction takes place with holes and electrons. Many other types of transistors are designed and fabricated to suit specific requirements. They are used in almost all semiconductor devices.
4. **Photodiode:** It conducts when illuminated with light.
5. **LED:** Light Emitting Diode: Emits light when current passes through it. House hold LED lamps use similar technology. They consume less power, are smaller in size and have a longer life and are cost effective.
6. **Solid State Laser:** It is a special type of LED. It emits light of specific frequency. It is smaller in size and consumes less power.
7. **Integrated Circuits (ICs):** A small device having hundreds of diodes and transistors performs the work of a large number of electronic circuits.

14.10 Thermistor:

Thermistor is a temperature sensitive resistor. Its resistance changes with change in its temperature. There are two types of thermistors,

the Negative Temperature Coefficient (NTC) and the Positive Temperature Coefficient (PTC).

Resistance of a NTC thermistor decreases with increase in its temperature. Its temperature coefficient is negative. They are commonly used as temperature sensors and also in temperature control circuits.

Resistance of a PTC thermistor increases with increase in its temperature. They are commonly used in series with a circuit. They are generally used as a reusable fuse to limit current passing through a circuit to protect against *over current* conditions, as resettable fuses.

Thermistors are made from thermally sensitive metal oxide semiconductors. Thermistors are very sensitive to changes in temperature. A small change in surrounding temperature causes a large change in their resistance. They can measure temperature variations of a small area due to their small size. Both types of thermistors have many applications in industry.



Do you know ?

Electric and electronic devices

Electric devices: These devices convert electrical energy into some other form. Fan, refrigerator, geyser etc. are some examples. Fan converts electrical energy into mechanical energy. A geyser converts it into heat energy. They use good conductors (mostly metals) for conduction of electricity. Common working range of currents for electric circuits is milli amperes (mA) to amperes. Their energy consumption is also moderate to high. A typical geyser consumes about 2.0 to 2.50 kW of power. They are moderate to large in size and are costly.

Electronic devices: Electronic circuits work with control or sequential changes in current through a cell. A calculator, a cell phone a smart watch or the remote control of a TV set are some of the electronic devices. Semiconductors are used to fabricate such devices. Common working range of currents for electronic circuits it is nano-ampere to μA . They consume very low energy. They are very compact, and cost effective.



Internet my friend

1. <https://www.electronics-tutorials.ws>diode>
2. <https://www.hitachi-hightech.com>
3. <https://ntpel.ac.in>courses>
4. <https://physics.info>semiconductors>
5. <https://www.hyperphysics.phy-astr.gsu.edu>semcn>



Exercises

1. Choose the correct option.

- i) Electric conduction through a semiconductor is due to:
 - (A) electrons
 - (B) holes
 - (C) none of these
 - (D) both electrons and holes
- ii) The energy levels of holes are:
 - (A) in the valence band
 - (B) in the conduction band
 - (C) in the band gap but close to valence band
 - (D) in the band gap but close to conduction band
- iii) Current through a reverse biased p-n junction, increases abruptly at:
 - (A) breakdown voltage
 - (B) 0.0 V
 - (C) 0.3V
 - (D) 0.7V
- iv) A reverse biased diode, is equivalent to:
 - (A) an off switch
 - (B) an on switch
 - (C) a low resistance
 - (D) none of the above
- v) The potential barrier in p-n diode is due to:
 - (A) depletion of positive charges near the junction
 - (B) accumulation of positive charges near the junction
 - (C) depletion of negative charges near the junction,
 - (D) accumulation of positive and negative charges near the junction

2. Answer the following questions.

- i) What is the importance of energy gap in a semiconductor?
- ii) Which element would you use as an impurity to make germanium an n-type semiconductor?
- iii) What causes a larger current through a p-n junction diode when forward biased?
- iv) On which factors does the electrical conductivity of a pure semiconductor depend at a given temperature?
- v) Why is the conductivity of a n-type semiconductor greater than that of p-type semiconductor even when both of these have same level of doping?

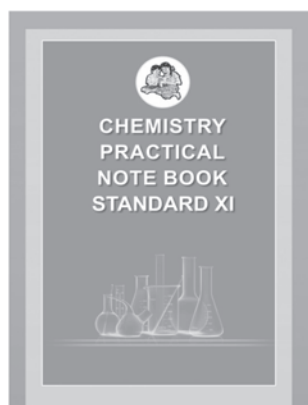
3. Answer in detail.

- i) Explain how solids are classified on the basis of band theory of solids.
- ii) Distinguish between intrinsic semiconductors and extrinsic semiconductors.
- iii) Explain the importance of the depletion region in a p-n junction diode.
- iv) Explain the I-V characteristic of a forward biased junction diode.
- v) Discuss the effect of external voltage on the width of depletion region of a p-n junction

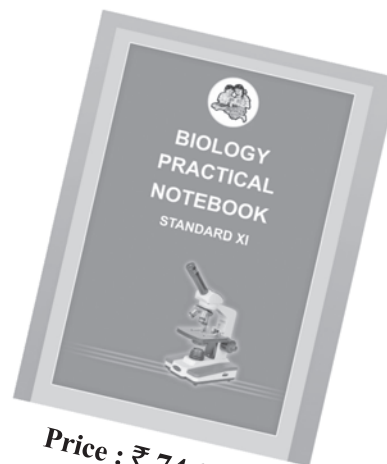
Practical Notebook for Standard XI ***Practical Notebook Cum Journal***



Price : ₹ 63.00



Price: ₹ 59.00



Price : ₹ 74.00

- Based on Government approved syllabus and textbook
- Inclusion of all practicals as per Evaluation scheme.
- Inclusion of various activities and objective questions
- Inclusion of useful questions for oral examination
- Separate space for writing answers

Practical notebooks are available for sale in the regional depots of the Textbook Bureau.

(1) Maharashtra State Textbook Stores and Distribution Centre, Senapati Bapat Marg, Pune 411004 ☎ 25659465
 (2) Maharashtra State Textbook Stores and Distribution Centre, P-41, Industrial Estate, Mumbai - Bengaluru Highway, Opposite Sakal Office, Kolhapur 416122 ☎ 2468576 (3) Maharashtra State Textbook Stores and Distribution Centre, 10, Udyognagar, S. V. Road, Goregaon (West), Mumbai 400062 ☎ 28771842
 (4) Maharashtra State Textbook Stores and Distribution Centre, CIDCO, Plot no. 14, W-Sector 12, Wavanja Road, New Panvel, Dist. Rajgad, Panvel 410206 ☎ 274626465 (5) Maharashtra State Textbook Stores and Distribution Centre, Near Lekhanagar, Plot no. 24, 'MAGH' Sector, CIDCO, New Mumbai-Agra Road, Nashik 422009 ☎ 2391511 (6) Maharashtra State Textbook Stores and Distribution Centre, M.I.D.C. Shed no. 2 and 3, Near Railway Station, Aurangabad 431001 ☎ 2332171 (7) Maharashtra State Textbook Stores and Distribution Centre, Opposite Rabindranath Tagore Science College, Maharaj Baug Road, Nagpur 440001 ☎ 2547716/2523078 (8) Maharashtra State Textbook Stores and Distribution Centre, Plot no. F-91, M.I.D.C., Latur 413531 ☎ 220930 (9) Maharashtra State Textbook Stores and Distribution Centre, Shakuntal Colony, Behind V.M.V. College, Amravati 444604 ☎ 2530965



ebalbharati

E-learning material (Audio-Visual) for Standards One to Twelve is available through Textbook Bureau, Balbharati...

- Register your demand by scanning the Q.R. Code given alongside.
- Register your demand for E-learning material by using Google play store and downloading ebalbharati app.
www.ebalbharati.in, www.balbharati.in





**Maharashtra State Bureau of Textbook Production and
Curriculum Research, Pune.**

भौतिकशास्त्र इयत्ता ११ वी (इंग्रजी माध्यम)

₹ 145.00