

PHYSICS

PART – II

TEXTBOOK FOR CLASS XII

© NCERT
not to be republished

© NCERT
not to be republished

PHYSICS

PART – II

TEXTBOOK FOR CLASS XII



12090



एन सी ई आर टी
NCERT

राष्ट्रीय शैक्षिक अनुसंधान और प्रशिक्षण परिषद्
NATIONAL COUNCIL OF EDUCATIONAL RESEARCH AND TRAINING

12090 – PHYSICS PART-II
Textbook for Class XII

ISBN 81-7450-631-4 (Part I)
ISBN 81-7450-671-3 (Part II)

First Edition

March 2007 Chaitra 1928

Reprinted

*December 2007, December 2008,
December 2009, January 2011,
January 2012, December 2012,
November 2013, December 2014,
December 2015, February 2017,
December 2017, February 2018,
February 2019, October 2019,
January 2021 and November 2021*

Revised Edition

November 2022 Agrahayana 1944

Reprinted

*March 2024 Chaitra 1946
December 2024 Pausha 1946
December 2025 Agrahayana 1947*

PD 165T HK

**© National Council of Educational
Research and Training, 2007, 2022**

₹ 105.00

*Printed on 80 GSM paper with NCERT
watermark*

Published at the Publication Division by the
Secretary, National Council of Educational
Research and Training, Sri Aurobindo
Marg, New Delhi 110 016 and printed at
Shagun Offset Press, Plot No. 82, Ecotech
12, Greater Noida 201 009 (UP)

ALL RIGHTS RESERVED

- ❑ No part of this publication may be reproduced, stored in a retrieval system or transmitted, in any form or by any means, electronic, mechanical, photocopying, recording or otherwise without the prior permission of the publisher.
- ❑ This book is sold subject to the condition that it shall not, by way of trade, be lent, re-sold, hired out or otherwise disposed of without the publisher's consent, in any form of binding or cover other than that in which it is published.
- ❑ The correct price of this publication is the price printed on this page. Any revised price indicated by a rubber stamp or by a sticker or by any other means is incorrect and should be unacceptable.

**OFFICES OF THE PUBLICATION
DIVISION, NCERT**

NCERT Campus
Sri Aurobindo Marg
New Delhi 110 016

Phone : 011-26562708

108, 100 Feet Road
Hosdakere Halli Extension
Banashankari III Stage
Bengaluru 560 085

Phone : 080-26725740

Navjivan Trust Building
P.O. Navjivan
Ahmedabad 380 014

Phone : 079-27541446

CWC Campus
Opp. Dhankal Bus Stop
Panihati
Kolkata 700 114

Phone : 033-25530454

CWC Complex
Maligaon
Guwahati 781 021

Phone : 0361-2674869

Publication Team

Head, Publication Division	: <i>M.V. Srinivasan</i>
Chief Editor	: <i>Bijnan Sutar</i>
Chief Business Manager	: <i>Amitabh Kumar</i>
Chief Production Officer (In charge)	: <i>Deepak Jaiswal</i>
Assistant Editor	: <i>R. N. Bhardwaj</i>
Production Assistant	: <i>Manoj Kumar</i>

Cover, Layout and Illustrations
Shweta Rao

FOREWORD

The National Curriculum Framework (NCF), 2005 recommends that children's life at school must be linked to their life outside the school. This principle marks a departure from the legacy of bookish learning which continues to shape our system and causes a gap between the school, home and community. The syllabi and textbooks developed on the basis of NCF signify an attempt to implement this basic idea. They also attempt to discourage rote learning and the maintenance of sharp boundaries between different subject areas. We hope these measures will take us significantly further in the direction of a child-centred system of education outlined in the National Policy on Education (NPE), 1986.

The success of this effort depends on the steps that school principals and teachers will take to encourage children to reflect on their own learning and to pursue imaginative activities and questions. We must recognise that, given space, time and freedom, children generate new knowledge by engaging with the information passed on to them by adults. Treating the prescribed textbook as the sole basis of examination is one of the key reasons why other resources and sites of learning are ignored. Inculcating creativity and initiative is possible if we perceive and treat children as participants in learning, not as receivers of a fixed body of knowledge.

These aims imply considerable change in school routines and mode of functioning. Flexibility in the daily time-table is as necessary as rigour in implementing the annual calendar so that the required number of teaching days are actually devoted to teaching. The methods used for teaching and evaluation will also determine how effective this textbook proves for making children's life at school a happy experience, rather than a source of stress or boredom. Syllabus designers have tried to address the problem of curricular burden by restructuring and reorienting knowledge at different stages with greater consideration for child psychology and the time available for teaching. The textbook attempts to enhance this endeavour by giving higher priority and space to opportunities for contemplation and wondering, discussion in small groups, and activities requiring hands-on experience.

The National Council of Educational Research and Training (NCERT) appreciates the hard work done by the textbook development committee responsible for this book. We wish to thank the Chairperson of the advisory group in science and mathematics, Professor J.V. Narlikar and the Chief Advisor for this book, Professor A.W. Joshi for guiding the work of this committee. Several teachers contributed to the development of this textbook; we are grateful to their principals for making this possible. We are indebted to the institutions and organisations which have generously permitted us to draw upon their resources, material and personnel. We are especially grateful to the members of the National Monitoring Committee, appointed by the Department of Secondary and Higher Education, Ministry of Human Resource Development under the Chairpersonship of Professor Mrinal Miri and Professor G.P. Deshpande, for their valuable time and contribution. As an organisation committed to systemic reform and continuous improvement in the quality of its products, NCERT welcomes comments and suggestions which will enable us to undertake further revision and refinement.

New Delhi
20 November 2006

Director
National Council of Educational
Research and Training

© NCERT
not to be republished

RATIONALISATION OF CONTENT IN THE TEXTBOOKS

In view of the COVID-19 pandemic, it is imperative to reduce content load on students. The National Education Policy 2020, also emphasises reducing the content load and providing opportunities for experiential learning with creative mindset. In this background, the NCERT has undertaken the exercise to rationalise the textbooks across all classes. Learning Outcomes already developed by the NCERT across classes have been taken into consideration in this exercise.

Contents of the textbooks have been rationalised in view of the following:

- Overlapping with similar content included in other subject areas in the same class
- Similar content included in the lower or higher class in the same subject
- Difficulty level
- Content, which is easily accessible to students without much interventions from teachers and can be learned by children through self-learning or peer-learning
- Content, which is irrelevant in the present context

This present edition, is a reformatted version after carrying out the changes given above.

© NCERT
not to be republished

PREFACE

It gives me pleasure to place this book in the hands of the students, teachers and the public at large (whose role cannot be overlooked). It is a natural sequel to the Class XI textbook which was brought out in 2006. This book is also a trimmed version of the textbooks which existed so far. The chapter on thermal and chemical effects of current has been cut out. This topic has also been dropped from the CBSE syllabus. Similarly, the chapter on communications has been substantially curtailed. It has been rewritten in an easily comprehensible form.

Although most other chapters have been based on the earlier versions, several parts and sections in them have been rewritten. The Development Team has been guided by the feedback received from innumerable teachers across the country.

In producing these books, Class XI as well as Class XII, there has been a basic change of emphasis. Both the books present physics to students without assuming that they would pursue this subject beyond the higher secondary level. This new view has been prompted by the various observations and suggestions made in the National Curriculum Framework (NCF), 2005. Similarly, in today's educational scenario where students can opt for various combinations of subjects, we cannot assume that a physics student is also studying mathematics. Therefore, physics has to be presented, so to say, in a stand-alone form.

As in Class XI textbook, some interesting box items have been inserted in many chapters. They are not meant for teaching or examinations. Their purpose is to catch the attention of the reader, to show some applications in daily life or in other areas of science and technology, to suggest a simple experiment, to show connection of concepts in different areas of physics, and in general, to break the monotony and enliven the book.

Features like Summary, Points to Ponder, Exercises and Additional Exercises at the end of each chapter, and Examples have been retained. Several concept-based Exercises have been transferred from end-of-chapter Exercises to Examples with Solutions in the text. It is hoped that this will make the concepts discussed in the chapter more comprehensible. Several new examples and exercises have been added. Students wishing to pursue physics further would find Points to Ponder and Additional Exercises very useful and thoughtful. To provide *resources beyond the textbook* and to encourage *eLearning*, each chapter has been provided with some relevant website addresses under the title *ePhysics*. These sites provide additional materials on specific topics and also provide learners the opportunities for interactive demonstrations/experiments.

The intricate concepts of physics must be understood, comprehended and appreciated. Students must learn to ask questions like 'why', 'how', 'how do we know it'. They will find almost always that the question 'why' has no answer within the domain of physics and science in general. But that itself is a learning experience, is it not? On the other hand, the question 'how' has been reasonably well answered by physicists in the case of most natural phenomena. In fact, with the understanding of how things happen, it has been possible to make use of many phenomena to create technological applications for the use of humans.

For example, consider statements in a book, like 'A negatively charged electron is attracted by the positively charged plate', or 'In this experiment, light (or electron) behaves like a wave'. You will realise that it is not possible to answer 'why'. This question belongs to the domain of philosophy or metaphysics. But we can answer 'how', we can find the force acting, we can find the wavelength of the photon (or electron), we can determine how things behave under different conditions, and we can develop instruments which will use these phenomena to our advantage.

It has been a pleasure to work for these books at the higher secondary level, along with a team of members. The Textbook Development Team, the Review Team and Editing Teams involved college and university teachers, teachers from Indian Institutes of Technology, scientists from national institutes and laboratories, as well as higher secondary teachers. The feedback and critical look provided by higher secondary teachers in the various teams are highly laudable. Most box items were generated by members of one or the other team, but three of them were generated by friends and well-wishers not part of any team. We are thankful to Dr P.N. Sen of Pune, Professor Roopmanjari Ghosh of Delhi and Dr Rajesh B Khaparde of Mumbai for allowing us to use their box items, respectively in Chapters 3, 4 (Part I) and 9 (Part II). We are very thankful to the members of the Review and Editing Workshops to discuss and refine the first draft of the textbook. We also express our gratitude to Prof. Krishna Kumar, Director, NCERT, for entrusting us with the task of presenting this textbook as a part of the national effort for improving science education. I also thank Prof. G. Ravindra, Joint Director, NCERT, for his help from time-to-time. Prof. Hukum Singh, Head, Department of Education in Science and Mathematics, NCERT, was always willing to help us in our endeavour in every possible way.

We welcome suggestions and comments from our valued users, especially students and teachers. We wish our young readers a happy journey into the exciting realm of physics.

A. W. JOSHI
Chief Advisor
Textbook Development Committee

TEXTBOOK DEVELOPMENT COMMITTEE

CHAIRPERSON, ADVISORY COMMITTEE FOR TEXTBOOKS IN SCIENCE AND MATHEMATICS

J.V. Narlikar, *Emeritus Professor*, Inter-University Centre for Astronomy and Astrophysics (IUCAA), Ganeshkhind, Pune University Campus, Pune

CHIEF ADVISOR

A.W. Joshi, *Honorary Visiting Scientist*, National Centre for Radio Astrophysics (NCRA), Pune University Campus, Pune (Formerly *Professor* at Department of Physics, University of Pune)

MEMBERS

A.K. Ghatak, *Emeritus Professor*, Department of Physics, Indian Institute of Technology, New Delhi

Alika Khare, *Professor*, Department of Physics, Indian Institute of Technology, Guwahati

Anjali Kshirsagar, *Reader*, Department of Physics, University of Pune, Pune

Anuradha Mathur, *PGT*, Modern School, Vasant Vihar, New Delhi

Atul Mody, *Lecturer (S.G.)*, VES College of Arts, Science and Commerce, Mumbai

B.K. Sharma, *Professor*, DESM, NCERT, New Delhi

Chitra Goel, *PGT*, Rajkiya Pratibha Vikas Vidyalaya, Tyagraj Nagar, New Delhi

Gagan Gupta, *Reader*, DESM, NCERT, New Delhi

H.C. Pradhan, *Professor*, Homi Bhabha Centre of Science Education (TIFR), Mumbai

N. Panchapakesan, *Professor (Retd.)*, Department of Physics and Astrophysics, University of Delhi, Delhi

R. Joshi, *Lecturer (S.G.)*, DESM, NCERT, New Delhi

S.K. Dash, *Reader*, DESM, NCERT, New Delhi

S. Rai Choudhary, *Professor*, Department of Physics and Astrophysics, University of Delhi, Delhi

S.K. Upadhyay, *PGT*, Jawahar Navodaya Vidyalaya, Muzaffar Nagar

S.N. Prabhakara, *PGT*, DM School, Regional Institute of Education (NCERT), Mysore

V.H. Raybagkar, *Reader*, Nowrosjee Wadia College, Pune

Vishwajeet Kulkarni, *Teacher (Grade I)*, Higher Secondary Section, Smt. Parvatibai Chowgule College, Margao, Goa

MEMBER-COORDINATOR

V.P. Srivastava, *Reader*, DESM, NCERT, New Delhi

THE CONSTITUTION OF INDIA

PREAMBLE

WE, PEOPLE OF INDIA, having solemnly resolved to constitute India into a ¹**[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 ²[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.**

1. Subs. by the Constitution (Forty-second Amendment) Act, 1976, Sec.2, for "Sovereign Democratic Republic" (w.e.f. 3.1.1977)
2. Subs. by the Constitution (Forty-second Amendment) Act, 1976, Sec.2, for "Unity of the Nation" (w.e.f. 3.1.1977)

ACKNOWLEDGEMENTS

The National Council of Educational Research and Training acknowledges the valuable contribution of the individuals and organisations involved in the development of Physics Textbook for Class XII. The Council also acknowledges the valuable contribution of the following academics for reviewing and refining the manuscripts of this book:

Anu Venugopalan, *Lecturer*, School of Basic and Applied Sciences, GGSIP University, Delhi; A.K. Das, *PGT*, St. Xavier's Senior Secondary School, Delhi; Bharati Kukkal, *PGT*, Kendriya Vidyalaya, Pushp Vihar, New Delhi; D.A. Desai, *Lecturer (Retd.)*, Ruparel College, Mumbai; Devendra Kumar, *PGT*, Rajkiya Pratibha Vikas Vidyalaya, Yamuna Vihar, Delhi; I.K. Gogia, *PGT*, Kendriya Vidyalaya, Gole Market, New Delhi; K.C. Sharma, *Reader*, Regional Institute of Education (NCERT), Ajmer; M.K. Nandy, *Associate Professor*, Department of Physics, Indian Institute of Technology, Guwahati; M.N. Bapat, *Reader*, Regional Institute of Education (NCERT), Mysore; R. Bhattacharjee, *Asstt. Professor*, Department of Electronics and Communication Engineering, Indian Institute of Technology, Guwahati; R.S. Das, *Vice-Principal (Retd.)*, Balwant Ray Mehta Senior Secondary School, Lajpat Nagar, New Delhi; Sangeeta D. Gadre, *Reader*, Kirori Mal College, Delhi; Suresh Kumar, *PGT*, Delhi Public School, Dwarka, New Delhi; Sushma Jaireth, *Reader*, Department of Women's Studies, NCERT, New Delhi; Shyama Rath, *Reader*, Department of Physics and Astrophysics, University of Delhi, Delhi; Yashu Kumar, *PGT*, Kulachi Hans Raj Model School, Ashok Vihar, Delhi.

The Council also gratefully acknowledges the valuable contribution of the following academics for the editing and finalisation of this book: B.B. Tripathi, *Professor (Retd.)*, Department of Physics, Indian Institute of Technology, New Delhi; Dipan K. Ghosh, *Professor*, Department of Physics, Indian Institute of Technology, Mumbai; Dipanjan Mitra, *Scientist*, National Centre for Radio Astrophysics (TIFR), Pune; G.K. Mehta, *Raja Ramanna Fellow*, Inter-University Accelerator Centre, New Delhi; G.S. Visweswaran, *Professor*, Department of Electrical Engineering, Indian Institute of Technology, New Delhi; H.C. Kandpal, *Head*, Optical Radiation Standards, National Physical Laboratory, New Delhi; H.S. Mani, *Raja Ramanna Fellow*, Institute of Mathematical Sciences, Chennai; K. Thyagarajan, *Professor*, Department of Physics, Indian Institute of Technology, New Delhi; P.C. Vinod Kumar, *Professor*, Department of Physics, Sardar Patel University, Vallabh Vidyanagar, Gujarat; S. Annapoorni, *Professor*, Department of Physics and Astrophysics, University of Delhi, Delhi; S.C. Dutta Roy, *Emeritus Professor*, Department of Electrical Engineering, Indian Institute of Technology, New Delhi; S.D. Joglekar, *Professor*, Department of Physics, Indian Institute of Technology, Kanpur; V. Sundara Raja, *Professor*, Sri Venkateswara University, Tirupati.

The Council also acknowledges the valuable contributions of the following academics for refining the text in 2017: A.K. Srivastava, *Assistant Professor*, DESM, NCERT, New Delhi; Arnab Sen, *Assistant Professor*, NERIE, Shillong; L.S. Chauhan, *Assistant Professor*, RIE, Bhopal; O.N. Awasthi, *Professor (Retd.)*, RIE, Bhopal; Rachna Garg, *Professor*, DESM, NCERT, New Delhi; Raman Namboodiri, *Assistant Professor*, RIE, Mysuru; R.R. Koireng, *Assistant Professor*, DCS, NCERT, New Delhi; Shashi Prabha, *Professor*, DESM, NCERT, New Delhi; and S.V. Sharma, *Professor*, RIE, Ajmer.

Special thanks are due to Hukum Singh, *Professor and Head*, DESM, NCERT for his support.

The Council also acknowledges the support provided by the APC office and the administrative staff of the DESM; Deepak Kapoor, *Incharge*, Computer Station; Inder Kumar, *DTP Operator*; Mohd. Qamar Tabrez and Hari Darshan Lodhi *Copy Editor*; Rishi Pal Singh, *Sr. Proof Reader*, NCERT and Ashima Srivastava, *Proof Reader* in shaping this book.

The contributions of the Publication Department in bringing out this book are also duly acknowledged.

Contents of Physics Part I

Class XII

CHAPTER ONE	
ELECTRIC CHARGES AND FIELDS	1
CHAPTER TWO	
ELECTROSTATIC POTENTIAL AND CAPACITANCE	45
CHAPTER THREE	
CURRENT ELECTRICITY	81
CHAPTER FOUR	
MOVING CHARGES AND MAGNETISM	107
CHAPTER FIVE	
MAGNETISM AND MATTER	136
CHAPTER SIX	
ELECTROMAGNETIC INDUCTION	154
CHAPTER SEVEN	
ALTERNATING CURRENT	177
CHAPTER EIGHT	
ELECTROMAGNETIC WAVES	201
ANSWERS	217

CONTENTS

FOREWORD	<i>v</i>
RATIONALISATION OF CONTENT IN THE TEXTBOOK	<i>vii</i>
PREFACE	<i>ix</i>
 CHAPTER NINE	
RAY OPTICS AND OPTICAL INSTRUMENTS	
9.1 Introduction	221
9.2 Reflection of Light by Spherical Mirrors	222
9.3 Refraction	228
9.4 Total Internal Reflection	229
9.5 Refraction at Spherical Surfaces and by Lenses	232
9.6 Refraction through a Prism	239
9.7 Optical Instruments	240
 CHAPTER TEN	
WAVE OPTICS	
10.1 Introduction	255
10.2 Huygens Principle	257
10.3 Refraction and Reflection of Plane Waves using Huygens Principle	258
10.4 Coherent and Incoherent Addition of Waves	262
10.5 Interference of Light Waves and Young's Experiment	265
10.6 Diffraction	266
10.7 Polarisation	269
 CHAPTER ELEVEN	
DUAL NATURE OF RADIATION AND MATTER	
11.1 Introduction	274
11.2 Electron Emission	275

11.3	Photoelectric Effect	276
11.4	Experimental Study of Photoelectric Effect	277
11.5	Photoelectric Effect and Wave Theory of Light	280
11.6	Einstein's Photoelectric Equation: Energy Quantum of Radiation	281
11.7	Particle Nature of Light: The Photon	283
11.8	Wave Nature of Matter	284

CHAPTER TWELVE

ATOMS

12.1	Introduction	290
12.2	Alpha-particle Scattering and Rutherford's Nuclear Model of Atom	291
12.3	Atomic Spectra	296
12.4	Bohr Model of the Hydrogen Atom	297
12.5	The Line Spectra of the Hydrogen Atom	300
12.6	DE Broglie's Explanation of Bohr's Second Postulate of Quantisation	301

CHAPTER THIRTEEN

NUCLEI

13.1	Introduction	306
13.2	Atomic Masses and Composition of Nucleus	306
13.3	Size of the Nucleus	309
13.4	Mass-Energy and Nuclear Binding Energy	310
13.5	Nuclear Force	313
13.6	Radioactivity	314
13.7	Nuclear Energy	314

CHAPTER FOURTEEN

SEMICONDUCTOR ELECTRONICS: MATERIALS, DEVICES AND SIMPLE CIRCUITS

14.1	Introduction	323
14.2	Classification of Metals, Conductors and Semiconductors	324
14.3	Intrinsic Semiconductor	327
14.4	Extrinsic Semiconductor	329

14.5	p-n Junction	333
14.6	Semiconductor Diode	334
14.7	Application of Junction Diode as a Rectifier	338
APPENDICES		344
ANSWERS		346
BIBLIOGRAPHY		353

© NCERT
not to be republished

COVER DESIGN

(Adapted from <http://nobelprize.org> and
the Nobel Prize in Physics 2006)

Different stages in the evolution of
the universe.

BACK COVER

(Adapted from <http://www.iter.org> and
<http://www.dae.gov.in>)

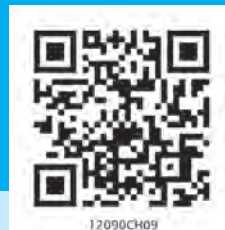
Cut away view of *International Thermonuclear Experimental Reactor* (ITER) device. The man in the bottom shows the scale.

ITER is a joint international research and development project that aims to demonstrate the scientific and technical feasibility of fusion power.

India is one of the seven full partners in the project, the others being the European Union (represented by EURATOM), Japan, the People's Republic of China, the Republic of Korea, the Russian Federation and the USA. ITER will be constructed in Europe, at Cadarache in the South of France and will provide 500 MW of fusion power.

Fusion is the energy source of the sun and the stars. On earth, fusion research is aimed at demonstrating that this energy source can be used to produce electricity in a safe and environmentally benign way, with abundant fuel resources, to meet the needs of a growing world population.

For details of India's role, see *Nuclear India*, Vol. 39, No. 11-12/ May-June 2006, issue available at Department of Atomic Energy (DAE) website mentioned above.



Chapter Nine

RAY OPTICS AND OPTICAL INSTRUMENTS



9.1 INTRODUCTION

Nature has endowed the human eye (retina) with the sensitivity to detect electromagnetic waves within a small range of the electromagnetic spectrum. Electromagnetic radiation belonging to this region of the spectrum (wavelength of about 400 nm to 750 nm) is called light. It is mainly through light and the sense of vision that we know and interpret the world around us.

There are two things that we can intuitively mention about light from common experience. First, that it travels with enormous speed and second, that it travels in a straight line. It took some time for people to realise that the speed of light is finite and measurable. Its presently accepted value in vacuum is $c = 2.99792458 \times 10^8 \text{ m s}^{-1}$. For many purposes, it suffices to take $c = 3 \times 10^8 \text{ m s}^{-1}$. The speed of light in vacuum is the highest speed attainable in nature.

The intuitive notion that light travels in a straight line seems to contradict what we have learnt in Chapter 8, that light is an electromagnetic wave of wavelength belonging to the visible part of the spectrum. How to reconcile the two facts? The answer is that the wavelength of light is very small compared to the size of ordinary objects that we encounter commonly (generally of the order of a few cm or larger). In this situation, as you will learn in Chapter 10, a light wave can be considered to travel from one point to another, along a straight line joining

them. The path is called a *ray* of light, and a bundle of such rays constitutes a *beam* of light.

In this chapter, we consider the phenomena of reflection, refraction and dispersion of light, using the ray picture of light. Using the basic laws of reflection and refraction, we shall study the image formation by plane and spherical reflecting and refracting surfaces. We then go on to describe the construction and working of some important optical instruments, including the human eye.

9.2 REFLECTION OF LIGHT BY SPHERICAL MIRRORS

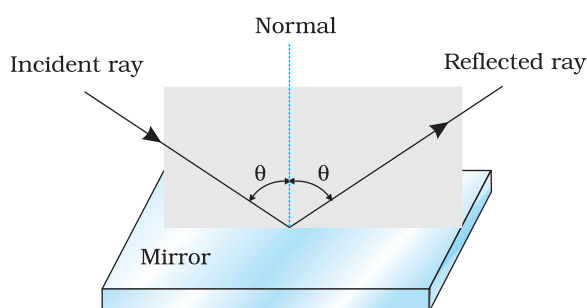


FIGURE 9.1 The incident ray, reflected ray and the normal to the reflecting surface lie in the same plane.

We are familiar with the laws of reflection. The angle of reflection (i.e., the angle between reflected ray and the normal to the reflecting surface or the mirror) equals the angle of incidence (angle between incident ray and the normal). Also that the incident ray, reflected ray and the normal to the reflecting surface at the point of incidence lie in the same plane (Fig. 9.1). These laws are valid at each point on any reflecting surface whether plane or curved. However, we shall restrict our discussion to the special case of curved surfaces, that is, spherical surfaces. The normal in this case is to be taken as normal to the tangent to surface at the point of incidence. That is, the normal is

along the radius, the line joining the centre of curvature of the mirror to the point of incidence.

We have already studied that the geometric centre of a spherical mirror is called its pole while that of a spherical lens is called its optical centre. The line joining the pole and the centre of curvature of the spherical mirror is known as the *principal axis*. In the case of spherical lenses, the principal axis is the line joining the optical centre with its principal focus as you will see later.

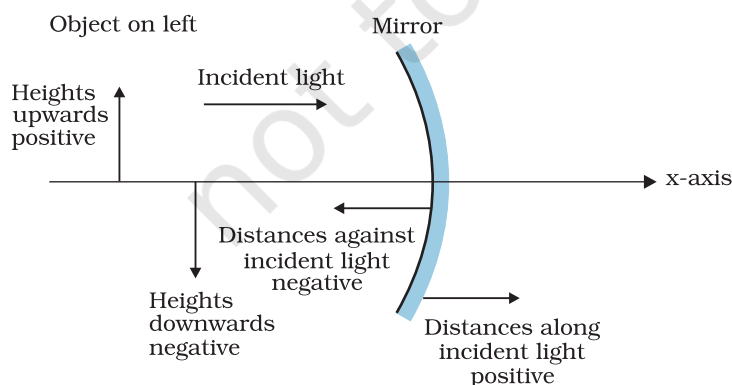


FIGURE 9.2 The Cartesian Sign Convention.

9.2.1 Sign convention

To derive the relevant formulae for reflection by spherical mirrors and refraction by spherical lenses, we must first adopt a sign convention for measuring distances. In this book, we shall follow the *Cartesian sign convention*. According to this convention, all distances are measured from the pole of the mirror or the optical centre of the lens. The distances measured in the same direction as the incident light are taken as positive and those measured in the direction

opposite to the direction of incident light are taken as negative (Fig. 9.2). The heights measured upwards with respect to x-axis and normal to the

principal axis (x -axis) of the mirror/lens are taken as positive (Fig. 9.2). The heights measured downwards are taken as negative.

With a common accepted convention, it turns out that a single formula for spherical mirrors and a single formula for spherical lenses can handle all different cases.

9.2.2 Focal length of spherical mirrors

Figure 9.3 shows what happens when a parallel beam of light is incident on (a) a concave mirror, and (b) a convex mirror. We assume that the rays are *paraxial*, i.e., they are incident at points close to the pole P of the mirror and make small angles with the principal axis. The reflected rays converge at a point F on the principal axis of a concave mirror [Fig. 9.3(a)]. For a convex mirror, the reflected rays appear to diverge from a point F on its principal axis [Fig. 9.3(b)]. The point F is called the *principal focus* of the mirror. If the parallel paraxial beam of light were incident, making some angle with the principal axis, the reflected rays would converge (or appear to diverge) from a point in a plane through F normal to the principal axis. This is called the *focal plane* of the mirror [Fig. 9.3(c)].

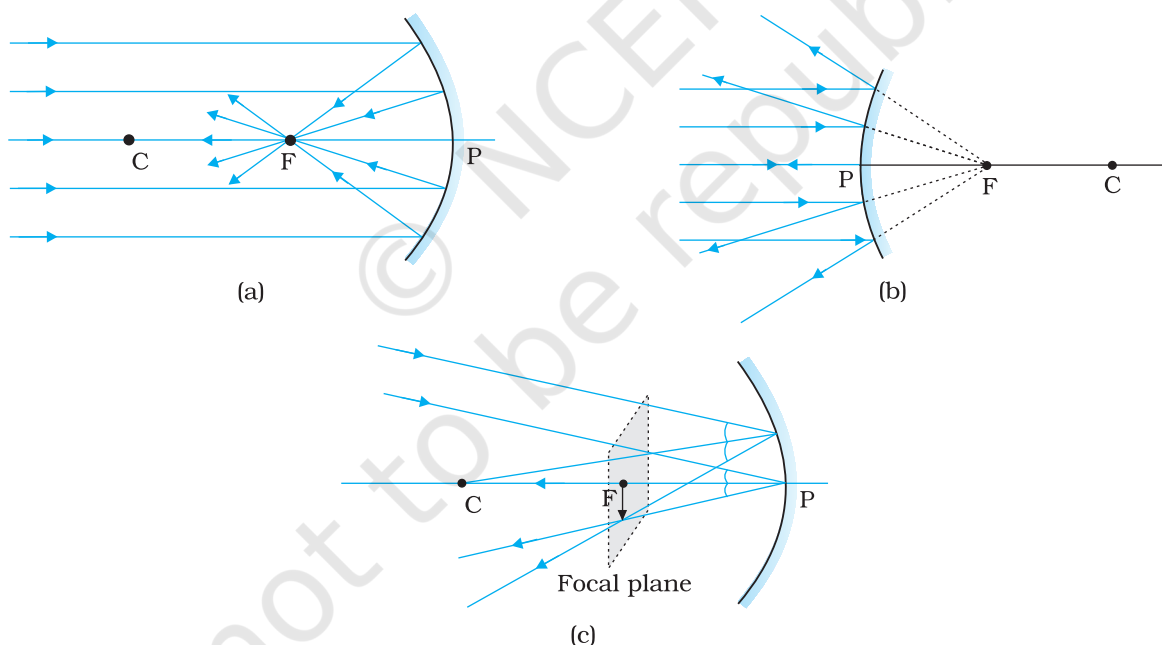


FIGURE 9.3 Focus of a concave and convex mirror.

The distance between the focus F and the pole P of the mirror is called the *focal length* of the mirror, denoted by f . We now show that $f = R/2$, where R is the radius of curvature of the mirror. The geometry of reflection of an incident ray is shown in Fig. 9.4.

Let C be the centre of curvature of the mirror. Consider a ray parallel to the principal axis striking the mirror at M . Then CM will be perpendicular to the mirror at M . Let θ be the angle of incidence, and MD

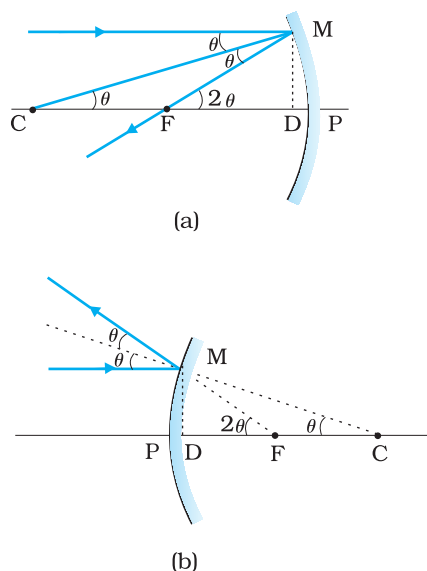


FIGURE 9.4 Geometry of reflection of an incident ray on (a) concave spherical mirror, and (b) convex spherical mirror.

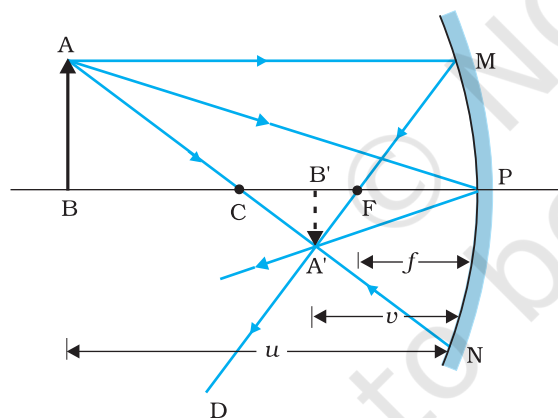


FIGURE 9.5 Ray diagram for image formation by a concave mirror.

- (i) The ray from the point which is parallel to the principal axis. The reflected ray goes through the focus of the mirror.
- (ii) The ray passing through the centre of curvature of a concave mirror or appearing to pass through it for a convex mirror. The reflected ray simply retraces the path.
- (iii) The ray passing through (or directed towards) the focus of the concave mirror or appearing to pass through (or directed towards) the focus of a convex mirror. The reflected ray is parallel to the principal axis.
- (iv) The ray incident at any angle at the pole. The reflected ray follows laws of reflection.

Figure 9.5 shows the ray diagram considering three rays. It shows the image $A'B'$ (in this case, real) of an object AB formed by a concave mirror. It does not mean that only three rays emanate from the point A . An infinite number of rays emanate from any source, in all directions. Thus, point A' is image point of A if every ray originating at point A and falling on the concave mirror after reflection passes through the point A' .

be the perpendicular from M on the principal axis. Then,
 $\angle MCP = \theta$ and $\angle MFP = 2\theta$

Now,

$$\tan \theta = \frac{MD}{CD} \text{ and } \tan 2\theta = \frac{MD}{FD} \quad (9.1)$$

For small θ , which is true for paraxial rays, $\tan \theta \approx \theta$, $\tan 2\theta \approx 2\theta$. Therefore, Eq. (9.1) gives

$$\frac{MD}{FD} = 2 \frac{MD}{CD}$$

$$\text{or, } FD = \frac{CD}{2} \quad (9.2)$$

Now, for small θ , the point D is very close to the point P . Therefore, $FD = f$ and $CD = R$. Equation (9.2) then gives

$$f = R/2 \quad (9.3)$$

9.2.3 The mirror equation

If rays emanating from a point actually meet at another point after reflection and/or refraction, that point is called the *image* of the first point. The image is *real* if the rays actually converge to the point; it is *virtual* if the rays do not actually meet but appear to diverge from the point when produced backwards. An image is thus a point-to-point correspondence with the object established through reflection and/or refraction.

In principle, we can take any two rays emanating from a point on an object, trace their paths, find their point of intersection and thus, obtain the image of the point due to reflection at a spherical mirror. In practice, however, it is convenient to choose any two of the following rays:

- (i) The ray from the point which is parallel to the principal axis. The reflected ray goes through the focus of the mirror.
- (ii) The ray passing through the centre of curvature of a concave mirror or appearing to pass through it for a convex mirror. The reflected ray simply retraces the path.

We now derive the mirror equation or the relation between the object distance (u), image distance (v) and the focal length (f).

From Fig. 9.5, the two right-angled triangles $A'B'F$ and MPF are similar. (For paraxial rays, MP can be considered to be a straight line perpendicular to CP .) Therefore,

$$\frac{B'A'}{PM} = \frac{B'F}{FP}$$

$$\text{or } \frac{B'A'}{BA} = \frac{B'F}{FP} \quad (\because PM = AB) \quad (9.4)$$

Since $\angle APB = \angle A'PB'$, the right angled triangles $A'B'P$ and ABP are also similar. Therefore,

$$\frac{B'A'}{BA} = \frac{B'P}{BP} \quad (9.5)$$

Comparing Eqs. (9.4) and (9.5), we get

$$\frac{B'F}{FP} = \frac{B'P - FP}{FP} = \frac{B'P}{BP} \quad (9.6)$$

Equation (9.6) is a relation involving magnitude of distances. We now apply the sign convention. We note that light travels from the object to the mirror MPN . Hence this is taken as the positive direction. To reach the object AB , image $A'B'$ as well as the focus F from the pole P , we have to travel opposite to the direction of incident light. Hence, all the three will have negative signs. Thus,

$$B'P = -v, FP = -f, BP = -u$$

Using these in Eq. (9.6), we get

$$\frac{-v + f}{-f} = \frac{-v}{-u}$$

$$\text{or } \frac{v - f}{f} = \frac{v}{u}$$

$$\frac{v}{f} = 1 + \frac{v}{u}$$

Dividing it by v , we get

$$\frac{1}{v} + \frac{1}{u} = \frac{1}{f} \quad (9.7)$$

This relation is known as the *mirror equation*.

The size of the image relative to the size of the object is another important quantity to consider. We define linear *magnification* (m) as the ratio of the height of the image (h') to the height of the object (h):

$$m = \frac{h'}{h} \quad (9.8)$$

h and h' will be taken positive or negative in accordance with the accepted sign convention. In triangles $A'B'P$ and ABP , we have,

$$\frac{B'A'}{BA} = \frac{B'P}{BP}$$

With the sign convention, this becomes

$$\frac{-h'}{h} = \frac{-v}{-u}$$

so that

$$m = \frac{h'}{h} = -\frac{v}{u} \quad (9.9)$$

We have derived here the mirror equation, Eq. (9.7), and the magnification formula, Eq. (9.9), for the case of real, inverted image formed by a concave mirror. With the proper use of sign convention, these are, in fact, valid for all the cases of reflection by a spherical mirror (concave or convex) whether the image formed is real or virtual. Figure 9.6 shows the ray diagrams for virtual image formed by a concave and convex mirror. You should verify that Eqs. (9.7) and (9.9) are valid for these cases as well.

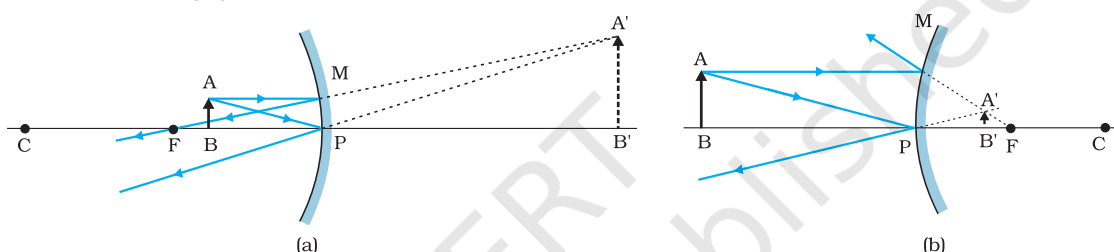


FIGURE 9.6 Image formation by (a) a concave mirror with object between P and F, and (b) a convex mirror.

EXAMPLE 9.1

Example 9.1 Suppose that the lower half of the concave mirror's reflecting surface in Fig. 9.6 is covered with an opaque (non-reflective) material. What effect will this have on the image of an object placed in front of the mirror?

Solution You may think that the image will now show only half of the object, but taking the laws of reflection to be true for all points of the remaining part of the mirror, the image will be that of the whole object. However, as the area of the reflecting surface has been reduced, the intensity of the image will be low (in this case, half).

EXAMPLE 9.2

Example 9.2 A mobile phone lies along the principal axis of a concave mirror, as shown in Fig. 9.7. Show by suitable diagram, the formation of its image. Explain why the magnification is not uniform. Will the distortion of image depend on the location of the phone with respect to the mirror?

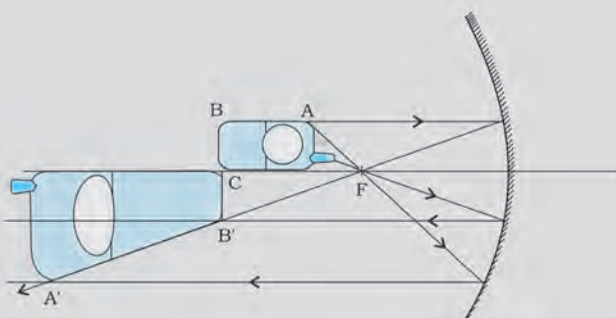


FIGURE 9.7

Solution

The ray diagram for the formation of the image of the phone is shown in Fig. 9.7. The image of the part which is on the plane perpendicular to principal axis will be on the same plane. It will be of the same size, i.e., $B'C = BC$. You can yourself realise why the image is distorted.

EXAMPLE 9.2

Example 9.3 An object is placed at (i) 10 cm, (ii) 5 cm in front of a concave mirror of radius of curvature 15 cm. Find the position, nature, and magnification of the image in each case.

Solution

The focal length $f = -15/2 \text{ cm} = -7.5 \text{ cm}$

(i) The object distance $u = -10 \text{ cm}$. Then Eq. (9.7) gives

$$\frac{1}{v} + \frac{1}{-10} = \frac{1}{-7.5}$$

$$\text{or } v = \frac{10 \times 7.5}{-2.5} = -30 \text{ cm}$$

The image is 30 cm from the mirror on the same side as the object.

$$\text{Also, magnification } m = -\frac{v}{u} = -\frac{(-30)}{(-10)} = -3$$

The image is magnified, real and inverted.

(ii) The object distance $u = -5 \text{ cm}$. Then from Eq. (9.7),

$$\frac{1}{v} + \frac{1}{-5} = \frac{1}{-7.5}$$

$$\text{or } v = \frac{5 \times 7.5}{(7.5 - 5)} = 15 \text{ cm}$$

This image is formed at 15 cm behind the mirror. It is a virtual image.

$$\text{Magnification } m = -\frac{v}{u} = -\frac{15}{(-5)} = 3$$

The image is magnified, virtual and erect.

EXAMPLE 9.3

Example 9.4 Suppose while sitting in a parked car, you notice a jogger approaching towards you in the side view mirror of $R = 2 \text{ m}$. If the jogger is running at a speed of 5 m s^{-1} , how fast the image of the jogger appear to move when the jogger is (a) 39 m, (b) 29 m, (c) 19 m, and (d) 9 m away.

Solution

From the mirror equation, Eq. (9.7), we get

$$v = \frac{fu}{u - f}$$

For convex mirror, since $R = 2 \text{ m}$, $f = 1 \text{ m}$. Then

$$\text{for } u = -39 \text{ m, } v = \frac{(-39) \times 1}{-39 - 1} = \frac{39}{40} \text{ m}$$

Since the jogger moves at a constant speed of 5 m s^{-1} , after 1 s the position of the image v (for $u = -39 + 5 = -34$) is $(34/35) \text{ m}$.

EXAMPLE 9.4

The shift in the position of image in 1 s is

$$\frac{39}{40} - \frac{34}{35} = \frac{1365 - 1360}{1400} = \frac{5}{1400} = \frac{1}{280} \text{ m}$$

Therefore, the average speed of the image when the jogger is between 39 m and 34 m from the mirror, is $(1/280) \text{ m s}^{-1}$

Similarly, it can be seen that for $u = -29 \text{ m}$, -19 m and -9 m , the speed with which the image appears to move is

$$\frac{1}{150} \text{ m s}^{-1}, \frac{1}{60} \text{ m s}^{-1} \text{ and } \frac{1}{10} \text{ m s}^{-1}, \text{ respectively.}$$

Although the jogger has been moving with a constant speed, the speed of his/her image appears to increase substantially as he/she moves closer to the mirror. This phenomenon can be noticed by any person sitting in a stationary car or a bus. In case of moving vehicles, a similar phenomenon could be observed if the vehicle in the rear is moving closer with a constant speed.

9.3 REFRACTION

When a beam of light encounters another transparent medium, a part of light gets reflected back into the first medium while the rest enters the other. A ray of light represents a beam. The direction of propagation of an obliquely incident ($0^\circ < i < 90^\circ$) ray of light that enters the other medium, changes at the interface of the two media. This phenomenon is called *refraction of light*. Snell experimentally obtained the following laws of refraction:

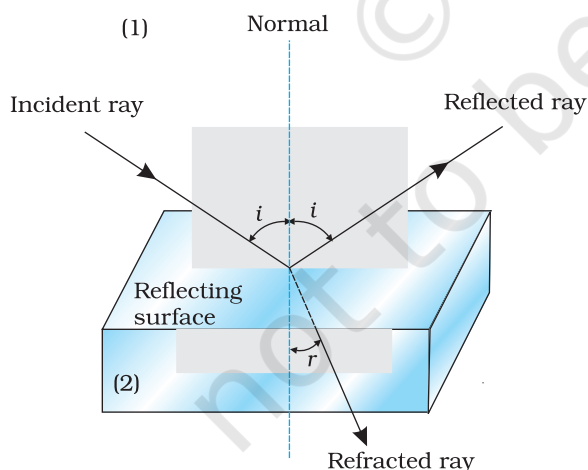


FIGURE 9.8 Refraction and reflection of light.

- (i) The incident ray, the refracted ray and the normal to the interface at the point of incidence, all lie in the same plane.
- (ii) The ratio of the sine of the angle of incidence to the sine of angle of refraction is constant. Remember that the angles of incidence (i) and refraction (r) are the angles that the incident and its refracted ray make with the normal, respectively. We have

$$\frac{\sin i}{\sin r} = n_{21} \quad (9.10)$$

where n_{21} is a constant, called the *refractive index* of the second medium with respect to the first medium. Equation (9.10) is the well-known

Snell's law of refraction. We note that n_{21} is a characteristic of the pair of media (and also depends on the wavelength of light), but is independent of the angle of incidence.

From Eq. (9.10), if $n_{21} > 1$, $r < i$, i.e., the refracted ray bends towards the normal. In such a case medium 2 is said to be *optically denser* (or *denser*, in short) than medium 1. On the other hand, if $n_{21} < 1$, $r > i$, the

refracted ray bends away from the normal. This is the case when incident ray in a denser medium refracts into a rarer medium.

Note: Optical density should not be confused with mass density, which is mass per unit volume. It is possible that mass density of an optically denser medium may be less than that of an optically rarer medium (optical density is the ratio of the speed of light in two media). For example, turpentine and water. Mass density of turpentine is less than that of water but its optical density is higher.

If n_{21} is the refractive index of medium 2 with respect to medium 1 and n_{12} the refractive index of medium 1 with respect to medium 2, then it should be clear that

$$n_{12} = \frac{1}{n_{21}} \quad (9.11)$$

It also follows that if n_{32} is the refractive index of medium 3 with respect to medium 2 then $n_{32} = n_{31} \times n_{12}$, where n_{31} is the refractive index of medium 3 with respect to medium 1.

Some elementary results based on the laws of refraction follow immediately. For a rectangular slab, refraction takes place at two interfaces (air-glass and glass-air). It is easily seen from Fig. 9.9 that $r_2 = i_1$, i.e., the emergent ray is parallel to the incident ray—there is no deviation, but it does suffer lateral displacement/shift with respect to the incident ray. Another familiar observation is that the bottom of a tank filled with water appears to be raised (Fig. 9.10). For viewing near the normal direction, it can be shown that the apparent depth (h_1) is real depth (h_2) divided by the refractive index of the medium (water).

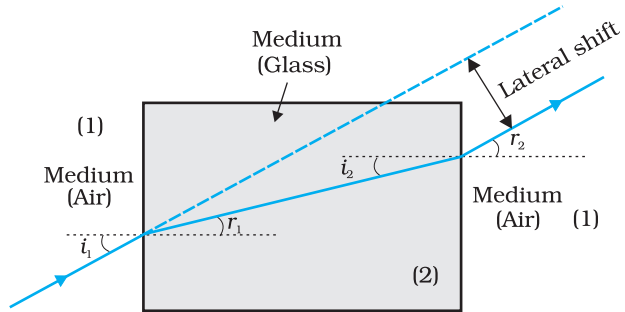


FIGURE 9.9 Lateral shift of a ray refracted through a parallel-sided slab.

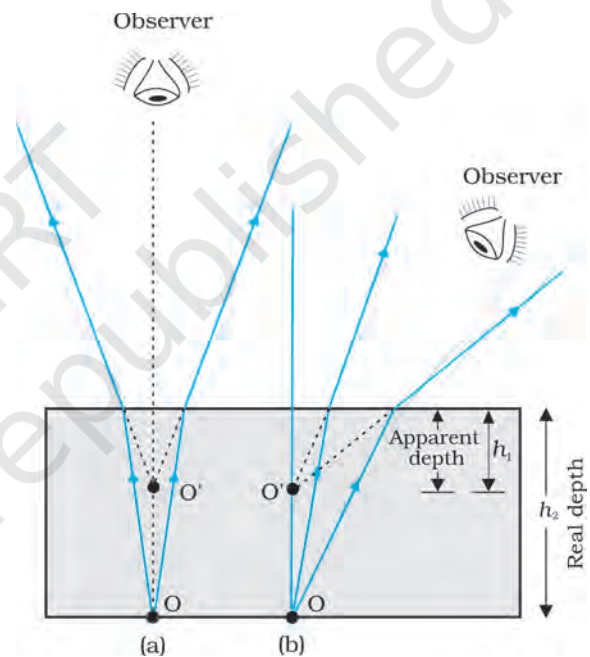


FIGURE 9.10 Apparent depth for (a) normal, and (b) oblique viewing.

9.4 TOTAL INTERNAL REFLECTION

When light travels from an optically denser medium to a rarer medium at the interface, it is partly reflected back into the same medium and partly refracted to the second medium. This reflection is called the *internal reflection*.

When a ray of light enters from a denser medium to a rarer medium, it bends away from the normal, for example, the ray AO_1B in Fig. 9.11. The incident ray AO_1 is partially reflected (O_1C) and partially transmitted (O_1B) or refracted, the angle of refraction (r) being larger than the angle of incidence (i). As the angle of incidence increases, so does the angle of

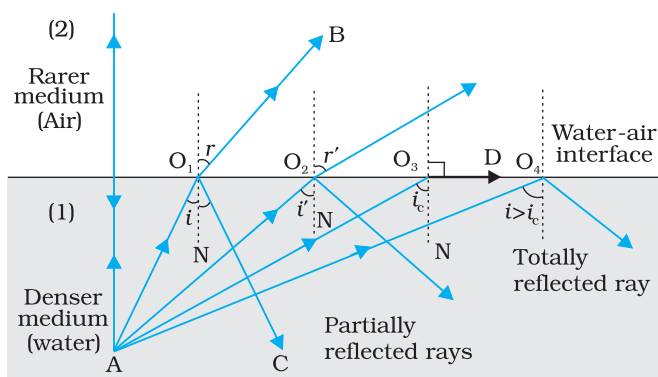


FIGURE 9.11 Refraction and internal reflection of rays from a point A in the denser medium (water) incident at different angles at the interface with a rarer medium (air).

refraction, till for the ray AO₃, the angle of refraction is $\pi/2$. The refracted ray is bent so much away from the normal that it grazes the surface at the interface between the two media. This is shown by the ray AO₃D in Fig. 9.11. If the angle of incidence is increased still further (e.g., the ray AO₄), refraction is not possible, and the incident ray is totally reflected. This is called *total internal reflection*. When light gets reflected by a surface, normally some fraction of it gets transmitted. The reflected ray, therefore, is always less intense than the incident ray, howsoever smooth the reflecting surface may be. In total internal reflection, on the other hand,

no transmission of light takes place.

The angle of incidence corresponding to an angle of refraction 90°, say $\angle AO_3N$, is called the *critical angle* (i_c) for the given pair of media. We see from Snell's law [Eq. (9.10)] that if the relative refractive index of the refracting medium is less than one then, since the maximum value of $\sin r$ is unity, there is an upper limit to the value of $\sin i$ for which the law can be satisfied, that is, $i = i_c$ such that

$$\sin i_c = n_{21} \quad (9.12)$$

For values of i larger than i_c , Snell's law of refraction cannot be satisfied, and hence no refraction is possible.

The refractive index of denser medium 1 with respect to rarer medium 2 will be $n_{12} = 1/\sin i_c$. Some typical critical angles are listed in Table 9.1.

TABLE 9.1 CRITICAL ANGLE OF SOME TRANSPARENT MEDIA WITH RESPECT TO AIR

Substance medium	Refractive index	Critical angle
Water	1.33	48.75
Crown glass	1.52	41.14
Dense flint glass	1.62	37.31
Diamond	2.42	24.41

A demonstration for total internal reflection

All optical phenomena can be demonstrated very easily with the use of a laser torch or pointer, which is easily available nowadays. Take a glass beaker with clear water in it. Add a few drops of milk or any other suspension to water and stir so that water becomes a little turbid. Take a laser pointer and shine its beam through the turbid water. You will find that the path of the beam inside the water shines brightly.

Shine the beam from below the beaker such that it strikes at the upper water surface at the other end. Do you find that it undergoes partial reflection (which is seen as a spot on the table below) and partial refraction [which comes out in the air and is seen as a spot on the roof; Fig. 9.12(a)]? Now direct the laser beam from one side of the beaker such that it strikes the upper surface of water more obliquely [Fig. 9.12(b)]. Adjust the direction of laser beam until you find the angle for which the refraction above the water surface is totally absent and the beam is totally reflected back to water. This is total internal reflection at its simplest.

Pour this water in a long test tube and shine the laser light from top, as shown in Fig. 9.12(c). Adjust the direction of the laser beam such that it is totally internally reflected every time it strikes the walls of the tube. This is similar to what happens in optical fibres.

Take care not to look into the laser beam directly and not to point it at anybody's face.

9.4.1 Total internal reflection in nature and its technological applications

- (i) **Prism:** Prisms designed to bend light by 90° or by 180° make use of total internal reflection [Fig. 9.13(a) and (b)]. Such a prism is also used to invert images without changing their size [Fig. 9.13(c)]. In the first two cases, the critical angle i_c for the material of the prism must be less than 45° . We see from Table 9.1 that this is true for both crown glass and dense flint glass.
- (ii) **Optical fibres:** Nowadays optical fibres are extensively used for transmitting audio and video signals through long distances. Optical fibres too make use of the phenomenon of total internal reflection. Optical fibres are fabricated with high quality composite glass/quartz fibres. Each fibre consists of a core and cladding. The refractive index of the material of the core is higher than that of the cladding.

When a signal in the form of light is directed at one end of the fibre at a suitable angle, it undergoes repeated total internal reflections along the length of the fibre and finally comes out at the other end (Fig. 9.14). Since light undergoes total internal reflection at each stage, there is no appreciable loss in the intensity of the light signal. Optical fibres are fabricated such that light reflected at one side of inner surface strikes the other at an angle larger than the critical angle. Even if the fibre is bent, light can easily travel along its length. Thus, an optical fibre can be used to act as an optical pipe.

A bundle of optical fibres can be put to several uses. Optical fibres are extensively used for transmitting and receiving

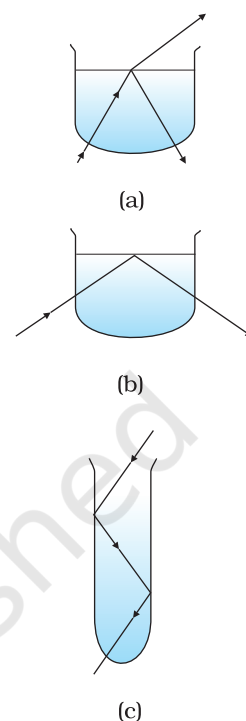


FIGURE 9.12 Observing total internal reflection in water with a laser beam (refraction due to glass of beaker neglected being very thin).

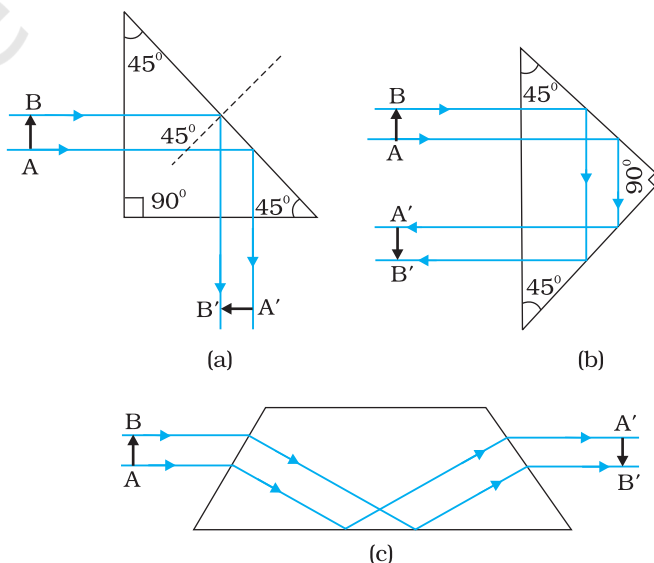


FIGURE 9.13 Prisms designed to bend rays by 90° and 180° or to invert image without changing its size make use of total internal reflection.

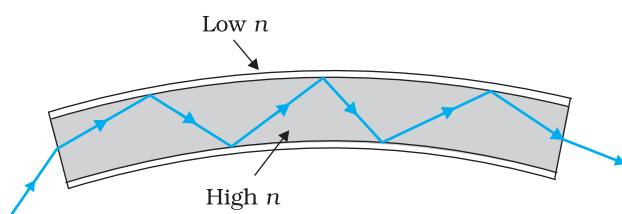


FIGURE 9.14 Light undergoes successive total internal reflections as it moves through an optical fibre.

electrical signals which are converted to light by suitable transducers. Obviously, optical fibres can also be used for transmission of optical signals. For example, these are used as a 'light pipe' to facilitate visual examination of internal organs like esophagus, stomach and intestines. You might have seen a commonly available decorative lamp with fine plastic fibres with their free ends forming a fountain like structure. The other end of the fibres is fixed over an electric lamp. When the

lamp is switched on, the light travels from the bottom of each fibre and appears at the tip of its free end as a dot of light. The fibres in such decorative lamps are optical fibres.

The main requirement in fabricating optical fibres is that there should be very little absorption of light as it travels for long distances inside them. This has been achieved by purification and special preparation of materials such as quartz. In silica glass fibres, it is possible to transmit more than 95% of the light over a fibre length of 1 km. (Compare with what you expect for a block of ordinary window glass 1 km thick.)

9.5 REFRACTION AT SPHERICAL SURFACES AND BY LENSES

We have so far considered refraction at a plane interface. We shall now consider refraction at a spherical interface between two transparent media. An infinitesimal part of a spherical surface can be regarded as planar and the same laws of refraction can be applied at every point on the surface. Just as for reflection by a spherical mirror, the normal at the point of incidence is perpendicular to the tangent plane to the spherical surface at that point and, therefore, passes through its centre of curvature. We first consider refraction by a single spherical surface and follow it by thin lenses. A thin lens is a transparent optical medium bounded by two surfaces; at least one of which should be spherical. Applying the formula for image formation by a single spherical surface successively at the two surfaces of a lens, we shall obtain the lens maker's formula and then the lens formula.

9.5.1 Refraction at a spherical surface

Figure 9.15 shows the geometry of formation of image I of an object O on the principal axis of a spherical surface with centre of curvature C , and radius of curvature R . The rays are incident from a medium of refractive index n_1 , to another of refractive index n_2 . As before, we take the aperture (or the lateral size) of the surface to be small compared to other distances involved, so that small angle approximation can be made. In particular, NM will be taken to be nearly equal to the length of the perpendicular from the point N on the principal axis. We have, for small angles,

$$\tan \angle NOM = \frac{MN}{OM}$$

$$\tan \angle NCM = \frac{MN}{MC}$$

$$\tan \angle NIM = \frac{MN}{MI}$$

Now, for $\triangle NOC$, i is the exterior angle. Therefore, $i = \angle NOM + \angle NCM$

$$i = \frac{MN}{OM} + \frac{MN}{MC} \quad (9.13)$$

Similarly,

$$r = \angle NCM - \angle NIM$$

$$\text{i.e., } r = \frac{MN}{MC} - \frac{MN}{MI} \quad (9.14)$$

Now, by Snell's law

$$n_1 \sin i = n_2 \sin r$$

or for small angles

$$n_1 i = n_2 r$$

Substituting i and r from Eqs. (9.13) and (9.14), we get

$$\frac{n_1}{OM} + \frac{n_2}{MI} = \frac{n_2 - n_1}{MC} \quad (9.15)$$

Here, OM, MI and MC represent magnitudes of distances. Applying the Cartesian sign convention,

$$OM = -u, MI = +v, MC = +R$$

Substituting these in Eq. (9.15), we get

$$\frac{n_2}{v} - \frac{n_1}{u} = \frac{n_2 - n_1}{R} \quad (9.16)$$

Equation (9.16) gives us a relation between object and image distance in terms of refractive index of the medium and the radius of curvature of the curved spherical surface. It holds for any curved spherical surface.

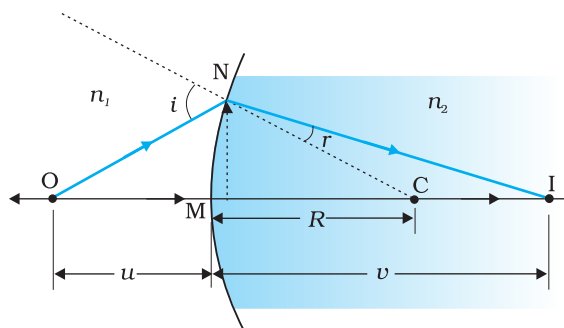


FIGURE 9.15 Refraction at a spherical surface separating two media.

Example 9.5 Light from a point source in air falls on a spherical glass surface ($n = 1.5$ and radius of curvature = 20 cm). The distance of the light source from the glass surface is 100 cm. At what position the image is formed?

Solution

We use the relation given by Eq. (9.16). Here

$u = -100$ cm, $v = ?$, $R = +20$ cm, $n_1 = 1$, and $n_2 = 1.5$.

We then have

$$\frac{1.5}{v} + \frac{1}{100} = \frac{0.5}{20}$$

or $v = +100$ cm

The image is formed at a distance of 100 cm from the glass surface, in the direction of incident light.

9.5.2 Refraction by a lens

Figure 9.16(a) shows the geometry of image formation by a double convex lens. The image formation can be seen in terms of two steps: (i) The first refracting surface forms the image I_1 of the object O [Fig. 9.16(b)]. The image I_1 acts as a virtual object for the second surface that forms the image at I [Fig. 9.16(c)]. Applying Eq. (9.15) to the first interface ABC , we get

$$\frac{n_1}{OB} + \frac{n_2}{BI_1} = \frac{n_2 - n_1}{BC_1} \quad (9.17)$$

A similar procedure applied to the second interface* ADC gives,

$$-\frac{n_2}{DI_1} + \frac{n_1}{DI} = \frac{n_2 - n_1}{DC_2} \quad (9.18)$$

For a thin lens, $BI_1 = DI_1$. Adding Eqs. (9.17) and (9.18), we get

$$\frac{n_1}{OB} + \frac{n_1}{DI} = (n_2 - n_1) \left(\frac{1}{BC_1} + \frac{1}{DC_2} \right) \quad (9.19)$$

Suppose the object is at infinity, i.e., $OB \rightarrow \infty$ and $DI = f$, Eq. (9.19) gives

$$\frac{n_1}{f} = (n_2 - n_1) \left(\frac{1}{BC_1} + \frac{1}{DC_2} \right) \quad (9.20)$$

The point where image of an object placed at infinity is formed is called the *focus* F , of the lens and the distance f gives its *focal length*. A lens has two foci, F and F' , on either side of it (Fig. 9.17). By the sign convention,

$$BC_1 = +R_1,$$

$$DC_2 = -R_2$$

So Eq. (9.20) can be written as

$$\frac{1}{f} = (n_{21} - 1) \left(\frac{1}{R_1} - \frac{1}{R_2} \right) \quad \left[\because n_{21} = \frac{n_2}{n_1} \right] \quad (9.21)$$

Equation (9.21) is known as the *lens maker's formula*. It is useful to design lenses of desired focal length using surfaces of suitable radii of curvature. Note that the formula is true for a concave lens also. In that case R_1 is negative, R_2 positive and therefore, f is negative.

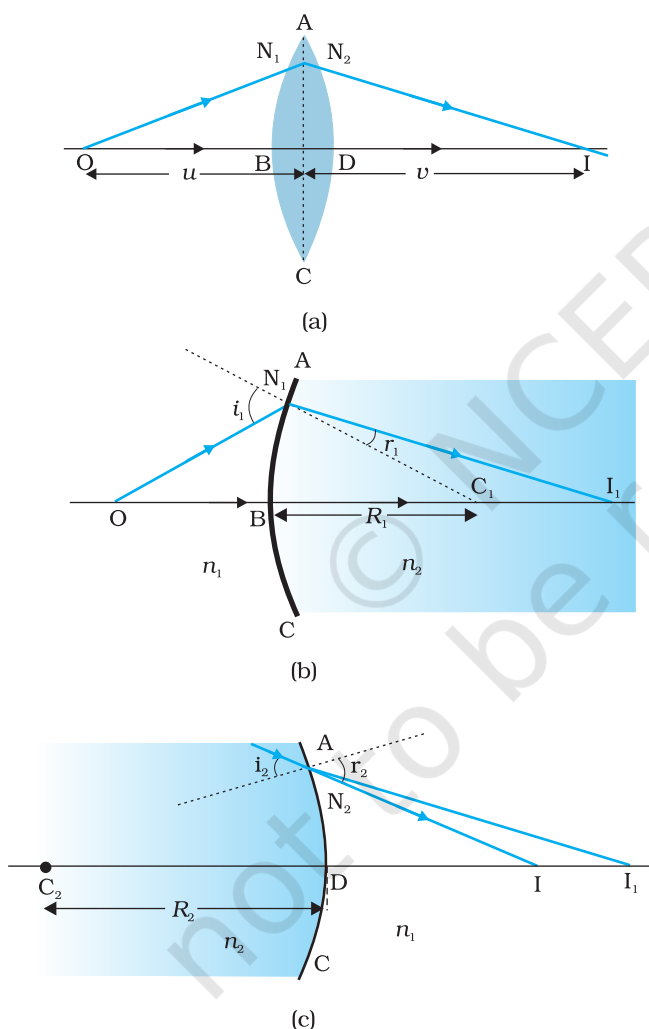


FIGURE 9.16 (a) The position of object, and the image formed by a double convex lens, (b) Refraction at the first spherical surface and (c) Refraction at the second spherical surface.

* Note that now the refractive index of the medium on the right side of ADC is n_1 while on its left it is n_2 . Further DI_1 is negative as the distance is measured against the direction of incident light.

From Eqs. (9.19) and (9.20), we get

$$\frac{n_1}{OB} + \frac{n_1}{DI} = \frac{n_1}{f} \quad (9.22)$$

Again, in the thin lens approximation, B and D are both close to the optical centre of the lens. Applying the sign convention,

BO = -u, DI = +v, we get

$$\frac{1}{v} - \frac{1}{u} = \frac{1}{f} \quad (9.23)$$

Equation (9.23) is the familiar *thin lens formula*. Though we derived it for a real image formed by a convex lens, the formula is valid for both convex as well as concave lenses and for both real and virtual images.

It is worth mentioning that the two foci, F and F', of a double convex or concave lens are equidistant from the optical centre. The focus on the side of the (original) source of light is called the *first focal point*, whereas the other is called the *second focal point*.

To find the image of an object by a lens, we can, in principle, take any two rays emanating from a point on an object; trace their paths using the laws of refraction and find the point where the refracted rays meet (or appear to meet). In practice, however, it is convenient to choose any two of the following rays:

- (i) A ray emanating from the object parallel to the principal axis of the lens after refraction passes through the second principal focus F' (in a convex lens) or appears to diverge (in a concave lens) from the first principal focus F.
- (ii) A ray of light, passing through the optical centre of the lens, emerges without any deviation after refraction.
- (iii) (a) A ray of light passing through the first principal focus of a convex lens [Fig. 9.17(a)] emerges parallel to the principal axis after refraction.
(b) A ray of light incident on a concave lens appearing to meet the principal axis at second focus point emerges parallel to the principal axis after refraction [Fig. 9.17(b)].

Figures 9.17(a) and (b) illustrate these rules for a convex and a concave lens, respectively. You should practice drawing similar ray diagrams for different positions of the object with respect to the lens and also verify that the lens formula, Eq. (9.23), holds good for all cases.

Here again it must be remembered that each point on an object gives out infinite number of rays. All these rays will pass through the same image point after refraction at the lens.

Magnification (*m*) produced by a lens is defined, like that for a mirror, as the ratio of the size of the image to that of the object. Proceeding

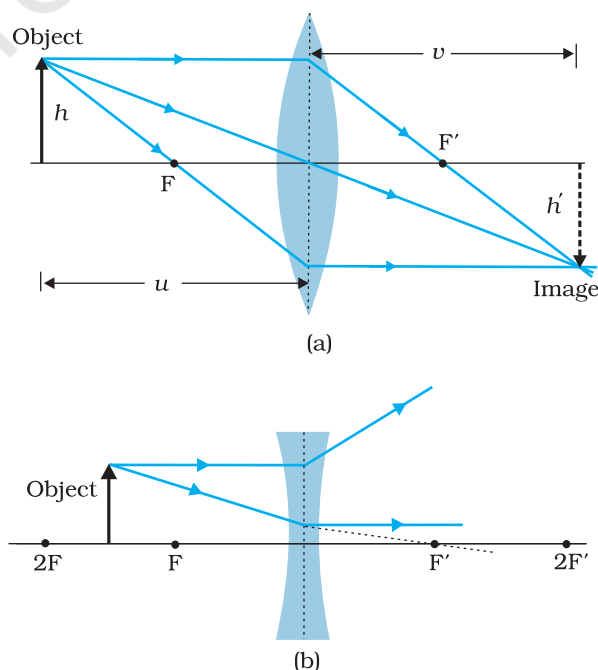


FIGURE 9.17 Tracing rays through (a) convex lens (b) concave lens.

in the same way as for spherical mirrors, it is easily seen that for a lens

$$m = \frac{h'}{h} = \frac{v}{u} \quad (9.24)$$

When we apply the sign convention, we see that, for erect (and virtual) image formed by a convex or concave lens, m is positive, while for an inverted (and real) image, m is negative.

EXAMPLE 9.6

Example 9.6 A magician during a show makes a glass lens with $n = 1.47$ disappear in a trough of liquid. What is the refractive index of the liquid? Could the liquid be water?

Solution

The refractive index of the liquid must be equal to 1.47 in order to make the lens disappear. This means $n_1 = n_2$. This gives $1/f = 0$ or $f \rightarrow \infty$. The lens in the liquid will act like a plane sheet of glass. No, the liquid is not water. It could be glycerine.

9.5.3 Power of a lens

Power of a lens is a measure of the convergence or divergence, which a lens introduces in the light falling on it. Clearly, a lens of shorter focal length bends the incident light more, while converging it in case of a convex lens and diverging it in case of a concave lens. The *power* P of a lens is defined as the tangent of the angle by which it converges or diverges a beam of light parallel to the principal axis falling at unit distance from the optical centre (Fig. 9.18).

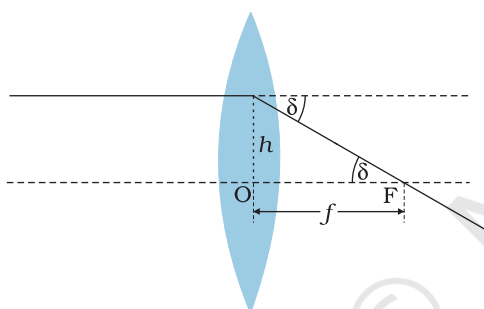


FIGURE 9.18 Power of a lens.

$$\tan \delta = \frac{h}{f}; \text{ if } h = 1, \quad \tan \delta = \frac{1}{f} \quad \text{or} \quad \delta = \frac{1}{f} \quad \text{for small}$$

value of δ . Thus,

$$P = \frac{1}{f} \quad (9.25)$$

The SI unit for power of a lens is dioptre (D): $1D = 1\text{m}^{-1}$. The power of a lens of focal length of 1 metre is one dioptre. Power of a lens is positive for a converging lens and negative for a diverging lens. Thus, when an optician prescribes a corrective lens of power + 2.5 D, the required lens is a convex lens of focal length + 40 cm. A lens of power of - 4.0 D means a concave lens of focal length - 25 cm.

EXAMPLE 9.7

Example 9.7 (i) If $f = 0.5$ m for a glass lens, what is the power of the lens? (ii) The radii of curvature of the faces of a double convex lens are 10 cm and 15 cm. Its focal length is 12 cm. What is the refractive index of glass? (iii) A convex lens has 20 cm focal length in air. What is focal length in water? (Refractive index of air-water = 1.33, refractive index for air-glass = 1.5.)

Solution

(i) Power = +2 dioptre.

(ii) Here, we have $f = +12$ cm, $R_1 = +10$ cm, $R_2 = -15$ cm.

Refractive index of air is taken as unity.

We use the lens formula of Eq. (9.22). The sign convention has to be applied for f , R_1 and R_2 .

Substituting the values, we have

$$\frac{1}{12} = (n - 1) \left(\frac{1}{10} - \frac{1}{-15} \right)$$

This gives $n = 1.5$.

(iii) For a glass lens in air, $n_2 = 1.5$, $n_1 = 1$, $f = +20$ cm. Hence, the lens formula gives

$$\frac{1}{20} = 0.5 \left[\frac{1}{R_1} - \frac{1}{R_2} \right]$$

For the same glass lens in water, $n_2 = 1.5$, $n_1 = 1.33$. Therefore,

$$\frac{1.33}{f} = (1.5 - 1.33) \left[\frac{1}{R_1} - \frac{1}{R_2} \right] \quad (9.26)$$

Combining these two equations, we find $f = +78.2$ cm.

EXAMPLE 9.7

9.5.4 Combination of thin lenses in contact

Consider two lenses A and B of focal length f_1 and f_2 placed in contact with each other. Let the object be placed at a point O beyond the focus of the first lens A (Fig. 9.19). The first lens produces an image at I_1 . Since image I_1 is real, it serves as a virtual object for the second lens B, producing the final image at I. It must, however, be borne in mind that formation of image by the first lens is presumed only to facilitate determination of the position of the final image. In fact, the direction of rays emerging from the first lens gets modified in accordance with the angle at which they strike the second lens. Since the lenses are thin, we assume the optical centres of the lenses to be coincident. Let this central point be denoted by P.

For the image formed by the first lens A, we get

$$\frac{1}{v_1} - \frac{1}{u} = \frac{1}{f_1} \quad (9.27)$$

For the image formed by the second lens B, we get

$$\frac{1}{v} - \frac{1}{v_1} = \frac{1}{f_2} \quad (9.28)$$

Adding Eqs. (9.27) and (9.28), we get

$$\frac{1}{v} - \frac{1}{u} = \frac{1}{f_1} + \frac{1}{f_2} \quad (9.29)$$

If the two lens-system is regarded as equivalent to a single lens of focal length f , we have

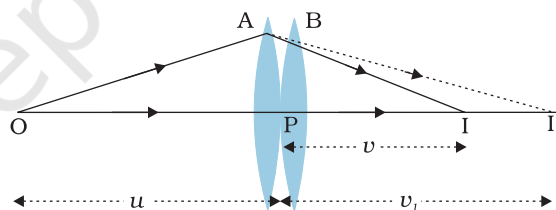


FIGURE 9.19 Image formation by a combination of two thin lenses in contact.

$$\frac{1}{v} - \frac{1}{u} = \frac{1}{f}$$

so that we get

$$\frac{1}{f} = \frac{1}{f_1} + \frac{1}{f_2} \quad (9.30)$$

The derivation is valid for any number of thin lenses in contact. If several thin lenses of focal length $f_1, f_2, f_3, \dots$ are in contact, the effective focal length of their combination is given by

$$\frac{1}{f} = \frac{1}{f_1} + \frac{1}{f_2} + \frac{1}{f_3} + \dots \quad (9.31)$$

In terms of power, Eq. (9.31) can be written as

$$P = P_1 + P_2 + P_3 + \dots \quad (9.32)$$

where P is the net power of the lens combination. Note that the sum in Eq. (9.32) is an algebraic sum of individual powers, so some of the terms on the right side may be positive (for convex lenses) and some negative (for concave lenses). Combination of lenses helps to obtain diverging or converging lenses of desired magnification. It also enhances sharpness of the image. Since the image formed by the first lens becomes the object for the second, Eq. (9.25) implies that the total magnification m of the combination is a product of magnification ($m_1, m_2, m_3, \dots$) of individual lenses

$$m = m_1 m_2 m_3 \dots \quad (9.33)$$

Such a system of combination of lenses is commonly used in designing lenses for cameras, microscopes, telescopes and other optical instruments.

Example 9.8 Find the position of the image formed by the lens combination given in the Fig. 9.20.

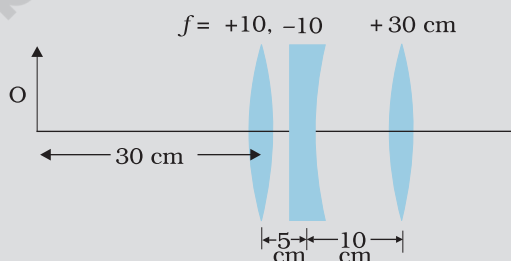


FIGURE 9.20

Solution Image formed by the first lens

$$\frac{1}{v_1} - \frac{1}{u_1} = \frac{1}{f_1}$$

$$\frac{1}{v_1} - \frac{1}{-30} = \frac{1}{10}$$

$$\text{or } v_1 = 15 \text{ cm}$$

The image formed by the first lens serves as the object for the second. This is at a distance of $(15 - 5) \text{ cm} = 10 \text{ cm}$ to the right of the second lens. Though the image is real, it serves as a virtual object for the second lens, which means that the rays appear to come from it for the second lens.

$$\frac{1}{v_2} - \frac{1}{10} = \frac{1}{-10}$$

$$\text{or } v_2 = \infty$$

The virtual image is formed at an infinite distance to the left of the second lens. This acts as an object for the third lens.

$$\frac{1}{v_3} - \frac{1}{u_3} = \frac{1}{f_3}$$

$$\text{or } \frac{1}{v_3} = \frac{1}{\infty} + \frac{1}{30}$$

$$\text{or } v_3 = 30 \text{ cm}$$

The final image is formed 30 cm to the right of the third lens.

EXAMPLE 9.8

9.6 REFRACTION THROUGH A PRISM

Figure 9.21 shows the passage of light through a triangular prism ABC. The angles of incidence and refraction at the first face AB are i and r_1 , while the angle of incidence (from glass to air) at the second face AC is r_2 and the angle of refraction or emergence e . The angle between the emergent ray RS and the direction of the incident ray PQ is called the *angle of deviation*, δ .

In the quadrilateral AQNR, two of the angles (at the vertices Q and R) are right angles. Therefore, the sum of the other angles of the quadrilateral is 180° .

$$\angle A + \angle QNR = 180^\circ$$

From the triangle QNR,

$$r_1 + r_2 + \angle QNR = 180^\circ$$

Comparing these two equations, we get

$$r_1 + r_2 = A \quad (9.34)$$

The total deviation δ is the sum of deviations at the two faces,

$$\delta = (i - r_1) + (e - r_2)$$

that is,

$$\delta = i + e - A \quad (9.35)$$

Thus, the angle of deviation depends on the angle of incidence. A plot between the angle of deviation and angle of incidence is shown in Fig. 9.22. You can see that, in general, any given value of δ , except for $i = e$, corresponds to two values i and hence of e . This, in fact, is expected from the symmetry of i and e in Eq. (9.35), i.e., δ remains the same if i

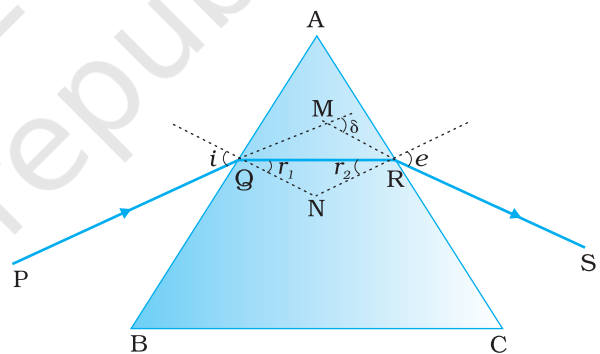


FIGURE 9.21 A ray of light passing through a triangular glass prism.

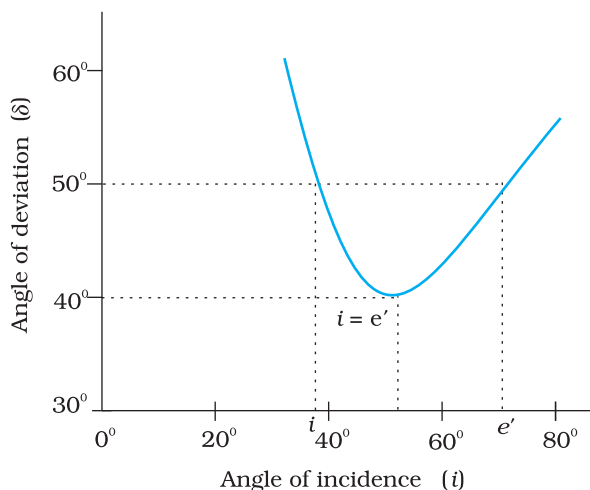


FIGURE 9.22 Plot of angle of deviation (δ) versus angle of incidence (i) for a triangular prism.

and e are interchanged. Physically, this is related to the fact that the path of ray in Fig. 9.21 can be traced back, resulting in the same angle of deviation. At the minimum deviation D_m , the refracted ray inside the prism becomes parallel to its base. We have

$$\delta = D_m, i = e \text{ which implies } r_1 = r_2.$$

Equation (9.34) gives

$$2r = A \text{ or } r = \frac{A}{2} \quad (9.36)$$

In the same way, Eq. (9.35) gives

$$D_m = 2i - A, \text{ or } i = (A + D_m)/2 \quad (9.37)$$

The refractive index of the prism is

$$n_{21} = \frac{n_2}{n_1} = \frac{\sin[(A + D_m)/2]}{\sin[A/2]} \quad (9.38)$$

The angles A and D_m can be measured experimentally. Equation (9.38) thus provides a method of determining refractive index of the material of the prism.

For a small angle prism, i.e., a thin prism, D_m is also very small, and we get

$$n_{21} = \frac{\sin[(A + D_m)/2]}{\sin[A/2]} \approx \frac{(A + D_m)/2}{A/2}$$

$$D_m = (n_{21} - 1)A$$

It implies that, thin prisms do not deviate light much.

9.7 OPTICAL INSTRUMENTS

A number of optical devices and instruments have been designed utilising reflecting and refracting properties of mirrors, lenses and prisms. Periscope, kaleidoscope, binoculars, telescopes, microscopes are some examples of optical devices and instruments that are in common use. Our eye is, of course, one of the most important optical device the nature has endowed us with. We have already studied about the human eye in Class X. We now go on to describe the principles of working of the microscope and the telescope.

9.7.1 The microscope

A simple magnifier or microscope is a converging lens of small focal length (Fig. 9.23). In order to use such a lens as a microscope, the lens is held near the object, one focal length away or less, and the eye is positioned close to the lens on the other side. The idea is to get an erect, magnified and virtual image of the object at a distance so that it can be viewed comfortably, i.e., at 25 cm or more. If the object is at a distance f , the image is at infinity. However, if the object is at a distance slightly less

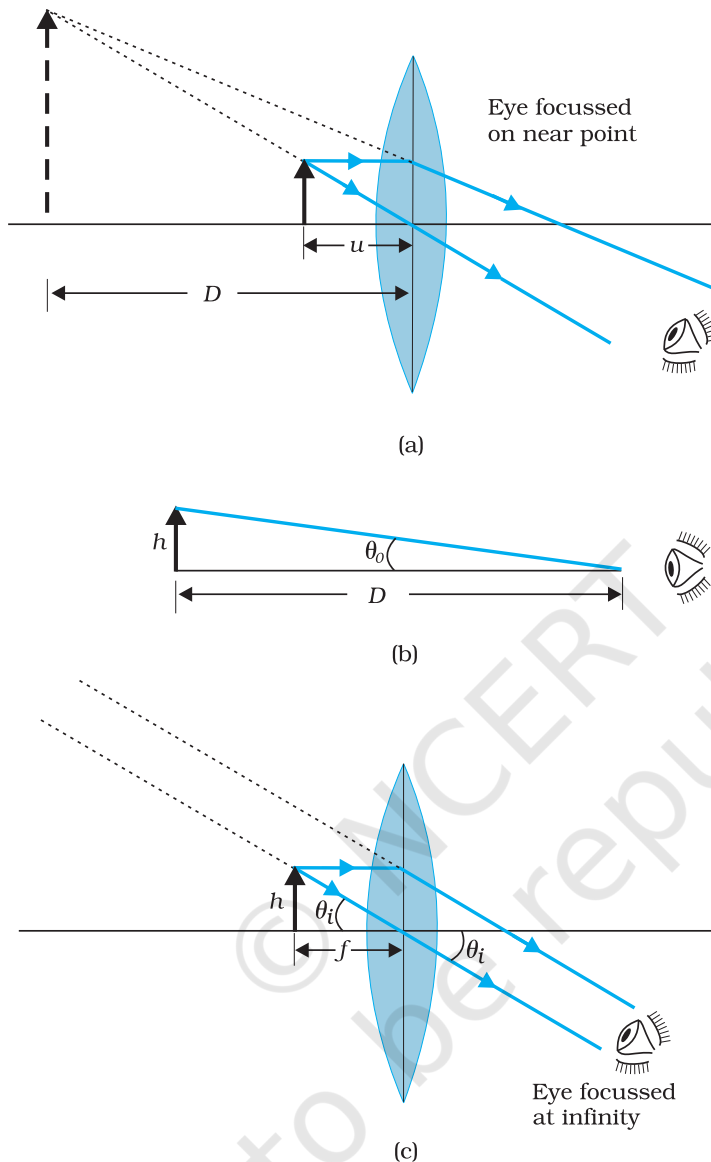


FIGURE 9.23 A simple microscope; (a) the magnifying lens is located such that the image is at the near point, (b) the angle subtended by the object, is the same as that at the near point, and (c) the object near the focal point of the lens; the image is far off but closer than infinity.

than the focal length of the lens, the image is virtual and closer than infinity. Although the closest comfortable distance for viewing the image is when it is at the near point (distance $D \cong 25$ cm), it causes some strain on the eye. Therefore, the image formed at infinity is often considered most suitable for viewing by the relaxed eye. We show both cases, the first in Fig. 9.23(a), and the second in Fig. 9.23(b) and (c).

The linear magnification m , for the image formed at the near point D , by a simple microscope can be obtained by using the relation

$$m = \frac{v}{u} = v \left(\frac{1}{v} - \frac{1}{f} \right) = \left(1 - \frac{v}{f} \right)$$

Now according to our sign convention, v is negative, and is equal in magnitude to D . Thus, the magnification is

$$m = \left(1 + \frac{D}{f} \right) \quad (9.39)$$

Since D is about 25 cm, to have a magnification of six, one needs a convex lens of focal length, $f = 5$ cm.

Note that $m = h'/h$ where h is the size of the object and h' the size of the image. This is also the ratio of the angle subtended by the image to that subtended by the object, if placed at D for comfortable viewing. (Note that this is not the angle actually subtended by the object at the eye, which is h/u .) What a single-lens simple magnifier achieves is that it allows the object to be brought closer to the eye than D .

We will now find the magnification when the image is at infinity. In this case we will have to obtain the *angular* magnification. Suppose the object has a height h . The maximum angle it can subtend, and be clearly visible (without a lens), is when it is at the near point, i.e., a distance D . The angle subtended is then given by

$$\tan \theta_o = \left(\frac{h}{D} \right) \approx \theta_o \quad (9.40)$$

We now find the angle subtended at the eye by the image when the object is at u . From the relations

$$\frac{h'}{h} = m = \frac{v}{u}$$

we have the angle subtended by the image

$$\tan \theta_i = \frac{h'}{-v} = \frac{h}{-v} \cdot \frac{v}{u} = \frac{h}{-u} \approx \theta. \text{ The angle subtended by the object, when it is at } u = -f.$$

$$\theta_i = \left(\frac{h}{f} \right) \quad (9.41)$$

as is clear from Fig. 9.23(c). The angular magnification is, therefore

$$m = \left(\frac{\theta_i}{\theta_o} \right) = \frac{D}{f} \quad (9.42)$$

This is one less than the magnification when the image is at the near point, Eq. (9.39), but the viewing is more comfortable and the difference in magnification is usually small. In subsequent discussions of optical instruments (microscope and telescope) we shall assume the image to be at infinity.

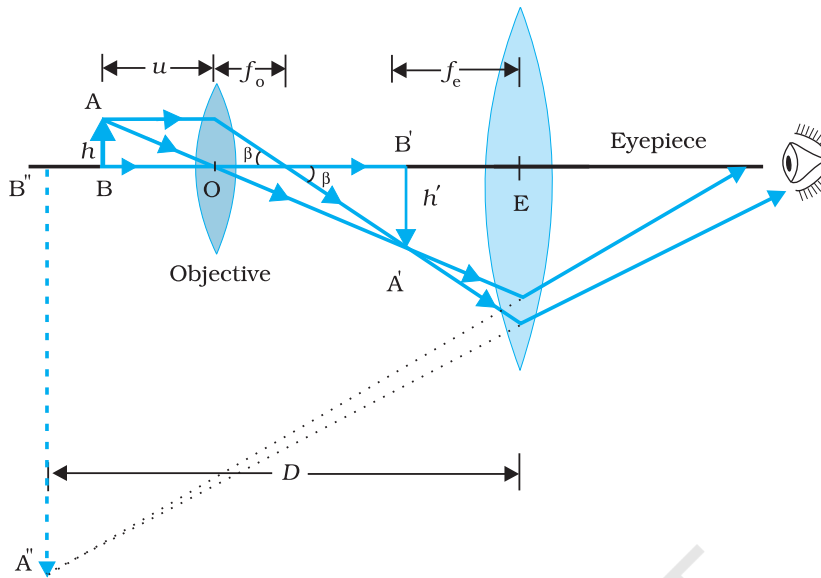


FIGURE 9.24 Ray diagram for the formation of image by a compound microscope.

A simple microscope has a limited maximum magnification (≤ 9) for realistic focal lengths. For much larger magnifications, one uses two lenses, one compounding the effect of the other. This is known as a *compound microscope*. A schematic diagram of a compound microscope is shown in Fig. 9.24. The lens nearest the object, called the *objective*, forms a real, inverted, magnified image of the object. This serves as the object for the second lens, the *eyepiece*, which functions essentially like a simple microscope or magnifier, produces the final image, which is enlarged and virtual. The first inverted image is thus near (at or within) the focal plane of the eyepiece, at a distance appropriate for final image formation at infinity, or a little closer for image formation at the near point. Clearly, the final image is inverted with respect to the original object.

We now obtain the magnification due to a compound microscope. The ray diagram of Fig. 9.24 shows that the (linear) magnification due to the objective, namely h'/h , equals

$$m_o = \frac{h'}{h} = \frac{L}{f_o} \quad (9.43)$$

where we have used the result

$$\tan \beta = \left(\frac{h}{f_o} \right) = \left(\frac{h'}{L} \right)$$

Here h' is the size of the first image, the object size being h and f_o being the focal length of the objective. The first image is formed near the focal point of the eyepiece. The distance L , i.e., the distance between the second focal point of the objective and the first focal point of the eyepiece (focal length f_e) is called the tube length of the compound microscope.

PHYSICS

The world's largest optical telescopes
<http://astro.nineplanets.org/bigeyes.html>

As the first inverted image is near the focal point of the eyepiece, we use the result from the discussion above for the simple microscope to obtain the (angular) magnification m_e due to it [Eq. (9.39)], when the final image is formed at the near point, is

$$m_e = \left(1 + \frac{D}{f_e} \right) \quad [9.44(a)]$$

When the final image is formed at infinity, the angular magnification due to the eyepiece [Eq. (9.42)] is

$$m_e = (D/f_e) \quad [9.44(b)]$$

Thus, the total magnification [(according to Eq. (9.33)], when the image is formed at infinity, is

$$m = m_o m_e = \left(\frac{L}{f_o} \right) \left(\frac{D}{f_e} \right) \quad (9.45)$$

Clearly, to achieve a large magnification of a *small* object (hence the name microscope), the objective and eyepiece should have small focal lengths. In practice, it is difficult to make the focal length much smaller than 1 cm. Also large lenses are required to make L large.

For example, with an objective with $f_o = 1.0$ cm, and an eyepiece with focal length $f_e = 2.0$ cm, and a tube length of 20 cm, the magnification is

$$\begin{aligned} m = m_o m_e &= \left(\frac{L}{f_o} \right) \left(\frac{D}{f_e} \right) \\ &= \frac{20}{1} \times \frac{25}{2} = 250 \end{aligned}$$

Various other factors such as illumination of the object, contribute to the quality and visibility of the image. In modern microscopes, multi-component lenses are used for both the objective and the eyepiece to improve image quality by minimising various optical aberrations (defects) in lenses.

9.7.2 Telescope

The telescope is used to provide angular magnification of distant objects (Fig. 9.25). It also has an objective and an eyepiece. But here, the objective has a large focal length and a much larger aperture than the eyepiece. Light from a distant object enters the objective and a real image is formed in the tube at its second focal point. The eyepiece magnifies this image producing a final inverted image. The magnifying power m is the ratio of the angle β subtended at the eye by the final image to the angle α which the object subtends at the lens or the eye. Hence

$$m \approx \frac{\beta}{\alpha} \approx \frac{h}{f_e} \cdot \frac{f_o}{h} = \frac{f_o}{f_e} \quad (9.46)$$

In this case, the length of the telescope tube is $f_o + f_e$.

Terrestrial telescopes have, in addition, a pair of inverting lenses to make the final image erect. Refracting telescopes can be used both for terrestrial and astronomical observations. For example, consider a telescope whose objective has a focal length of 100 cm and the eyepiece a focal length of 1 cm. The magnifying power of this telescope is $m = 100/1 = 100$.

Let us consider a pair of stars of actual separation $1'$ (one minute of arc). The stars appear as though they are separated by an angle of $100 \times 1' = 100' = 1.67^\circ$.

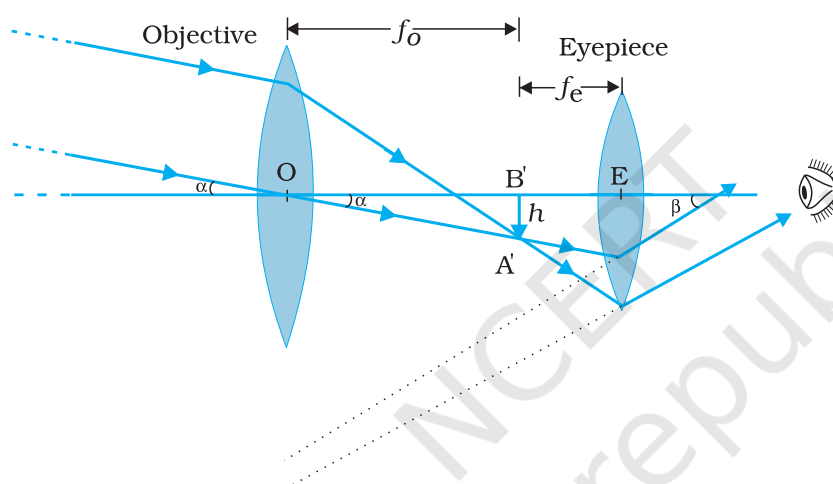


FIGURE 9.25 A refracting telescope.

The main considerations with an astronomical telescope are its light gathering power and its resolution or resolving power. The former clearly depends on the area of the objective. With larger diameters, fainter objects can be observed. The resolving power, or the ability to observe two objects distinctly, which are in very nearly the same direction, also depends on the diameter of the objective. So, the desirable aim in optical telescopes is to make them with objective of large diameter. The largest lens objective in use has a diameter of 40 inch (~ 1.02 m). It is at the Yerkes Observatory in Wisconsin, USA. Such big lenses tend to be very heavy and therefore, difficult to make and support by their edges. Further, it is rather difficult and expensive to make such large sized lenses which form images that are free from any kind of chromatic aberration and distortions.

For these reasons, modern telescopes use a concave mirror rather than a lens for the objective. Telescopes with mirror objectives are called *reflecting* telescopes. There is no chromatic aberration in a mirror. Mechanical support is much less of a problem since a mirror weighs much less than a lens of equivalent optical quality, and can be supported over its entire back surface, not just over its rim. One obvious problem with a reflecting telescope is that the objective mirror focusses light inside

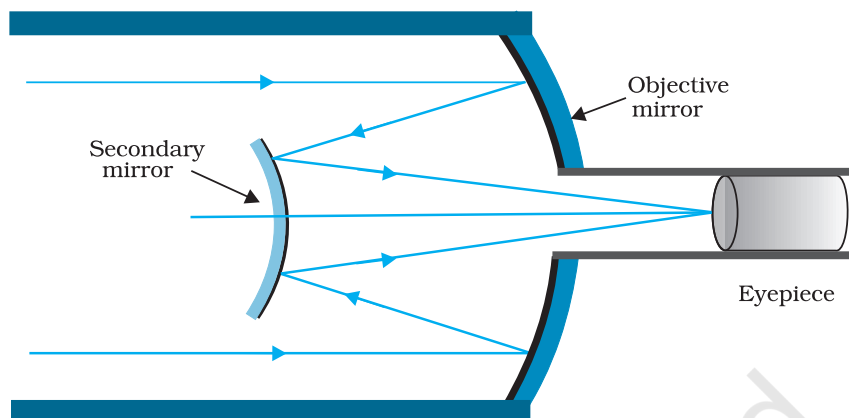


FIGURE 9.26 Schematic diagram of a reflecting telescope (Cassegrain).

the telescope tube. One must have an eyepiece and the observer right there, obstructing some light (depending on the size of the observer cage). This is what is done in the very large 200 inch (~ 5.08 m) diameters, Mt. Palomar telescope, California. The viewer sits near the focal point of the mirror, in a small cage. Another solution to the problem is to deflect the light being focussed by another mirror. One such arrangement using a convex secondary mirror to focus the incident light, which now passes through a hole in the objective primary mirror, is shown in Fig. 9.26. This is known as a *Cassegrain* telescope, after its inventor. It has the advantages of a large focal length in a short telescope. The largest telescope in India is in Kavalur, Tamil Nadu. It is a 2.34 m diameter reflecting telescope (Cassegrain). It was ground, polished, set up, and is being used by the Indian Institute of Astrophysics, Bangalore. The largest reflecting telescopes in the world are the pair of Keck telescopes in Hawaii, USA, with a reflector of 10 metre in diameter.

SUMMARY

1. Reflection is governed by the equation $\angle i = \angle r'$ and refraction by the Snell's law, $\sin i / \sin r = n$, where the incident ray, reflected ray, refracted ray and normal lie in the same plane. Angles of incidence, reflection and refraction are i , r' and r , respectively.
2. The *critical angle of incidence* i_c for a ray incident from a denser to rarer medium, is that angle for which the angle of refraction is 90° . For $i > i_c$, total internal reflection occurs. Multiple internal reflections in diamond ($i_c \cong 24.4^\circ$), totally reflecting prisms and mirage, are some examples of total internal reflection. Optical fibres consist of glass fibres coated with a thin layer of material of *lower* refractive index. Light incident at an angle at one end comes out at the other, after multiple internal reflections, even if the fibre is bent.

3. *Cartesian sign convention:* Distances measured in the same direction as the incident light are positive; those measured in the opposite direction are negative. All distances are measured from the pole/optic centre of the mirror/lens on the principal axis. The heights measured upwards above x -axis and normal to the principal axis of the mirror/lens are taken as positive. The heights measured downwards are taken as negative.

4. *Mirror equation:*

$$\frac{1}{v} + \frac{1}{u} = \frac{1}{f}$$

where u and v are object and image distances, respectively and f is the focal length of the mirror. f is (approximately) half the radius of curvature R . f is negative for concave mirror; f is positive for a convex mirror.

5. For a prism of the angle A , of refractive index n_2 placed in a medium of refractive index n_1 ,

$$n_{21} = \frac{n_2}{n_1} = \frac{\sin[(A + D_m)/2]}{\sin(A/2)}$$

where D_m is the angle of minimum deviation.

6. *For refraction through a spherical interface* (from medium 1 to 2 of refractive index n_1 and n_2 , respectively)

$$\frac{n_2}{v} - \frac{n_1}{u} = \frac{n_2 - n_1}{R}$$

Thin lens formula

$$\frac{1}{v} - \frac{1}{u} = \frac{1}{f}$$

Lens maker's formula

$$\frac{1}{f} = \frac{(n_2 - n_1)}{n_1} \left(\frac{1}{R_1} - \frac{1}{R_2} \right)$$

R_1 and R_2 are the radii of curvature of the lens surfaces. f is positive for a converging lens; f is negative for a diverging lens. The power of a lens $P = 1/f$.

The SI unit for power of a lens is dioptre (D): $1 \text{ D} = 1 \text{ m}^{-1}$.

If several thin lenses of focal length $f_1, f_2, f_3, \dots$ are in contact, the effective focal length of their combination, is given by

$$\frac{1}{f} = \frac{1}{f_1} + \frac{1}{f_2} + \frac{1}{f_3} + \dots$$

The total power of a combination of several lenses is

$$P = P_1 + P_2 + P_3 + \dots$$

7. *Dispersion* is the splitting of light into its constituent colour.

8. *Magnifying power m of a simple microscope* is given by $m = 1 + (D/f)$, where $D = 25$ cm is the least distance of distinct vision and f is the focal length of the convex lens. If the image is at infinity, $m = D/f$. For a compound microscope, the magnifying power is given by $m = m_e \times m_o$ where $m_e = 1 + (D/f_e)$, is the magnification due to the eyepiece and m_o is the magnification produced by the objective. *Approximately,*

$$m = \frac{L}{f_o} \times \frac{D}{f_e}$$

where f_o and f_e are the focal lengths of the objective and eyepiece, respectively, and L is the distance between their focal points.

9. *Magnifying power m of a telescope* is the ratio of the angle β subtended at the eye by the image to the angle α subtended at the eye by the object.

$$m = \frac{\beta}{\alpha} = \frac{f_o}{f_e}$$

where f_o and f_e are the focal lengths of the objective and eyepiece, respectively.

POINTS TO PONDER

1. The laws of reflection and refraction are true for all surfaces and pairs of media at the point of the incidence.
2. The real image of an object placed between f and $2f$ from a convex lens can be seen on a screen placed at the image location. If the screen is removed, is the image still there? This question puzzles many, because it is difficult to reconcile ourselves with an image suspended in air without a screen. But the image does exist. Rays from a given point on the object are converging to an image point in space and diverging away. The screen simply diffuses these rays, some of which reach our eye and we see the image. This can be seen by the images formed in air during a laser show.
3. Image formation needs regular reflection/refraction. In principle, all rays from a given point should reach the same image point. This is why you do not see your image by an irregular reflecting object, say the page of a book.
4. Thick lenses give coloured images due to dispersion. The variety in colour of objects we see around us is due to the constituent colours of the light incident on them. A monochromatic light may produce an entirely different perception about the colours on an object as seen in white light.
5. For a simple microscope, the angular size of the object equals the angular size of the image. Yet it offers magnification because we can keep the small object much closer to the eye than 25 cm and hence have it subtend a large angle. The image is at 25 cm which we can see. Without the microscope, you would need to keep the small object at 25 cm which would subtend a very small angle.

EXERCISES

- 9.1** A small candle, 2.5 cm in size is placed at 27 cm in front of a concave mirror of radius of curvature 36 cm. At what distance from the mirror should a screen be placed in order to obtain a sharp image? Describe the nature and size of the image. If the candle is moved closer to the mirror, how would the screen have to be moved?
- 9.2** A 4.5 cm needle is placed 12 cm away from a convex mirror of focal length 15 cm. Give the location of the image and the magnification. Describe what happens as the needle is moved farther from the mirror.
- 9.3** A tank is filled with water to a height of 12.5 cm. The apparent depth of a needle lying at the bottom of the tank is measured by a microscope to be 9.4 cm. What is the refractive index of water? If water is replaced by a liquid of refractive index 1.63 up to the same height, by what distance would the microscope have to be moved to focus on the needle again?
- 9.4** Figures 9.27(a) and (b) show refraction of a ray in air incident at 60° with the normal to a glass-air and water-air interface, respectively. Predict the angle of refraction in glass when the angle of incidence in water is 45° with the normal to a water-glass interface [Fig. 9.27(c)].

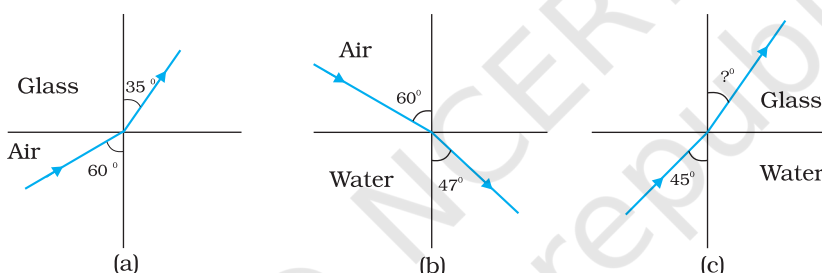


FIGURE 9.27

- 9.5** A small bulb is placed at the bottom of a tank containing water to a depth of 80 cm. What is the area of the surface of water through which light from the bulb can emerge out? Refractive index of water is 1.33. (Consider the bulb to be a point source.)
- 9.6** A prism is made of glass of unknown refractive index. A parallel beam of light is incident on a face of the prism. The angle of minimum deviation is measured to be 40° . What is the refractive index of the material of the prism? The refracting angle of the prism is 60° . If the prism is placed in water (refractive index 1.33), predict the new angle of minimum deviation of a parallel beam of light.
- 9.7** Double-convex lenses are to be manufactured from a glass of refractive index 1.55, with both faces of the same radius of curvature. What is the radius of curvature required if the focal length is to be 20 cm?
- 9.8** A beam of light converges at a point P. Now a lens is placed in the path of the convergent beam 12 cm from P. At what point does the beam converge if the lens is (a) a convex lens of focal length 20 cm, and (b) a concave lens of focal length 16 cm?
- 9.9** An object of size 3.0 cm is placed 14 cm in front of a concave lens of focal length 21 cm. Describe the image produced by the lens. What happens if the object is moved further away from the lens?

- 9.10** What is the focal length of a convex lens of focal length 30cm in contact with a concave lens of focal length 20cm? Is the system a converging or a diverging lens? Ignore thickness of the lenses.
- 9.11** A compound microscope consists of an objective lens of focal length 2.0cm and an eyepiece of focal length 6.25cm separated by a distance of 15cm. How far from the objective should an object be placed in order to obtain the final image at (a) the least distance of distinct vision (25cm), and (b) at infinity? What is the magnifying power of the microscope in each case?
- 9.12** A person with a normal near point (25cm) using a compound microscope with objective of focal length 8.0 mm and an eyepiece of focal length 2.5cm can bring an object placed at 9.0mm from the objective in sharp focus. What is the separation between the two lenses? Calculate the magnifying power of the microscope.
- 9.13** A small telescope has an objective lens of focal length 144cm and an eyepiece of focal length 6.0cm. What is the magnifying power of the telescope? What is the separation between the objective and the eyepiece?
- 9.14** (a) A giant refracting telescope at an observatory has an objective lens of focal length 15m. If an eyepiece of focal length 1.0cm is used, what is the angular magnification of the telescope?
(b) If this telescope is used to view the moon, what is the diameter of the image of the moon formed by the objective lens? The diameter of the moon is 3.48×10^6 m, and the radius of lunar orbit is 3.8×10^8 m.
- 9.15** Use the mirror equation to deduce that:
(a) an object placed between f and $2f$ of a concave mirror produces a real image beyond $2f$.
(b) a convex mirror always produces a virtual image independent of the location of the object.
(c) the virtual image produced by a convex mirror is always diminished in size and is located between the focus and the pole.
(d) an object placed between the pole and focus of a concave mirror produces a virtual and enlarged image.
[Note: This exercise helps you deduce algebraically properties of images that one obtains from explicit ray diagrams.]
- 9.16** A small pin fixed on a table top is viewed from above from a distance of 50cm. By what distance would the pin appear to be raised if it is viewed from the same point through a 15cm thick glass slab held parallel to the table? Refractive index of glass = 1.5. Does the answer depend on the location of the slab?
- 9.17** (a) Figure 9.28 shows a cross-section of a 'light pipe' made of a glass fibre of refractive index 1.68. The outer covering of the pipe is made of a material of refractive index 1.44. What is the range of the angles of the incident rays with the axis of the pipe for which total reflections inside the pipe take place, as shown in the figure.

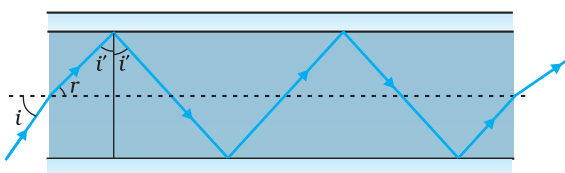


FIGURE 9.28

- (b) What is the answer if there is no outer covering of the pipe?
- 9.18** The image of a small electric bulb fixed on the wall of a room is to be obtained on the opposite wall 3m away by means of a large convex lens. What is the maximum possible focal length of the lens required for the purpose?
- 9.19** A screen is placed 90cm from an object. The image of the object on the screen is formed by a convex lens at two different locations separated by 20cm. Determine the focal length of the lens.
- 9.20** (a) Determine the 'effective focal length' of the combination of the two lenses in Exercise 9.10, if they are placed 8.0cm apart with their principal axes coincident. Does the answer depend on which side of the combination a beam of parallel light is incident? Is the notion of effective focal length of this system useful at all?
- (b) An object 1.5 cm in size is placed on the side of the convex lens in the arrangement (a) above. The distance between the object and the convex lens is 40cm. Determine the magnification produced by the two-lens system, and the size of the image.
- 9.21** At what angle should a ray of light be incident on the face of a prism of refracting angle 60° so that it just suffers total internal reflection at the other face? The refractive index of the material of the prism is 1.524.
- 9.22** A card sheet divided into squares each of size 1 mm^2 is being viewed at a distance of 9 cm through a magnifying glass (a converging lens of focal length 9 cm) held close to the eye.
- (a) What is the magnification produced by the lens? How much is the area of each square in the virtual image?
- (b) What is the angular magnification (magnifying power) of the lens?
- (c) Is the magnification in (a) equal to the magnifying power in (b)? Explain.
- 9.23** (a) At what distance should the lens be held from the card sheet in Exercise 9.22 in order to view the squares distinctly with the maximum possible magnifying power?
- (b) What is the magnification in this case?
- (c) Is the magnification equal to the magnifying power in this case? Explain.
- 9.24** What should be the distance between the object in Exercise 9.23 and the magnifying glass if the virtual image of each square in the figure is to have an area of 6.25 mm^2 . Would you be able to see the squares distinctly with your eyes very close to the magnifier?

[Note: Exercises 9.22 to 9.24 will help you clearly understand the difference between magnification in absolute size and the angular magnification (or magnifying power) of an instrument.]

9.25 Answer the following questions:

- (a) The angle subtended at the eye by an object is equal to the angle subtended at the eye by the virtual image produced by a magnifying glass. In what sense then does a magnifying glass provide angular magnification?
- (b) In viewing through a magnifying glass, one usually positions one's eyes very close to the lens. Does angular magnification change if the eye is moved back?
- (c) Magnifying power of a simple microscope is inversely proportional to the focal length of the lens. What then stops us from using a convex lens of smaller and smaller focal length and achieving greater and greater magnifying power?
- (d) Why must both the objective and the eyepiece of a compound microscope have short focal lengths?
- (e) When viewing through a compound microscope, our eyes should be positioned not on the eyepiece but a short distance away from it for best viewing. Why? How much should be that short distance between the eye and eyepiece?

9.26 An angular magnification (magnifying power) of 30X is desired using an objective of focal length 1.25 cm and an eyepiece of focal length 5 cm. How will you set up the compound microscope?

9.27 A small telescope has an objective lens of focal length 140 cm and an eyepiece of focal length 5.0 cm. What is the magnifying power of the telescope for viewing distant objects when

- (a) the telescope is in normal adjustment (i.e., when the final image is at infinity)?
- (b) the final image is formed at the least distance of distinct vision (25 cm)?

9.28 (a) For the telescope described in Exercise 9.27 (a), what is the separation between the objective lens and the eyepiece?

(b) If this telescope is used to view a 100 m tall tower 3 km away, what is the height of the image of the tower formed by the objective lens?

(c) What is the height of the final image of the tower if it is formed at 25 cm?

9.29 A Cassegrain telescope uses two mirrors as shown in Fig. 9.26. Such a telescope is built with the mirrors 20 mm apart. If the radius of curvature of the large mirror is 220 mm and the small mirror is 140 mm, where will the final image of an object at infinity be?

9.30 Light incident normally on a plane mirror attached to a galvanometer coil retraces backwards as shown in Fig. 9.29. A current in the coil produces a deflection of 3.5° of the mirror. What is the displacement of the reflected spot of light on a screen placed 1.5 m away?

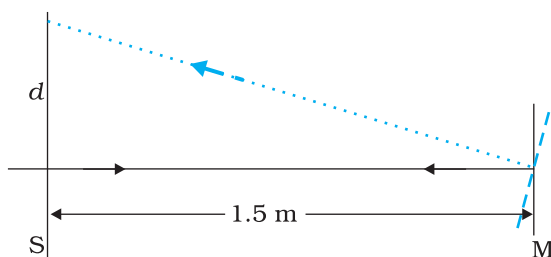


FIGURE 9.29

- 9.31** Figure 9.30 shows an equiconvex lens (of refractive index 1.50) in contact with a liquid layer on top of a plane mirror. A small needle with its tip on the principal axis is moved along the axis until its inverted image is found at the position of the needle. The distance of the needle from the lens is measured to be 45.0 cm. The liquid is removed and the experiment is repeated. The new distance is measured to be 30.0 cm. What is the refractive index of the liquid?

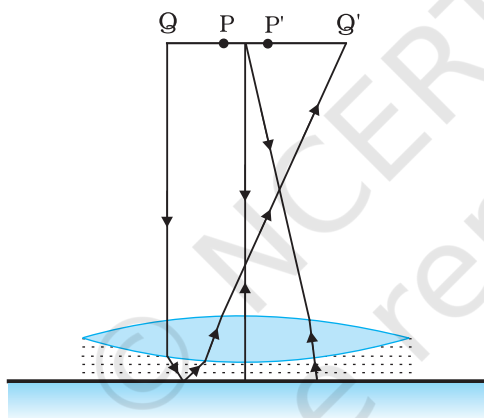
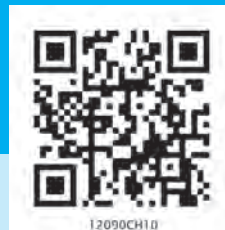


FIGURE 9.30

Notes

© NCERT
not to be republished



Chapter Ten

WAVE OPTICS

10.1 INTRODUCTION

In 1637 Descartes gave the corpuscular model of light and derived Snell's law. It explained the laws of reflection and refraction of light at an interface. The corpuscular model predicted that if the ray of light (on refraction) bends towards the normal then the speed of light would be greater in the second medium. This corpuscular model of light was further developed by Isaac Newton in his famous book entitled *OPTICKS* and because of the tremendous popularity of this book, the corpuscular model is very often attributed to Newton.

In 1678, the Dutch physicist Christiaan Huygens put forward the wave theory of light – it is this wave model of light that we will discuss in this chapter. As we will see, the wave model could satisfactorily explain the phenomena of reflection and refraction; however, it predicted that on refraction if the wave bends towards the normal then the speed of light would be less in the second medium. This is in contradiction to the prediction made by using the corpuscular model of light. It was much later confirmed by experiments where it was shown that the speed of light in water is less than the speed in air confirming the prediction of the wave model; Foucault carried out this experiment in 1850.

The wave theory was not readily accepted primarily because of Newton's authority and also because light could travel through vacuum

and it was felt that a wave would always require a medium to propagate from one point to the other. However, when Thomas Young performed his famous interference experiment in 1801, it was firmly established that light is indeed a wave phenomenon. The wavelength of visible light was measured and found to be extremely small; for example, the wavelength of yellow light is about $0.6 \mu\text{m}$. Because of the smallness of the wavelength of visible light (in comparison to the dimensions of typical mirrors and lenses), light can be assumed to approximately travel in straight lines. This is the field of geometrical optics, which we had discussed in the previous chapter. Indeed, the branch of optics in which one completely neglects the finiteness of the wavelength is called geometrical optics and a ray is defined as the path of energy propagation in the limit of wavelength tending to zero.

After the interference experiment of Young in 1801, for the next 40 years or so, many experiments were carried out involving the interference and diffraction of lightwaves; these experiments could only be satisfactorily explained by assuming a wave model of light. Thus, around the middle of the nineteenth century, the wave theory seemed to be very well established. The only major difficulty was that since it was thought that a wave required a medium for its propagation, how could light waves propagate through vacuum. This was explained when Maxwell put forward his famous electromagnetic theory of light. Maxwell had developed a set of equations describing the laws of electricity and magnetism and using these equations he derived what is known as the wave equation from which he *predicted* the existence of electromagnetic waves*. From the wave equation, Maxwell could calculate the speed of electromagnetic waves in free space and he found that the theoretical value was very close to the measured value of speed of light. From this, he propounded that *light must be an electromagnetic wave*. Thus, according to Maxwell, light waves are associated with changing electric and magnetic fields; changing electric field produces a time and space varying magnetic field and a changing magnetic field produces a time and space varying electric field. The changing electric and magnetic fields result in the propagation of electromagnetic waves (or light waves) even in vacuum.

In this chapter we will first discuss the original formulation of the *Huygens principle* and derive the laws of reflection and refraction. In Sections 10.4 and 10.5, we will discuss the phenomenon of interference which is based on the principle of superposition. In Section 10.6 we will discuss the phenomenon of diffraction which is based on Huygens-Fresnel principle. Finally in Section 10.7 we will discuss the phenomenon of polarisation which is based on the fact that the light waves are *transverse electromagnetic waves*.

* Maxwell had predicted the existence of electromagnetic waves around 1855; it was much later (around 1890) that Heinrich Hertz produced radiowaves in the laboratory. J.C. Bose and G. Marconi made practical applications of the *Hertzian waves*

10.2 HUYGENS PRINCIPLE

We would first define a wavefront: when we drop a small stone on a calm pool of water, waves spread out from the point of impact. Every point on the surface starts oscillating with time. At any instant, a photograph of the surface would show circular rings on which the disturbance is maximum. Clearly, all points on such a circle are oscillating in phase because they are at the same distance from the source. Such a locus of points, which oscillate in phase is called a *wavefront*; thus *a wavefront is defined as a surface of constant phase*. The speed with which the wavefront moves outwards from the source is called the speed of the wave. The energy of the wave travels in a direction perpendicular to the wavefront.

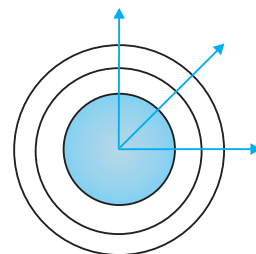


FIGURE 10.1 (a) A diverging spherical wave emanating from a point source. The wavefronts are spherical.

If we have a point source emitting waves uniformly in all directions, then the locus of points which have the same amplitude and vibrate in the same phase are spheres and we have what is known as a *spherical wave* as shown in Fig. 10.1(a). At a large distance from the source, a small portion of the sphere can be considered as a plane and we have what is known as a *plane wave* [Fig. 10.1(b)].

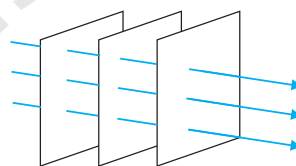


FIGURE 10.1 (b) At a large distance from the source, a small portion of the spherical wave can be approximated by a plane wave.

Now, if we know the shape of the wavefront at $t = 0$, then Huygens principle allows us to determine the shape of the wavefront at a later time τ . Thus, Huygens principle is essentially a geometrical construction, which given the shape of the wavefront at any time allows us to determine the shape of the wavefront at a later time. Let us consider a diverging wave and let F_1F_2 represent a portion of the spherical wavefront at $t = 0$ (Fig. 10.2). Now, according to Huygens principle, *each point of the wavefront is the source of a secondary disturbance and the wavelets emanating from these points spread out in all directions with the speed of the wave. These wavelets emanating from the wavefront are usually referred to as secondary wavelets and if we draw a common tangent to all these spheres, we obtain the new position of the wavefront at a later time.*

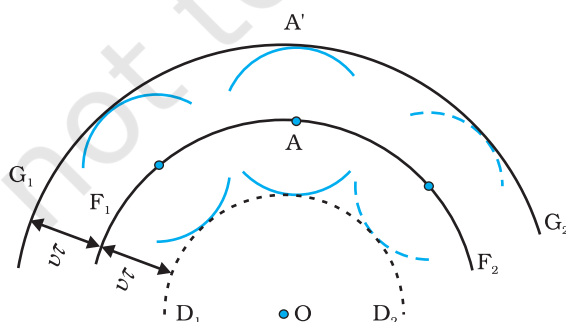


FIGURE 10.2 F_1F_2 represents the spherical wavefront (with O as centre) at $t = 0$. The envelope of the secondary wavelets emanating from F_1F_2 produces the forward moving wavefront G_1G_2 . The backwave D_1D_2 does not exist.

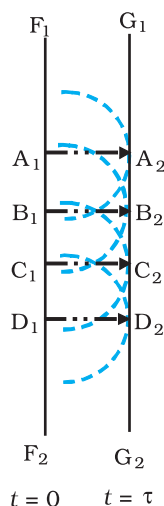


FIGURE 10.3

Huygens geometrical construction for a plane wave propagating to the right. $F_1 F_2$ is the plane wavefront at $t = 0$ and $G_1 G_2$ is the wavefront at a later time τ . The lines $A_1 A_2$, $B_1 B_2$... etc., are normal to both $F_1 F_2$ and $G_1 G_2$ and represent rays.

Thus, if we wish to determine the shape of the wavefront at $t = \tau$, we draw spheres of radius $v\tau$ from each point on the spherical wavefront where v represents the speed of the waves in the medium. If we now draw a common tangent to all these spheres, we obtain the new position of the wavefront at $t = \tau$. The new wavefront shown as $G_1 G_2$ in Fig. 10.2 is again spherical with point O as the centre.

The above model has one shortcoming: we also have a backwave which is shown as $D_1 D_2$ in Fig. 10.2. Huygens argued that the amplitude of the secondary wavelets is maximum in the forward direction and zero in the backward direction; by making this adhoc assumption, Huygens could explain the absence of the backwave. However, this adhoc assumption is not satisfactory and the absence of the backwave is really justified from more rigorous wave theory.

In a similar manner, we can use Huygens principle to determine the shape of the wavefront for a plane wave propagating through a medium (Fig. 10.3).

10.3 REFRACTION AND REFLECTION OF PLANE WAVES USING HUYGENS PRINCIPLE

10.3.1 Refraction of a plane wave

We will now use Huygens principle to derive the laws of refraction. Let PP' represent the surface separating medium 1 and medium 2, as shown in Fig. 10.4. Let v_1 and v_2 represent the speed of light in medium 1 and medium 2, respectively. We assume a plane wavefront AB propagating in the direction $A'A$ incident on the interface at an angle i as shown in the figure. Let τ be the time taken by the wavefront to travel the distance BC . Thus,

$$BC = v_1 \tau$$

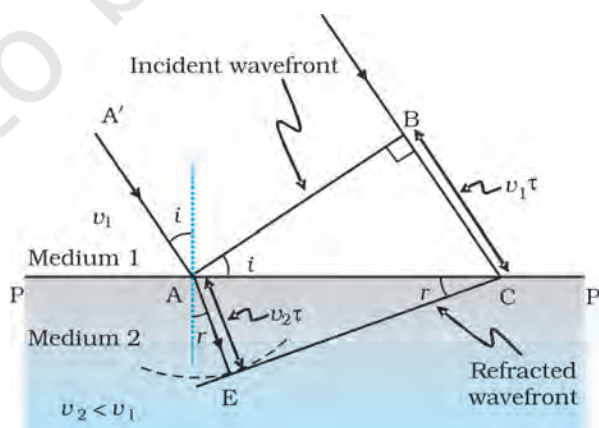


FIGURE 10.4 A plane wave AB is incident at an angle i on the surface PP' separating medium 1 and medium 2. The plane wave undergoes refraction and CE represents the refracted wavefront. The figure corresponds to $v_2 < v_1$ so that the refracted waves bends towards the normal.

In order to determine the shape of the refracted wavefront, we draw a sphere of radius $v_2\tau$ from the point A in the second medium (the speed of the wave in the second medium is v_2). Let CE represent a tangent plane drawn from the point C on to the sphere. Then, $AE = v_2\tau$ and CE would represent the refracted wavefront. If we now consider the triangles ABC and AEC, we readily obtain

$$\sin i = \frac{BC}{AC} = \frac{v_1\tau}{AC} \quad (10.1)$$

and

$$\sin r = \frac{AE}{AC} = \frac{v_2\tau}{AC} \quad (10.2)$$

where i and r are the angles of incidence and refraction, respectively. Thus we obtain

$$\frac{\sin i}{\sin r} = \frac{v_1}{v_2} \quad (10.3)$$

From the above equation, we get the important result that if $r < i$ (i.e., if the ray bends toward the normal), the speed of the light wave in the second medium (v_2) will be less than the speed of the light wave in the first medium (v_1). This prediction is opposite to the prediction from the corpuscular model of light and as later experiments showed, the prediction of the wave theory is correct. Now, if c represents the speed of light in vacuum, then,

$$n_1 = \frac{c}{v_1} \quad (10.4)$$

and

$$n_2 = \frac{c}{v_2} \quad (10.5)$$

are known as the refractive indices of medium 1 and medium 2, respectively. In terms of the refractive indices, Eq. (10.3) can be written as

$$n_1 \sin i = n_2 \sin r \quad (10.6)$$

This is the *Snell's law of refraction*. Further, if λ_1 and λ_2 denote the wavelengths of light in medium 1 and medium 2, respectively and if the distance BC is equal to λ_1 then the distance AE will be equal to λ_2 (because if the crest from B has reached C in time τ , then the crest from A should have also reached E in time τ); thus,

$$\frac{\lambda_1}{\lambda_2} = \frac{BC}{AE} = \frac{v_1}{v_2}$$

or

$$\frac{v_1}{\lambda_1} = \frac{v_2}{\lambda_2} \quad (10.7)$$



Christiaan Huygens (1629 – 1695) Dutch physicist, astronomer, mathematician and the founder of the wave theory of light. His book, *Treatise on light*, makes fascinating reading even today. He brilliantly explained the double refraction shown by the mineral calcite in this work in addition to reflection and refraction. He was the first to analyse circular and simple harmonic motion and designed and built improved clocks and telescopes. He discovered the true geometry of Saturn's rings.

CHRISTIAAN HUYGENS (1629 – 1695)

The above equation implies that when a wave gets refracted into a denser medium ($v_1 > v_2$) the wavelength and the speed of propagation decrease but the frequency $\nu (= v/\lambda)$ remains the same.

10.3.2 Refraction at a rarer medium

We now consider refraction of a plane wave at a rarer medium, i.e., $v_2 > v_1$. Proceeding in an exactly similar manner we can construct a refracted wavefront as shown in Fig. 10.5. The angle of refraction will now be greater than angle of incidence; however, we will still have $n_1 \sin i = n_2 \sin r$. We define an angle i_c by the following equation

$$\sin i_c = \frac{n_2}{n_1} \quad (10.8)$$

Thus, if $i = i_c$ then $\sin r = 1$ and $r = 90^\circ$. Obviously, for $i > i_c$, there can not be any refracted wave. The angle i_c is known as the *critical angle* and for all angles of incidence greater than the critical angle, we will not have any refracted wave and the wave will undergo what is known as *total internal reflection*. The phenomenon of total internal reflection and its applications was discussed in Section 9.4.

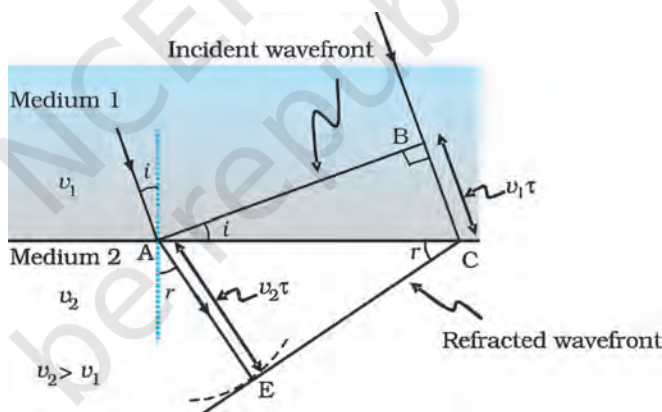


FIGURE 10.5 Refraction of a plane wave incident on a rarer medium for which $v_2 > v_1$. The plane wave bends away from the normal.

10.3.3 Reflection of a plane wave by a plane surface

We next consider a plane wave AB incident at an angle i on a reflecting surface MN. If v represents the speed of the wave in the medium and if τ represents the time taken by the wavefront to advance from the point B to C then the distance

$$BC = v\tau$$

In order to construct the reflected wavefront we draw a sphere of radius $v\tau$ from the point A as shown in Fig. 10.6. Let CE represent the tangent plane drawn from the point C to this sphere. Obviously

$$AE = BC = v\tau$$

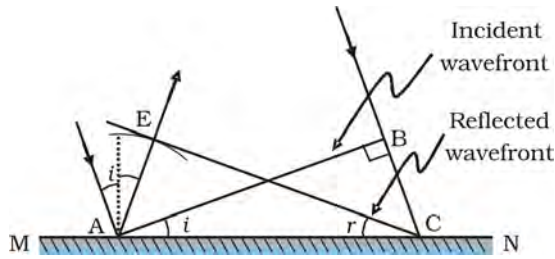


FIGURE 10.6 Reflection of a plane wave AB by the reflecting surface MN. AB and CE represent incident and reflected wavefronts.

If we now consider the triangles EAC and BAC we will find that they are congruent and therefore, the angles i and r (as shown in Fig. 10.6) would be equal. This is the *law of reflection*.

Once we have the laws of reflection and refraction, the behaviour of prisms, lenses, and mirrors can be understood. These phenomena were discussed in detail in Chapter 9 on the basis of rectilinear propagation of light. Here we just describe the behaviour of the wavefronts as they undergo reflection or refraction. In Fig. 10.7(a) we consider a plane wave passing through a thin prism. Clearly, since the speed of light waves is less in glass, the lower portion of the incoming wavefront (which travels through the greatest thickness of glass) will get delayed resulting in a tilt in the emerging wavefront as shown in the figure. In Fig. 10.7(b) we consider a plane wave incident on a thin convex lens; the central part of the incident plane wave traverses the thickest portion of the lens and is delayed the most. The emerging wavefront has a depression at the centre and therefore the wavefront becomes spherical and converges to the point F which is known as the *focus*. In Fig. 10.7(c) a plane wave is incident on a concave mirror and on reflection we have a spherical wave converging to the focal point F. In a similar manner, we can understand refraction and reflection by concave lenses and convex mirrors.

From the above discussion it follows that the total time taken from a point on the object to the corresponding point on the image is the same measured along any ray. For example, when a convex lens focusses light to form a real image, although the ray going through the centre traverses a shorter path, but because of the slower speed in glass, the time taken is the same as for rays travelling near the edge of the lens.

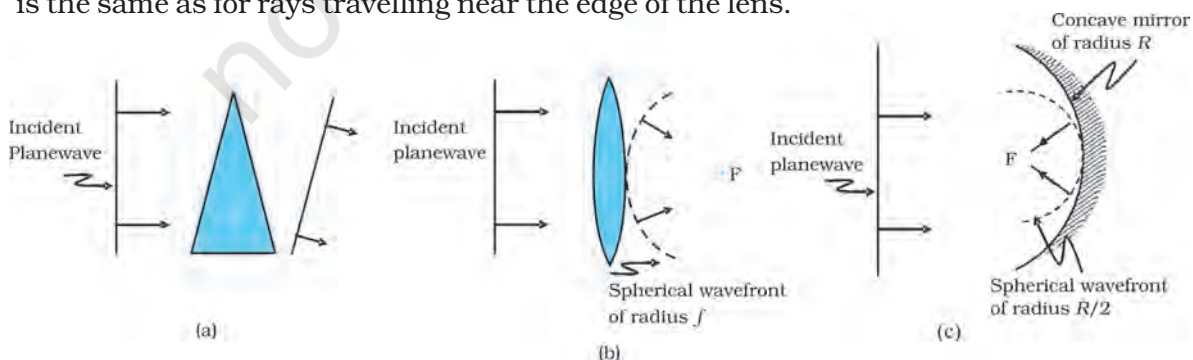


FIGURE 10.7 Refraction of a plane wave by (a) a thin prism, (b) a convex lens. (c) Reflection of a plane wave by a concave mirror.

Example 10.1

- When monochromatic light is incident on a surface separating two media, the reflected and refracted light both have the same frequency as the incident frequency. Explain why?
- When light travels from a rarer to a denser medium, the speed decreases. Does the reduction in speed imply a reduction in the energy carried by the light wave?
- In the wave picture of light, intensity of light is determined by the square of the amplitude of the wave. What determines the intensity of light in the photon picture of light.

Solution

- Reflection and refraction arise through interaction of incident light with the atomic constituents of matter. Atoms may be viewed as oscillators, which take up the frequency of the external agency (light) causing forced oscillations. The frequency of light emitted by a charged oscillator equals its frequency of oscillation. Thus, the frequency of scattered light equals the frequency of incident light.
- No. Energy carried by a wave depends on the amplitude of the wave, not on the speed of wave propagation.
- For a given frequency, intensity of light in the photon picture is determined by the number of photons crossing an unit area per unit time.

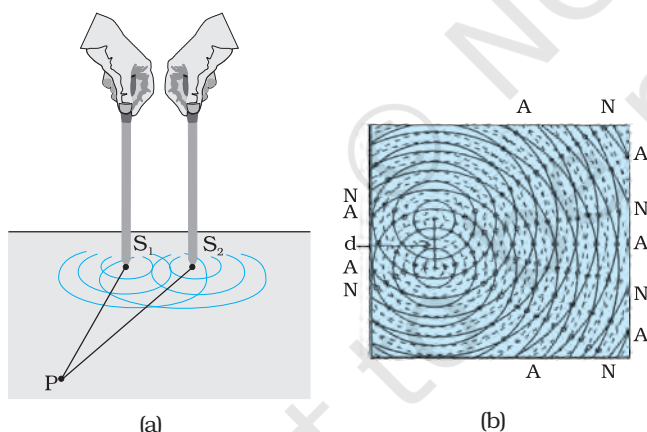


FIGURE 10.8 (a) Two needles oscillating in phase in water represent two coherent sources. (b) The pattern of displacement of water molecules at an instant on the surface of water showing nodal N (no displacement) and antinodal A (maximum displacement) lines.

fashion in a trough of water [Fig. 10.8(a)]. They produce two water waves, and at a particular point, the phase difference between the displacements produced by each of the waves does not change with time; when this happens the two sources are said to be *coherent*. Figure 10.8(b) shows the position of crests (solid circles) and troughs (dashed circles) at a given instant of time. Consider a point P for which

$$S_1 P = S_2 P$$

10.4 COHERENT AND INCOHERENT ADDITION OF WAVES

In this section we will discuss the interference pattern produced by the superposition of two waves. You may recall that we had discussed the superposition principle in Chapter 14 of your Class XI textbook. Indeed the entire field of interference is based on the *superposition principle* according to which *at a particular point in the medium, the resultant displacement produced by a number of waves is the vector sum of the displacements produced by each of the waves*.

Consider two needles S_1 and S_2 moving periodically up and down in an identical

Since the distances $S_1 P$ and $S_2 P$ are equal, waves from S_1 and S_2 will take the same time to travel to the point P and waves that emanate from S_1 and S_2 in phase will also arrive, at the point P , in phase.

Thus, if the displacement produced by the source S_1 at the point P is given by

$$y_1 = a \cos \omega t$$

then, the displacement produced by the source S_2 (at the point P) will also be given by

$$y_2 = a \cos \omega t$$

Thus, the resultant of displacement at P would be given by

$$y = y_1 + y_2 = 2 a \cos \omega t$$

Since the intensity is proportional to the square of the amplitude, the resultant intensity will be given by

$$I = 4 I_0$$

where I_0 represents the intensity produced by each one of the individual sources; I_0 is proportional to a^2 . In fact at any point on the perpendicular bisector of $S_1 S_2$, the intensity will be $4I_0$. The two sources are said to interfere constructively and we have what is referred to as *constructive interference*. We next consider a point Q [Fig. 10.9(a)] for which

$$S_2 Q - S_1 Q = 2\lambda$$

The waves emanating from S_1 will arrive exactly two cycles earlier than the waves from S_2 and will again be in phase [Fig. 10.9(a)]. Thus, if the displacement produced by S_1 is given by

$$y_1 = a \cos \omega t$$

then the displacement produced by S_2 will be given by

$$y_2 = a \cos (\omega t - 4\pi) = a \cos \omega t$$

where we have used the fact that a path difference of 2λ corresponds to a phase difference of 4π . The two displacements are once again in phase and the intensity will again be $4I_0$ giving rise to constructive interference. In the above analysis we have assumed that the distances $S_1 Q$ and $S_2 Q$ are much greater than d (which represents the distance between S_1 and S_2) so that although $S_1 Q$ and $S_2 Q$ are not equal, the amplitudes of the displacement produced by each wave are very nearly the same.

We next consider a point R [Fig. 10.9(b)] for which

$$S_2 R - S_1 R = -2.5\lambda$$

The waves emanating from S_1 will arrive exactly two and a half cycles later than the waves from S_2 [Fig. 10.10(b)]. Thus if the displacement produced by S_1 is given by

$$y_1 = a \cos \omega t$$

then the displacement produced by S_2 will be given by

$$y_2 = a \cos (\omega t + 5\pi) = -a \cos \omega t$$

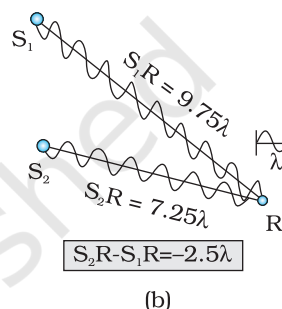
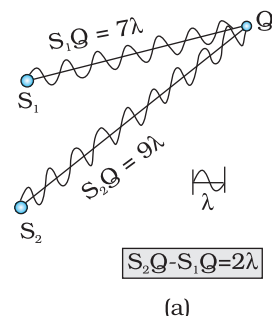


FIGURE 10.9

(a) Constructive interference at a point Q for which the path difference is 2λ .
(b) Destructive interference at a point R for which the path difference is 2.5λ .

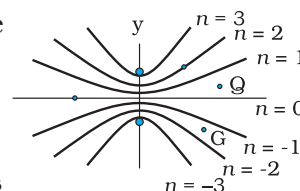
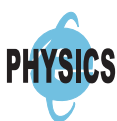


FIGURE 10.10 Locus of points for which $S_1P - S_2P$ is equal to zero, $\pm\lambda$, $\pm2\lambda$, $\pm3\lambda$.



where we have used the fact that a path difference of 2.5λ corresponds to a phase difference of 5π . The two displacements are now out of phase and the two displacements will cancel out to give zero intensity. This is referred to as *destructive interference*.

To summarise: If we have two coherent sources S_1 and S_2 vibrating in phase, then for an arbitrary point P whenever the path difference,

$$S_1P \sim S_2P = n\lambda \quad (n = 0, 1, 2, 3, \dots) \quad (10.9)$$

we will have constructive interference and the resultant intensity will be $4I_0$; the sign $\sim$ between S_1P and S_2P represents the difference between S_1P and S_2P . On the other hand, if the point P is such that the path difference,

$$S_1P \sim S_2P = (n + \frac{1}{2})\lambda \quad (n = 0, 1, 2, 3, \dots) \quad (10.10)$$

we will have *destructive interference* and the resultant intensity will be zero. Now, for any other arbitrary point G (Fig. 10.10) let the phase difference between the two displacements be ϕ . Thus, if the displacement produced by S_1 is given by

$$y_1 = a \cos \omega t$$

then, the displacement produced by S_2 would be

$$y_2 = a \cos (\omega t + \phi)$$

and the resultant displacement will be given by

$$\begin{aligned} y &= y_1 + y_2 \\ &= a [\cos \omega t + \cos (\omega t + \phi)] \\ &= 2a \cos (\phi/2) \cos (\omega t + \phi/2) \\ &\left[\because \cos A + \cos B = 2 \cos \left(\frac{A+B}{2} \right) \cos \left(\frac{A-B}{2} \right) \right] \end{aligned}$$

The amplitude of the resultant displacement is $2a \cos (\phi/2)$ and therefore the intensity at that point will be

$$I = 4 I_0 \cos^2 (\phi/2) \quad (10.11)$$

If $\phi = 0, \pm 2\pi, \pm 4\pi, \dots$ which corresponds to the condition given by Eq. (10.9) we will have constructive interference leading to maximum intensity. On the other hand, if $\phi = \pm \pi, \pm 3\pi, \pm 5\pi \dots$ [which corresponds to the condition given by Eq. (10.10)] we will have destructive interference leading to zero intensity.

Now if the two sources are coherent (i.e., if the two needles are going up and down regularly) then the phase difference ϕ at any point will not change with time and we will have a stable interference pattern; i.e., the positions of maxima and minima will not change with time. However, if the two needles do not maintain a constant phase difference, then the interference pattern will also change with time and, if the phase difference changes very rapidly with time, the positions of maxima and minima will also vary rapidly with time and we will see a “time-averaged” intensity distribution. When this happens, we will observe an average intensity that will be given by

$$I = 2 I_0 \quad (10.12)$$

at all points.

When the phase difference between the two vibrating sources changes rapidly with time, we say that the two sources are incoherent and when this happens the intensities just add up. This is indeed what happens when two separate light sources illuminate a wall.

10.5 INTERFERENCE OF LIGHT WAVES AND YOUNG'S EXPERIMENT

We will now discuss interference using light waves. If we use two sodium lamps illuminating two pinholes (Fig. 10.11) we will not observe any interference fringes. This is because of the fact that the light wave emitted from an ordinary source (like a sodium lamp) undergoes abrupt phase changes in times of the order of 10^{-10} seconds. Thus the light waves coming out from two independent sources of light will not have any fixed phase relationship and would be incoherent, when this happens, as discussed in the previous section, the intensities on the screen will add up.

The British physicist Thomas Young used an ingenious technique to “lock” the phases of the waves emanating from S_1 and S_2 . He made two pinholes S_1 and S_2 (very close to each other) on an opaque screen [Fig. 10.12(a)]. These were illuminated by another pinhole that was in turn, lit by a bright source. Light waves spread out from S and fall on both S_1 and S_2 . S_1 and S_2 then behave like two coherent sources because light waves coming out from S_1 and S_2 are derived from the same original source and any abrupt phase change in S will manifest in exactly similar phase changes in the light coming out from S_1 and S_2 . Thus, the two sources S_1 and S_2 will be *locked* in phase; i.e., they will be coherent like the two vibrating needle in our water wave example [Fig. 10.8(a)].

The spherical waves emanating from S_1 and S_2 will produce interference fringes on the screen GG' , as shown in Fig. 10.12(b). The positions of maximum and minimum intensities can be calculated by using the analysis given in Section 10.4.

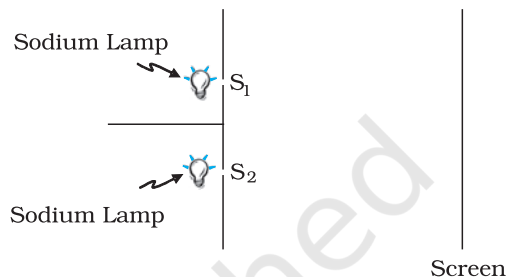


FIGURE 10.11 If two sodium lamps illuminate two pinholes S_1 and S_2 , the intensities will add up and no interference fringes will be observed on the screen.

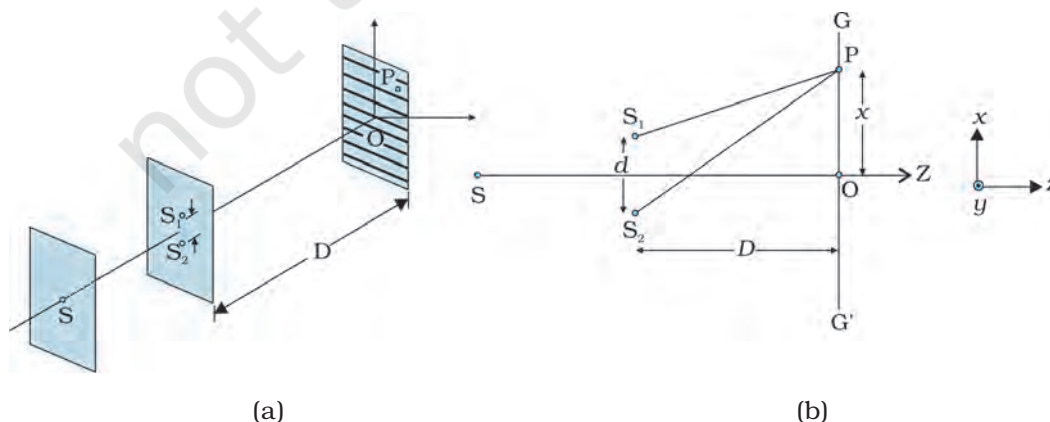


FIGURE 10.12 Young's arrangement to produce interference pattern.



Thomas Young (1773 – 1829) English physicist, physician and Egyptologist. Young worked on a wide variety of scientific problems, ranging from the structure of the eye and the mechanism of vision to the decipherment of the Rosetta stone. He revived the wave theory of light and recognised that interference phenomena provide proof of the wave properties of light.

We will have constructive interference resulting in a bright region when $\frac{xd}{D} = n\lambda$. That is,

$$x = x_n = \frac{n\lambda D}{d}; n = 0, \pm 1, \pm 2, \dots \quad (10.13)$$

On the other hand, we will have destructive interference resulting in a dark region when $\frac{xd}{D} = (n + \frac{1}{2})\lambda$ that is

$$x = x_n = (n + \frac{1}{2}) \frac{\lambda D}{d}; n = 0, \pm 1, \pm 2 \quad (10.14)$$

Thus dark and bright bands appear on the screen, as shown in Fig. 10.13. Such bands are called *fringes*. Equations (10.13) and (10.14) show that dark and bright fringes are equally spaced.

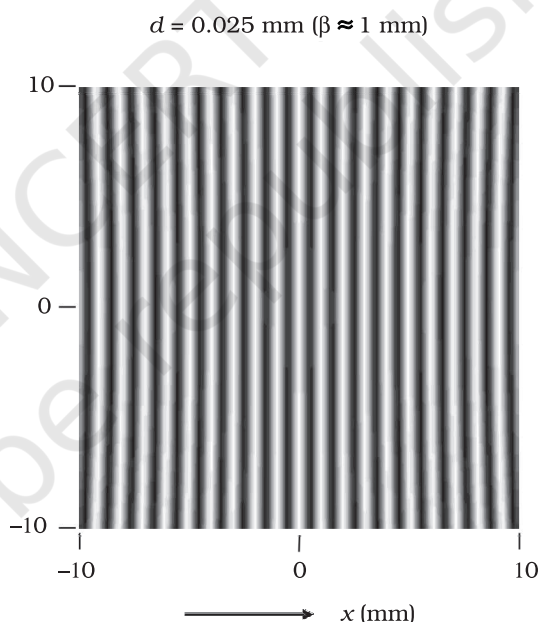


FIGURE 10.13 Computer generated fringe pattern produced by two point source S_1 and S_2 on the screen GG' (Fig. 10.12); $d = 0.025$ mm, $D = 5$ cm and $\lambda = 5 \times 10^{-5}$ cm.) (Adopted from OPTICS by A. Ghatak, Tata McGraw Hill Publishing Co. Ltd., New Delhi, 2000.)

10.6 DIFFRACTION

If we look clearly at the shadow cast by an opaque object, close to the region of geometrical shadow, there are alternate dark and bright regions just like in interference. This happens due to the phenomenon of diffraction. Diffraction is a general characteristic exhibited by all types of waves, be it sound waves, light waves, water waves or matter waves. Since the wavelength of light is much smaller than the dimensions of most obstacles; we do not encounter diffraction effects of light in everyday

observations. However, the finite resolution of our eye or of optical instruments such as telescopes or microscopes is limited due to the phenomenon of diffraction. Indeed the colours that you see when a CD is viewed is due to diffraction effects. We will now discuss the phenomenon of diffraction.

10.6.1 The single slit

In the discussion of Young's experiment, we stated that a single narrow slit acts as a new source from which light spreads out. Even before Young, early experimenters – including Newton – had noticed that light spreads out from narrow holes and slits. It seems to turn around corners and enter regions where we would expect a shadow. These effects, known as *diffraction*, can only be properly understood using wave ideas. After all, you are hardly surprised to hear sound waves from someone talking around a corner!

When the double slit in Young's experiment is replaced by a single narrow slit (illuminated by a monochromatic source), a broad pattern with a central bright region is seen. On both sides, there are alternate dark and bright regions, the intensity becoming weaker away from the centre (Fig. 10.15). To understand this, go to Fig. 10.14, which shows a parallel beam of light falling normally on a single slit LN of width a . The diffracted light goes on to meet a screen. The midpoint of the slit is M.

A straight line through M perpendicular to the slit plane meets the screen at C. We want the intensity at any point P on the screen. As before, straight lines joining P to the different points L, M, N, etc., can be treated as parallel, making an angle θ with the normal MC.

The basic idea is to divide the slit into much smaller parts, and add their contributions at P with the proper phase differences. We are treating different parts of the wavefront at the slit as secondary sources. Because the incoming wavefront is parallel to the plane of the slit, these sources are in phase.

It is observed that the intensity has a central maximum at $\theta = 0$ and other secondary maxima at $\theta \approx (n+1/2) \lambda/a$, which go on becoming weaker and weaker with increasing n . The minima (zero intensity) are at $\theta \approx n\lambda/a$, $n = \pm 1, \pm 2, \pm 3, \dots$

The photograph and intensity pattern corresponding to it is shown in Fig. 10.15.

There has been prolonged discussion about difference between interference and diffraction among

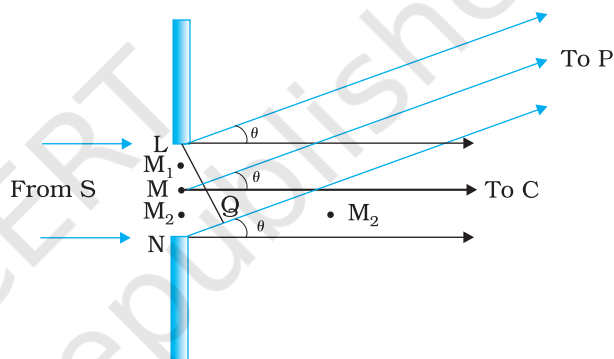


FIGURE 10.14 The geometry of path differences for diffraction by a single slit.

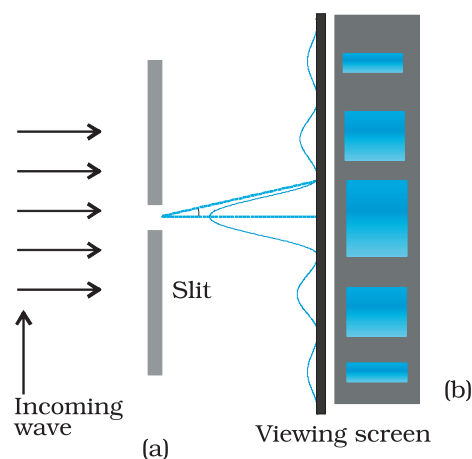


FIGURE 10.15 Intensity distribution and photograph of fringes due to diffraction at single slit.

scientists since the discovery of these phenomena. In this context, it is interesting to note what Richard Feynman* has said in his famous Feynman Lectures on Physics:

No one has ever been able to define the difference between interference and diffraction satisfactorily. It is just a question of usage, and there is no specific, important physical difference between them. The best we can do is, roughly speaking, is to say that when there are only a few sources, say two interfering sources, then the result is usually called interference, but if there is a large number of them, it seems that the word diffraction is more often used.

In the double-slit experiment, we must note that the pattern on the screen is actually a superposition of single-slit diffraction from each slit or hole, and the double-slit interference pattern.

10.6.2 Seeing the single slit diffraction pattern

It is surprisingly easy to see the single-slit diffraction pattern for oneself. The equipment needed can be found in most homes — two razor blades and one clear glass electric bulb preferably with a straight filament. One has to hold the two blades so that the edges are parallel and have a narrow slit in between. This is easily done with the thumb and forefingers (Fig. 10.16).

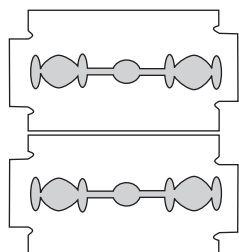


FIGURE 10.16

Holding two blades to form a single slit. A bulb filament viewed through this shows clear diffraction bands.

Keep the slit parallel to the filament, right in front of the eye. Use spectacles if you normally do. With slight adjustment of the width of the slit and the parallelism of the edges, the pattern should be seen with its bright and dark bands. Since the position of all the bands (except the central one) depends on wavelength, they will show some colours. Using a filter for red or blue will make the fringes clearer. With both filters available, the wider fringes for red compared to blue can be seen.

In this experiment, the filament plays the role of the first slit S in Fig. 10.15. The lens of the eye focuses the pattern on the screen (the retina of the eye).

With some effort, one can cut a double slit in an aluminium foil with a blade. The bulb filament can be viewed as before to repeat Young's experiment. In daytime, there is another suitable bright source subtending a small angle at the eye. This is the reflection of the Sun in any shiny convex surface (e.g., a cycle bell). Do not try direct sunlight – it can damage the eye and will not give fringes anyway as the Sun subtends an angle of $(1/2)^\circ$.

In interference and diffraction, light energy is redistributed. If it reduces in one region, producing a dark fringe, it increases in another region, producing a bright fringe. There is no gain or loss of energy, which is consistent with the principle of conservation of energy.

* Richard Feynman was one of the recipients of the 1965 Nobel Prize in Physics for his fundamental work in quantum electrodynamics.

10.7 POLARISATION

Consider holding a long string that is held horizontally, the other end of which is assumed to be fixed. If we move the end of the string up and down in a periodic manner, we will generate a wave propagating in the $+x$ direction (Fig. 10.17). Such a wave could be described by the following equation

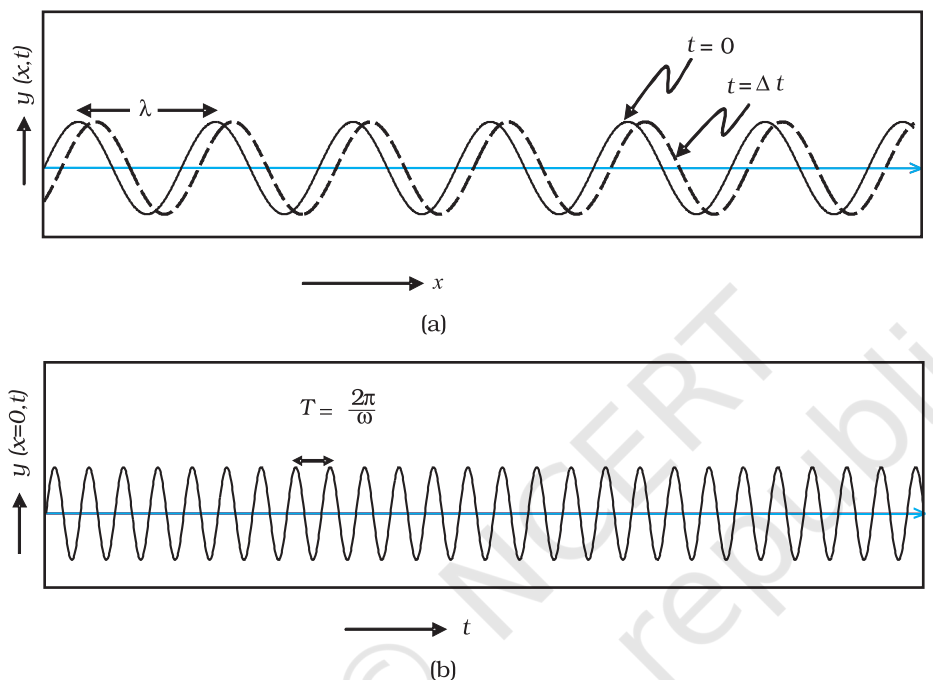


FIGURE 10.17 (a) The curves represent the displacement of a string at $t = 0$ and at $t = \Delta t$, respectively when a sinusoidal wave is propagating in the $+x$ -direction. (b) The curve represents the time variation of the displacement at $x = 0$ when a sinusoidal wave is propagating in the $+x$ -direction. At $x = \Delta x$, the time variation of the displacement will be slightly displaced to the right.

$$y(x, t) = a \sin(kx - \omega t) \quad (10.15)$$

where a and $\omega (= 2\pi\nu)$ represent the amplitude and the angular frequency of the wave, respectively; further,

$$\lambda = \frac{2\pi}{k} \quad (10.16)$$

represents the wavelength associated with the wave. We had discussed propagation of such waves in Chapter 14 of Class XI textbook. Since the displacement (which is along the y direction) is at right angles to the direction of propagation of the wave, we have what is known as a *transverse wave*. Also, since the displacement is in the y direction, it is often referred to as a y -polarised wave. Since each point on the string moves on a straight line, the wave is also referred to as a linearly polarised

wave. Further, the string always remains confined to the x - y plane and therefore it is also referred to as a *plane polarised wave*.

In a similar manner we can consider the vibration of the string in the x - z plane generating a z -polarised wave whose displacement will be given by

$$z(x, t) = a \sin(kx - \omega t) \quad (10.17)$$

It should be mentioned that the linearly polarised waves [described by Eqs. (10.15) and (10.17)] are all transverse waves; i.e., the displacement of each point of the string is always at right angles to the direction of propagation of the wave. Finally, if the plane of vibration of the string is changed randomly in very short intervals of time, then we have what is known as an *unpolarised wave*. Thus, for an unpolarised wave the displacement will be randomly changing with time though it will always be perpendicular to the direction of propagation.

Light waves are transverse in nature; i.e., the electric field associated with a propagating light wave is always at right angles to the direction of propagation of the wave. This can be easily demonstrated using a simple polaroid. You must have seen thin plastic like sheets, which are called *polaroids*. A polaroid consists of long chain molecules aligned in a particular direction. The electric vectors (associated with the propagating light wave) along the direction of the aligned molecules get absorbed. Thus, if an unpolarised light wave is incident on such a polaroid then the light wave will get linearly polarised with the electric vector oscillating along a direction perpendicular to the aligned molecules; this direction is known as the *pass-axis* of the polaroid.

Thus, if the light from an ordinary source (like a sodium lamp) passes through a polaroid sheet P_1 , it is observed that its intensity is reduced by half. Rotating P_1 has no effect on the transmitted beam and transmitted intensity remains constant. Now, let an identical piece of polaroid P_2 be placed before P_1 . As expected, the light from the lamp is reduced in intensity on passing through P_2 alone. But now rotating P_1 has a dramatic effect on the light coming from P_2 . In one position, the intensity transmitted by P_2 followed by P_1 is nearly zero. When turned by 90° from this position, P_1 transmits nearly the full intensity emerging from P_2 (Fig. 10.18).

The experiment at figure 10.18 can be easily understood by assuming that light passing through the polaroid P_2 gets polarised along the pass-axis of P_2 . If the pass-axis of P_2 makes an angle θ with the pass-axis of P_1 , then when the polarised beam passes through the polaroid P_2 , the component $E \cos \theta$ (along the pass-axis of P_2) will pass through P_2 . Thus, as we rotate the polaroid P_1 (or P_2), the intensity will vary as:

$$I = I_0 \cos^2 \theta \quad (10.18)$$

where I_0 is the intensity of the polarized light after passing through P_1 . This is known as *Malus' law*. The above discussion shows that the

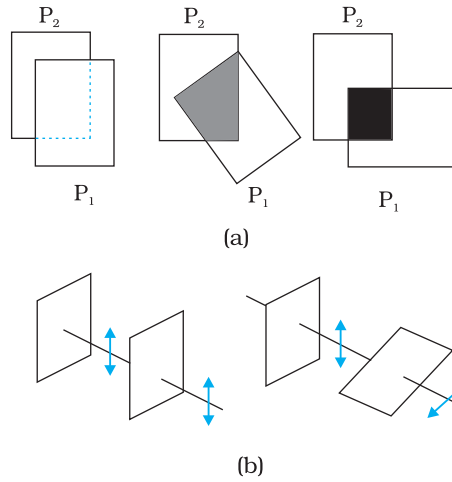


FIGURE 10.18 (a) Passage of light through two polaroids P_2 and P_1 . The transmitted fraction falls from 1 to 0 as the angle between them varies from 0° to 90° . Notice that the light seen through a single polaroid P_1 does not vary with angle. (b) Behaviour of the electric vector when light passes through two polaroids. The transmitted polarisation is the component parallel to the polaroid axis. The double arrows show the oscillations of the electric vector.

intensity coming out of a single polaroid is half of the incident intensity. By putting a second polaroid, the intensity can be further controlled from 50% to zero of the incident intensity by adjusting the angle between the pass-axes of two polaroids.

Polaroids can be used to control the intensity, in sunglasses, windowpanes, etc. Polaroids are also used in photographic cameras and 3D movie cameras.

Example 10.2 Discuss the intensity of transmitted light when a polaroid sheet is rotated between two crossed polaroids?

Solution Let I_0 be the intensity of polarised light after passing through the first polariser P_1 . Then the intensity of light after passing through second polariser P_2 will be

$$I = I_0 \cos^2 \theta,$$

where θ is the angle between pass axes of P_1 and P_2 . Since P_1 and P_3 are crossed the angle between the pass axes of P_2 and P_3 will be $(\pi/2 - \theta)$. Hence the intensity of light emerging from P_3 will be

$$\begin{aligned} I &= I_0 \cos^2 \theta \cos^2 \left(\frac{\pi}{2} - \theta \right) \\ &= I_0 \cos^2 \theta \sin^2 \theta = (I_0/4) \sin^2 2\theta \end{aligned}$$

Therefore, the transmitted intensity will be maximum when $\theta = \pi/4$.

SUMMARY

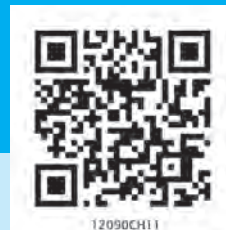
1. Huygens' principle tells us that each point on a wavefront is a source of secondary waves, which add up to give the wavefront at a later time.
2. Huygens' construction tells us that the new wavefront is the forward envelope of the secondary waves. When the speed of light is independent of direction, the secondary waves are spherical. The rays are then perpendicular to both the wavefronts and the time of travel is the same measured along any ray. This principle leads to the well known laws of reflection and refraction.
3. The principle of superposition of waves applies whenever two or more sources of light illuminate the same point. When we consider the intensity of light due to these sources at the given point, there is an interference term in addition to the sum of the individual intensities. But this term is important only if it has a non-zero average, which occurs only if the sources have the same frequency and a stable phase difference.
4. Young's double slit of separation d gives equally spaced interference fringes.
5. A single slit of width a gives a diffraction pattern with a central maximum. The intensity falls to zero at angles of $\pm \frac{\lambda}{a}, \pm \frac{2\lambda}{a}$, etc., with successively weaker secondary maxima in between.
6. Natural light, e.g., from the sun is unpolarised. This means the electric vector takes all possible directions in the transverse plane, rapidly and randomly, during a measurement. A polaroid transmits only one component (parallel to a special axis). The resulting light is called linearly polarised or plane polarised. When this kind of light is viewed through a second polaroid whose axis turns through 2π , two maxima and minima of intensity are seen.

POINTS TO PONDER

1. Waves from a point source spread out in all directions, while light was seen to travel along narrow rays. It required the insight and experiment of Huygens, Young and Fresnel to understand how a wave theory could explain all aspects of the behaviour of light.
2. The crucial new feature of waves is interference of amplitudes from different sources which can be both constructive and destructive, as shown in Young's experiment.
3. Diffraction phenomena define the limits of ray optics. The limit of the ability of microscopes and telescopes to distinguish very close objects is set by the wavelength of light.
4. Most interference and diffraction effects exist even for longitudinal waves like sound in air. But polarisation phenomena are special to transverse waves like light waves.

EXERCISES

- 10.1** Monochromatic light of wavelength 589 nm is incident from air on a water surface. What are the wavelength, frequency and speed of (a) reflected, and (b) refracted light? Refractive index of water is 1.33.
- 10.2** What is the shape of the wavefront in each of the following cases:
- (a) Light diverging from a point source.
 - (b) Light emerging out of a convex lens when a point source is placed at its focus.
 - (c) The portion of the wavefront of light from a distant star intercepted by the Earth.
- 10.3** (a) The refractive index of glass is 1.5. What is the speed of light in glass? (Speed of light in vacuum is $3.0 \times 10^8 \text{ m s}^{-1}$)
- (b) Is the speed of light in glass independent of the colour of light? If not, which of the two colours red and violet travels slower in a glass prism?
- 10.4** In a Young's double-slit experiment, the slits are separated by 0.28 mm and the screen is placed 1.4 m away. The distance between the central bright fringe and the fourth bright fringe is measured to be 1.2 cm. Determine the wavelength of light used in the experiment.
- 10.5** In Young's double-slit experiment using monochromatic light of wavelength λ , the intensity of light at a point on the screen where path difference is λ , is K units. What is the intensity of light at a point where path difference is $\lambda/3$?
- 10.6** A beam of light consisting of two wavelengths, 650 nm and 520 nm, is used to obtain interference fringes in a Young's double-slit experiment.
- (a) Find the distance of the third bright fringe on the screen from the central maximum for wavelength 650 nm.
 - (b) What is the least distance from the central maximum where the bright fringes due to both the wavelengths coincide?



Chapter Eleven

DUAL NATURE OF RADIATION AND MATTER



11.1 INTRODUCTION

The Maxwell's equations of electromagnetism and Hertz experiments on the generation and detection of electromagnetic waves in 1887 strongly established the wave nature of light. Towards the same period at the end of 19th century, experimental investigations on conduction of electricity (electric discharge) through gases at low pressure in a discharge tube led to many historic discoveries. The discovery of X-rays by Roentgen in 1895, and of electron by J. J. Thomson in 1897, were important milestones in the understanding of atomic structure. It was found that at sufficiently low pressure of about 0.001 mm of mercury column, a discharge took place between the two electrodes on applying the electric field to the gas in the discharge tube. A fluorescent glow appeared on the glass opposite to cathode. The colour of glow of the glass depended on the type of glass, it being yellowish-green for soda glass. The cause of this fluorescence was attributed to the radiation which appeared to be coming from the cathode. These *cathode rays* were discovered, in 1870, by William Crookes who later, in 1879, suggested that these rays consisted of streams of fast moving negatively charged particles. The British physicist J. J. Thomson (1856-1940) confirmed this hypothesis. By applying mutually perpendicular electric and magnetic fields across the discharge tube, J. J. Thomson was the first to determine experimentally the speed

and the specific charge [charge to mass ratio (e/m)] of the cathode ray particles. They were found to travel with speeds ranging from about 0.1 to 0.2 times the speed of light (3×10^8 m/s). The presently accepted value of e/m is 1.76×10^{11} C/kg. Further, the value of e/m was found to be independent of the nature of the material/metal used as the cathode (emitter), or the gas introduced in the discharge tube. This observation suggested the universality of the cathode ray particles.

Around the same time, in 1887, it was found that certain metals, when irradiated by ultraviolet light, emitted negatively charged particles having small speeds. Also, certain metals when heated to a high temperature were found to emit negatively charged particles. The value of e/m of these particles was found to be the same as that for cathode ray particles. These observations thus established that all these particles, although produced under different conditions, were identical in nature. J. J. Thomson, in 1897, named these particles as *electrons*, and suggested that they were fundamental, universal constituents of matter. For his epoch-making discovery of electron, through his theoretical and experimental investigations on conduction of electricity by gasses, he was awarded the Nobel Prize in Physics in 1906. In 1913, the American physicist R. A. Millikan (1868-1953) performed the pioneering oil-drop experiment for the precise measurement of the charge on an electron. He found that the charge on an oil-droplet was always an integral multiple of an elementary charge, 1.602×10^{-19} C. Millikan's experiment established that *electric charge is quantised*. From the values of charge (e) and specific charge (e/m), the mass (m) of the electron could be determined.

11.2 ELECTRON EMISSION

We know that metals have free electrons (negatively charged particles) that are responsible for their conductivity. However, the free electrons cannot normally escape out of the metal surface. If an electron attempts to come out of the metal, the metal surface acquires a positive charge and pulls the electron back to the metal. The free electron is thus held inside the metal surface by the attractive forces of the ions. Consequently, the electron can come out of the metal surface only if it has got sufficient energy to overcome the attractive pull. A certain minimum amount of energy is required to be given to an electron to pull it out from the surface of the metal. This minimum energy required by an electron to escape from the metal surface is called the *work function* of the metal. It is generally denoted by ϕ_0 and measured in eV (electron volt). One electron volt is the energy gained by an electron when it has been accelerated by a potential difference of 1 volt, so that $1 \text{ eV} = 1.602 \times 10^{-19} \text{ J}$.

This unit of energy is commonly used in atomic and nuclear physics. The work function (ϕ_0) depends on the properties of the metal and the nature of its surface.

The minimum energy required for the electron emission from the metal surface can be supplied to the free electrons by any one of the following physical processes:

- (i) *Thermionic emission*: By suitably heating, sufficient thermal energy can be imparted to the free electrons to enable them to come out of the metal.

- (ii) *Field emission*: By applying a very strong electric field (of the order of 10^8 V m^{-1}) to a metal, electrons can be pulled out of the metal, as in a spark plug.
- (iii) *Photoelectric emission*: When light of suitable frequency illuminates a metal surface, electrons are emitted from the metal surface. These photo(light)-generated electrons are called *photoelectrons*.

11.3 PHOTOELECTRIC EFFECT

11.3.1 Hertz's observations

The phenomenon of photoelectric emission was discovered in 1887 by Heinrich Hertz (1857-1894), during his electromagnetic wave experiments. In his experimental investigation on the production of electromagnetic waves by means of a spark discharge, Hertz observed that high voltage sparks across the detector loop were enhanced when the emitter plate was illuminated by ultraviolet light from an arc lamp.

Light shining on the metal surface somehow facilitated the escape of free, charged particles which we now know as electrons. When light falls on a metal surface, some electrons near the surface absorb enough energy from the incident radiation to overcome the attraction of the positive ions in the material of the surface. After gaining sufficient energy from the incident light, the electrons escape from the surface of the metal into the surrounding space.

11.3.2 Hallwachs' and Lenard's observations

Wilhelm Hallwachs and Philipp Lenard investigated the phenomenon of photoelectric emission in detail during 1886-1902.

Lenard (1862-1947) observed that when ultraviolet radiations were allowed to fall on the emitter plate of an evacuated glass tube enclosing two electrodes (metal plates), current flows in the circuit (Fig. 11.1). As soon as the ultraviolet radiations were stopped, the current flow also stopped. These observations indicate that when ultraviolet radiations fall on the emitter plate C, electrons are ejected from it which are attracted towards the positive, collector plate A by the electric field. The electrons flow through the evacuated glass tube, resulting in the current flow. Thus, light falling on the surface of the emitter causes current in the external circuit. Hallwachs and Lenard studied how this photo current varied with collector plate potential, and with frequency and intensity of incident light.

Hallwachs, in 1888, undertook the study further and connected a negatively charged zinc plate to an electroscope. He observed that the zinc plate lost its charge when it was illuminated by ultraviolet light. Further, the uncharged zinc plate became positively charged when it was irradiated by ultraviolet light. Positive charge on a positively charged zinc plate was found to be further enhanced when it was illuminated by ultraviolet light. From these observations he concluded that negatively charged particles were emitted from the zinc plate under the action of ultraviolet light.

After the discovery of the electron in 1897, it became evident that the incident light causes electrons to be emitted from the emitter plate. Due

to negative charge, the emitted electrons are pushed towards the collector plate by the electric field. Hallwachs and Lenard also observed that when ultraviolet light fell on the emitter plate, no electrons were emitted at all when the frequency of the incident light was smaller than a certain minimum value, called the *threshold frequency*. This minimum frequency depends on the nature of the material of the emitter plate.

It was found that certain metals like zinc, cadmium, magnesium, etc., responded only to ultraviolet light, having short wavelength, to cause electron emission from the surface. However, some alkali metals such as lithium, sodium, potassium, caesium and rubidium were sensitive even to visible light. All these *photosensitive substances* emit electrons when they are illuminated by light. After the discovery of electrons, these electrons were termed as *photoelectrons*. The phenomenon is called *photoelectric effect*.

11.4 EXPERIMENTAL STUDY OF PHOTOELECTRIC EFFECT

Figure 11.1 depicts a schematic view of the arrangement used for the experimental study of the photoelectric effect. It consists of an evacuated glass/quartz tube having a thin photosensitive plate C and another metal plate A. Monochromatic light from the source S of sufficiently short wavelength passes through the window W and falls on the photosensitive plate C (emitter). A transparent quartz window is sealed on to the glass tube, which permits ultraviolet radiation to pass through it and irradiate the photosensitive plate C. The electrons are emitted by the plate C and are collected by the plate A (collector), by the electric field created by the battery. The battery maintains the potential difference between the plates C and A, that can be varied. The polarity of the plates C and A can be reversed by a commutator. Thus, the plate A can be maintained at a desired positive or negative potential with respect to emitter C. When the collector plate A is positive with respect to the emitter plate C, the electrons are attracted to it. The emission of electrons causes flow of electric current in the circuit. The potential difference between the emitter and collector plates is measured by a voltmeter (V) whereas the resulting photo current flowing in the circuit is measured by a microammeter (μA). The photoelectric current can be increased or decreased by varying the potential of collector plate A with respect to the emitter plate C. The intensity and frequency of the incident light can be varied, as can the potential difference V between the emitter C and the collector A.

We can use the experimental arrangement of Fig. 11.1 to study the variation of photocurrent with (a) intensity of radiation, (b) frequency of incident radiation, (c) the potential difference between the plates A and C, and (d) the nature of the material of plate C. Light of different frequencies can be used by putting appropriate coloured filter or coloured glass in the path of light falling

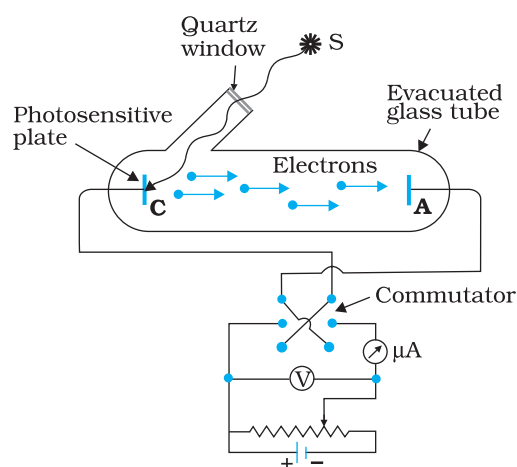


FIGURE 11.1 Experimental arrangement for study of photoelectric effect.

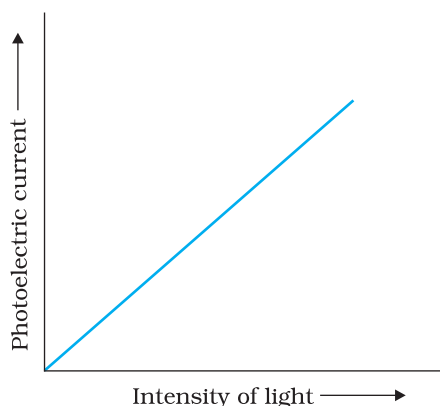


FIGURE 11.2 Variation of Photoelectric current with intensity of light.

on the emitter C. The intensity of light is varied by changing the distance of the light source from the emitter.

11.4.1 Effect of intensity of light on photocurrent

The collector A is maintained at a positive potential with respect to emitter C so that electrons ejected from C are attracted towards collector A. Keeping the frequency of the incident radiation and the potential fixed, the intensity of light is varied and the resulting photoelectric current is measured each time. It is found that the photocurrent increases linearly with intensity of incident light as shown graphically in Fig. 11.2. The photocurrent is directly proportional to the number of photoelectrons emitted per second. This implies that *the number of photoelectrons emitted per second is directly proportional to the intensity of incident radiation.*

11.4.2 Effect of potential on photoelectric current

We first keep the plate A at some positive potential with respect to the plate C and illuminate the plate C with light of fixed frequency ν and fixed intensity I_1 . We next vary the positive potential of plate A gradually and measure the resulting photocurrent each time. It is found that the photoelectric current increases with increase in positive (accelerating) potential. At some stage, for a certain positive potential of plate A, all the emitted electrons are collected by the plate A and the photoelectric current becomes maximum or saturates. If we increase the accelerating potential of plate A further, the photocurrent does not increase. This maximum value of the photoelectric current is called *saturation current*. Saturation current corresponds to the case when all the photoelectrons emitted by the emitter plate C reach the collector plate A.

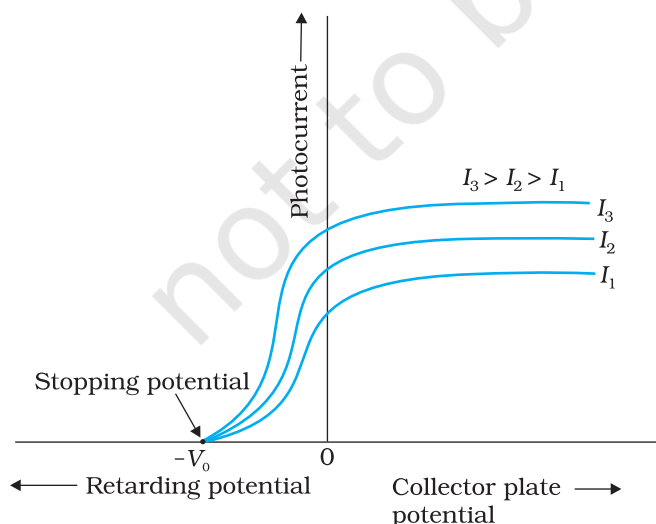


FIGURE 11.3 Variation of photocurrent with collector plate potential for different intensity of incident radiation.

We now apply a negative (retarding) potential to the plate A with respect to the plate C and make it increasingly negative gradually. When the polarity is reversed, the electrons are repelled and only the sufficiently energetic electrons are able to reach the collector A. The photocurrent is found to decrease rapidly until it drops to zero at a certain sharply defined, critical value of the negative potential V_0 on the plate A. For a particular frequency of incident radiation, *the minimum negative (retarding) potential V_0 given to the plate A for which the photocurrent stops or becomes zero is called the cut-off or stopping potential.*

The interpretation of the observation in terms of photoelectrons is straightforward. All the photoelectrons emitted from the metal do not have the

same energy. Photoelectric current is zero when the stopping potential is sufficient to repel even the most energetic photoelectrons, with the maximum kinetic energy ($K_{\max}$), so that

$$K_{\max} = e V_0 \quad (11.1)$$

We can now repeat this experiment with incident radiation of the same frequency but of higher intensity I_2 and I_3 ($I_3 > I_2 > I_1$). We note that the saturation currents are now found to be at higher values. This shows that more electrons are being emitted per second, proportional to the intensity of incident radiation. But the stopping potential remains the same as that for the incident radiation of intensity I_1 , as shown graphically in Fig. 11.3. Thus, *for a given frequency of the incident radiation, the stopping potential is independent of its intensity*. In other words, the maximum kinetic energy of photoelectrons depends on the light source and the emitter plate material, but is independent of intensity of incident radiation.

11.4.3 Effect of frequency of incident radiation on stopping potential

We now study the relation between the frequency ν of the incident radiation and the stopping potential V_0 . We suitably adjust the same intensity of light radiation at various frequencies and study the variation of photocurrent with collector plate potential. The resulting variation is shown in Fig. 11.4. We obtain different values of stopping potential but the same value of the saturation current for incident radiation of different frequencies. The energy of the emitted electrons depends on the frequency of the incident radiations. The stopping potential is more negative for higher frequencies of incident radiation. Note from Fig. 11.4 that the stopping potentials are in the order $V_{03} > V_{02} > V_{01}$ if the frequencies are in the order $\nu_3 > \nu_2 > \nu_1$. This implies that greater the frequency of incident light, greater is the maximum kinetic energy of the photoelectrons. Consequently, we need greater retarding potential to stop them completely. If we plot a graph between the frequency of incident radiation and the corresponding stopping potential for different metals we get a straight line, as shown in Fig. 11.5.

The graph shows that

- the stopping potential V_0 varies linearly with the frequency of incident radiation for a given photosensitive material.
- there exists a certain minimum cut-off frequency ν_0 for which the stopping potential is zero.

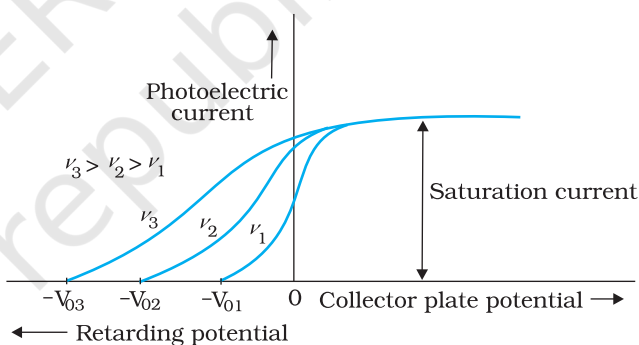


FIGURE 11.4 Variation of photoelectric current with collector plate potential for different frequencies of incident radiation.

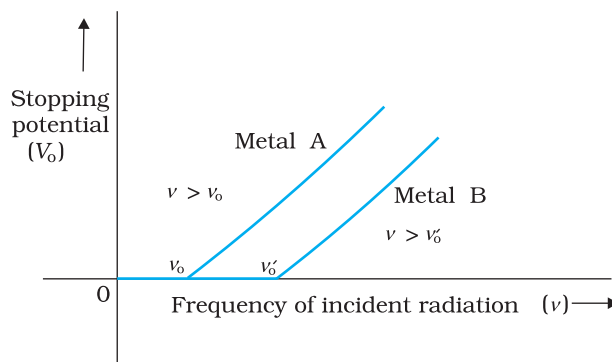


FIGURE 11.5 Variation of stopping potential V_0 with frequency ν of incident radiation for a given photosensitive material.

These observations have two implications:

- (i) *The maximum kinetic energy of the photoelectrons varies linearly with the frequency of incident radiation, but is independent of its intensity.*
- (ii) *For a frequency ν of incident radiation, lower than the cut-off frequency ν_0 , no photoelectric emission is possible even if the intensity is large.*

This minimum, cut-off frequency ν_0 , is called the *threshold frequency*. It is different for different metals.

Different photosensitive materials respond differently to light. Selenium is more sensitive than zinc or copper. The same photosensitive substance gives different response to light of different wavelengths. For example, ultraviolet light gives rise to photoelectric effect in copper while green or red light does not.

Note that in all the above experiments, it is found that, if frequency of the incident radiation exceeds the threshold frequency, the photoelectric emission starts instantaneously without any apparent time lag, even if the incident radiation is very dim. It is now known that emission starts in a time of the order of 10^{-9} s or less.

We now summarise the experimental features and observations described in this section.

- (i) For a given photosensitive material and frequency of incident radiation (above the threshold frequency), the photoelectric current is directly proportional to the intensity of incident light (Fig. 11.2).
- (ii) For a given photosensitive material and frequency of incident radiation, saturation current is found to be proportional to the intensity of incident radiation whereas the stopping potential is independent of its intensity (Fig. 11.3).
- (iii) For a given photosensitive material, there exists a certain minimum cut-off frequency of the incident radiation, called the *threshold frequency*, below which no emission of photoelectrons takes place, no matter how intense the incident light is. Above the threshold frequency, the stopping potential or equivalently the maximum kinetic energy of the emitted photoelectrons increases linearly with the frequency of the incident radiation, but is independent of its intensity (Fig. 11.5).
- (iv) The photoelectric emission is an instantaneous process without any apparent time lag ($\sim 10^{-9}$ s or less), even when the incident radiation is made exceedingly dim.

11.5 PHOTOELECTRIC EFFECT AND WAVE THEORY OF LIGHT

The wave nature of light was well established by the end of the nineteenth century. The phenomena of interference, diffraction and polarisation were explained in a natural and satisfactory way by the wave picture of light. According to this picture, light is an electromagnetic wave consisting of electric and magnetic fields with continuous distribution of energy over the region of space over which the wave is extended. Let us now see if this

wave picture of light can explain the observations on photoelectric emission given in the previous section.

According to the wave picture of light, the free electrons at the surface of the metal (over which the beam of radiation falls) absorb the radiant energy continuously. The greater the intensity of radiation, the greater are the amplitude of electric and magnetic fields. Consequently, the greater the intensity, the greater should be the energy absorbed by each electron. In this picture, the maximum kinetic energy of the photoelectrons on the surface is then expected to increase with increase in intensity. Also, no matter what the frequency of radiation is, a sufficiently intense beam of radiation (over sufficient time) should be able to impart enough energy to the electrons, so that they exceed the minimum energy needed to escape from the metal surface. A threshold frequency, therefore, should not exist. These expectations of the wave theory directly contradict observations (i), (ii) and (iii) given at the end of sub-section 11.4.3.

Further, we should note that in the wave picture, the absorption of energy by electron takes place continuously over the entire wavefront of the radiation. Since a large number of electrons absorb energy, the energy absorbed per electron per unit time turns out to be small. Explicit calculations estimate that it can take hours or more for a single electron to pick up sufficient energy to overcome the work function and come out of the metal. This conclusion is again in striking contrast to observation (iv) that the photoelectric emission is instantaneous. In short, the wave picture is unable to explain the most basic features of photoelectric emission.

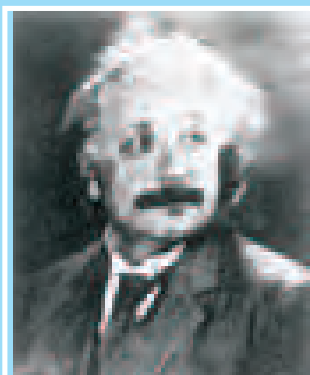
11.6 EINSTEIN'S PHOTOELECTRIC EQUATION: ENERGY QUANTUM OF RADIATION

In 1905, Albert Einstein (1879-1955) proposed a radically new picture of electromagnetic radiation to explain photoelectric effect. In this picture, photoelectric emission does not take place by continuous absorption of energy from radiation. Radiation energy is built up of discrete units – the so called *quanta of energy of radiation*. Each quantum of radiant energy has energy $h\nu$, where h is Planck's constant and ν the frequency of light. In photoelectric effect, an electron absorbs a quantum of energy ($h\nu$) of radiation. If this quantum of energy absorbed exceeds the minimum energy needed for the electron to escape from the metal surface (work function ϕ_0), the electron is emitted with maximum kinetic energy

$$K_{\max} = h\nu - \phi_0 \quad (11.2)$$

More tightly bound electrons will emerge with kinetic energies less than the maximum value. Note that the intensity of light of a given frequency is determined by the number of photons incident per second. Increasing the intensity will increase the number of emitted electrons per second. However, the maximum kinetic energy of the emitted photoelectrons is determined by the energy of each photon.

Equation (11.2) is known as *Einstein's photoelectric equation*. We now see how this equation accounts in a simple and elegant manner all the observations on photoelectric effect given at the end of sub-section 11.4.3.



Albert Einstein (1879 – 1955) Einstein, one of the greatest physicists of all time, was born in Ulm, Germany. In 1905, he published three path-breaking papers. In the first paper, he introduced the notion of light quanta (now called photons) and used it to explain the features of photoelectric effect. In the second paper, he developed a theory of Brownian motion, confirmed experimentally a few years later and provided a convincing evidence of the atomic picture of matter. The third paper gave birth to the special theory of relativity. In 1916, he published the general theory of relativity. Some of Einstein's most significant later contributions are: the notion of stimulated emission introduced in an alternative derivation of Planck's blackbody radiation law, static model of the universe which started modern cosmology, quantum statistics of a gas of massive bosons, and a critical analysis of the foundations of quantum mechanics. In 1921, he was awarded the Nobel Prize in physics for his contribution to theoretical physics and the photoelectric effect.

- According to Eq. (11.2), $K_{\max}$ depends linearly on ν , and is independent of intensity of radiation, in agreement with observation. This has happened because in Einstein's picture, photoelectric effect arises from the absorption of a single quantum of radiation by a single electron. The intensity of radiation (that is proportional to the number of energy quanta per unit area per unit time) is irrelevant to this basic process.
- Since $K_{\max}$ must be non-negative, Eq. (11.2) implies that photoelectric emission is possible only if $h\nu > \phi_0$ or $\nu > \nu_0$, where

$$\nu_0 = \frac{\phi_0}{h} \quad (11.3)$$

Equation (11.3) shows that the greater the work function ϕ_0 , the higher the minimum or threshold frequency ν_0 needed to emit photoelectrons. Thus, there exists a threshold frequency $\nu_0 (= \phi_0/h)$ for the metal surface, below which no photoelectric emission is possible, no matter how intense the incident radiation may be or how long it falls on the surface.

- In this picture, intensity of radiation as noted above, is proportional to the number of energy quanta per unit area per unit time. The greater the number of energy quanta available, the greater is the number of electrons absorbing the energy quanta and greater, therefore, is the number of electrons coming out of the metal (for $\nu > \nu_0$). This explains why, for $\nu > \nu_0$, photoelectric current is proportional to intensity.
- In Einstein's picture, the basic elementary process involved in photoelectric effect is the absorption of a light quantum by an electron. This process is instantaneous. Thus, whatever may be the intensity i.e., the number of quanta of radiation per unit area per unit time, photoelectric emission is instantaneous. Low intensity does not mean delay in emission, since the basic elementary process is the same. Intensity only determines how many electrons are able to participate in the elementary process (absorption of a light quantum by a single electron) and, therefore, the photoelectric current.

Using Eq. (11.1), the photoelectric equation, Eq. (11.2), can be written as

$$eV_0 = h\nu - \phi_0; \text{ for } \nu \geq \nu_0$$

$$\text{or } V_0 = \frac{h}{e} \nu - \frac{\phi_0}{e} \quad (11.4)$$

This is an important result. It predicts that the V_0 versus ν curve is a straight line with slope $= (h/e)$,

independent of the nature of the material. During 1906-1916, Millikan performed a series of experiments on photoelectric effect, aimed at disproving Einstein's photoelectric equation. He measured the slope of the straight line obtained for sodium, similar to that shown in Fig. 11.5. Using the known value of e , he determined the value of Planck's constant h . This value was close to the value of Planck's constant ($= 6.626 \times 10^{-34} \text{ J s}$) determined in an entirely different context. In this way, in 1916, Millikan proved the validity of Einstein's photoelectric equation, instead of disproving it.

The successful explanation of photoelectric effect using the hypothesis of light quanta and the experimental determination of values of h and ϕ_0 , in agreement with values obtained from other experiments, led to the acceptance of Einstein's picture of photoelectric effect. Millikan verified photoelectric equation with great precision, for a number of alkali metals over a wide range of radiation frequencies.

11.7 PARTICLE NATURE OF LIGHT: THE PHOTON

Photoelectric effect thus gave evidence to the strange fact that light in interaction with matter behaved as if it was made of quanta or packets of energy, each of energy $h\nu$.

Is the light quantum of energy to be associated with a particle? Einstein arrived at the important result, that the light quantum can also be associated with momentum ($h\nu/c$). A definite value of energy as well as momentum is a strong sign that the light quantum can be associated with a particle. This particle was later named *photon*. The particle-like behaviour of light was further confirmed, in 1924, by the experiment of A.H. Compton (1892-1962) on scattering of X-rays from electrons. In 1921, Einstein was awarded the Nobel Prize in Physics for his contribution to theoretical physics and the photoelectric effect. In 1923, Millikan was awarded the Nobel Prize in physics for his work on the elementary charge of electricity and on the photoelectric effect.

We can summarise the photon picture of electromagnetic radiation as follows:

- (i) In interaction of radiation with matter, radiation behaves as if it is made up of particles called photons.
- (ii) Each photon has energy $E (=h\nu)$ and momentum $p (= h\nu/c)$, and speed c , the speed of light.
- (iii) All photons of light of a particular frequency ν , or wavelength λ , have the same energy $E (=h\nu = hc/\lambda)$ and momentum $p (= h\nu/c = h/\lambda)$, whatever the intensity of radiation may be. By increasing the intensity of light of given wavelength, there is only an increase in the number of photons per second crossing a given area, with each photon having the same energy. Thus, photon energy is independent of intensity of radiation.
- (iv) Photons are electrically neutral and are not deflected by electric and magnetic fields.
- (v) In a photon-particle collision (such as photon-electron collision), the total energy and total momentum are conserved. However, the number of photons may not be conserved in a collision. The photon may be absorbed or a new photon may be created.

Example 11.1 Monochromatic light of frequency 6.0×10^{14} Hz is produced by a laser. The power emitted is 2.0×10^{-3} W. (a) What is the energy of a photon in the light beam? (b) How many photons per second, on an average, are emitted by the source?

Solution

(a) Each photon has an energy

$$E = h\nu = (6.63 \times 10^{-34} \text{ J s}) (6.0 \times 10^{14} \text{ Hz}) \\ = 3.98 \times 10^{-19} \text{ J}$$

(b) If N is the number of photons emitted by the source per second, the power P transmitted in the beam equals N times the energy per photon E , so that $P = NE$. Then

$$N = \frac{P}{E} = \frac{2.0 \times 10^{-3} \text{ W}}{3.98 \times 10^{-19} \text{ J}} \\ = 5.0 \times 10^{15} \text{ photons per second.}$$

Example 11.2 The work function of caesium is 2.14 eV. Find (a) the threshold frequency for caesium, and (b) the wavelength of the incident light if the photocurrent is brought to zero by a stopping potential of 0.60 V.

Solution

(a) For the cut-off or threshold frequency, the energy $h\nu_0$ of the incident radiation must be equal to work function ϕ_0 , so that

$$\nu_0 = \frac{\phi_0}{h} = \frac{2.14 \text{ eV}}{6.63 \times 10^{-34} \text{ J s}} \\ = \frac{2.14 \times 1.6 \times 10^{-19} \text{ J}}{6.63 \times 10^{-34} \text{ J s}} = 5.16 \times 10^{14} \text{ Hz}$$

Thus, for frequencies less than this threshold frequency, no photoelectrons are ejected.

(b) Photocurrent reduces to zero, when maximum kinetic energy of the emitted photoelectrons equals the potential energy eV_0 by the retarding potential V_0 . Einstein's Photoelectric equation is

$$eV_0 = h\nu - \phi_0 = \frac{hc}{\lambda} - \phi_0 \\ \text{or, } \lambda = hc/(eV_0 + \phi_0) \\ = \frac{(6.63 \times 10^{-34} \text{ J s}) \times (3 \times 10^8 \text{ m/s})}{(0.60 \text{ eV} + 2.14 \text{ eV})} \\ = \frac{19.89 \times 10^{-26} \text{ J m}}{(2.74 \text{ eV})} \\ \lambda = \frac{19.89 \times 10^{-26} \text{ J m}}{2.74 \times 1.6 \times 10^{-19} \text{ J}} = 454 \text{ nm}$$

11.8 WAVE NATURE OF MATTER

The dual (wave-particle) nature of light (electromagnetic radiation, in general) comes out clearly from what we have learnt in this and the preceding chapters. The wave nature of light shows up in the phenomena of interference, diffraction and polarisation. On the other hand, in

photoelectric effect and Compton effect which involve energy and momentum transfer, radiation behaves as if it is made up of a bunch of particles – the photons. Whether a particle or wave description is best suited for understanding an experiment depends on the nature of the experiment. For example, in the familiar phenomenon of seeing an object by our eye, both descriptions are important. The gathering and focussing mechanism of light by the eye-lens is well described in the wave picture. But its absorption by the rods and cones (of the retina) requires the photon picture of light.

A natural question arises: If radiation has a dual (wave-particle) nature, might not the particles of nature (the electrons, protons, etc.) also exhibit wave-like character? In 1924, the French physicist Louis Victor de Broglie (pronounced as de Broy) (1892-1987) put forward the bold hypothesis that moving particles of matter should display wave-like properties under suitable conditions. He reasoned that nature was symmetrical and that the two basic physical entities – matter and energy, must have symmetrical character. If radiation shows dual aspects, so should matter. De Broglie proposed that the wave length λ associated with a particle of momentum p is given as

$$\lambda = \frac{h}{p} = \frac{h}{mv} \quad (11.5)$$

where m is the mass of the particle and v its speed. Equation (11.5) is known as the *de Broglie relation* and the wavelength λ of the *matter wave* is called *de Broglie wavelength*. The dual aspect of matter is evident in the de Broglie relation. On the left hand side of Eq. (11.5), λ is the attribute of a wave while on the right hand side the momentum p is a typical attribute of a particle. Planck's constant h relates the two attributes.

Equation (11.5) for a material particle is basically a hypothesis whose validity can be tested only by experiment. However, it is interesting to see that it is satisfied also by a photon. For a photon, as we have seen,

$$p = hv / c \quad (11.6)$$

Therefore,

$$\frac{h}{p} = \frac{c}{v} = \lambda \quad (11.7)$$

That is, the de Broglie wavelength of a photon given by Eq. (11.5) equals the wavelength of electromagnetic radiation of which the photon is a quantum of energy and momentum.

Clearly, from Eq. (11.5), λ is smaller for a heavier particle (large m) or more energetic particle (large v). For example, the de Broglie wavelength of a ball of mass 0.12 kg moving with a speed of 20 m s⁻¹ is easily calculated:



Louis Victor de Broglie (1892 – 1987) French physicist who put forth revolutionary idea of wave nature of matter. This idea was developed by Erwin Schrödinger into a full-fledged theory of quantum mechanics commonly known as wave mechanics. In 1929, he was awarded the Nobel Prize in Physics for his discovery of the wave nature of electrons.

LOUIS VICTOR DE BROGLIE (1892 – 1987)

$$p = m v = 0.12 \text{ kg} \times 20 \text{ m s}^{-1} = 2.40 \text{ kg m s}^{-1}$$

$$\lambda = \frac{h}{p} = \frac{6.63 \times 10^{-34} \text{ J s}}{2.40 \text{ kg m s}^{-1}} = 2.76 \times 10^{-34} \text{ m}$$

This wavelength is so small that it is beyond any measurement. This is the reason why macroscopic objects in our daily life do not show wave-like properties. On the other hand, in the sub-atomic domain, the wave character of particles is significant and measurable.

Example 11.3 What is the de Broglie wavelength associated with (a) an electron moving with a speed of $5.4 \times 10^6 \text{ m/s}$, and (b) a ball of mass 150 g travelling at 30.0 m/s?

Solution

(a) For the electron:

Mass $m = 9.11 \times 10^{-31} \text{ kg}$, speed $v = 5.4 \times 10^6 \text{ m/s}$. Then, momentum

$$p = m v = 9.11 \times 10^{-31} \text{ (kg)} \times 5.4 \times 10^6 \text{ (m/s)}$$

$$p = 4.92 \times 10^{-24} \text{ kg m/s}$$

de Broglie wavelength, $\lambda = h/p$

$$= \frac{6.63 \times 10^{-34} \text{ J s}}{4.92 \times 10^{-24} \text{ kg m/s}}$$

$$\lambda = 0.135 \text{ nm}$$

(b) For the ball:

Mass $m' = 0.150 \text{ kg}$, speed $v' = 30.0 \text{ m/s}$.

Then momentum $p' = m' v' = 0.150 \text{ (kg)} \times 30.0 \text{ (m/s)}$

$$p' = 4.50 \text{ kg m/s}$$

de Broglie wavelength $\lambda' = h/p'$.

$$= \frac{6.63 \times 10^{-34} \text{ J s}}{4.50 \times \text{kg m/s}}$$

$$\lambda' = 1.47 \times 10^{-34} \text{ m}$$

The de Broglie wavelength of electron is comparable with X-ray wavelengths. However, for the ball it is about 10^{-19} times the size of the proton, quite beyond experimental measurement.

SUMMARY

1. The minimum energy needed by an electron to come out from a metal surface is called the work function of the metal. Energy (greater than the work function (ϕ_0)) required for electron emission from the metal surface can be supplied by suitably heating or applying strong electric field or irradiating it by light of suitable frequency.
2. Photoelectric effect is the phenomenon of emission of electrons by metals when illuminated by light of suitable frequency. Certain metals respond to ultraviolet light while others are sensitive even to the visible light. Photoelectric effect involves conversion of light energy into electrical energy. It follows the law of conservation of energy. The photoelectric emission is an instantaneous process and possesses certain special features.

3. Photoelectric current depends on (i) the intensity of incident light, (ii) the potential difference applied between the two electrodes, and (iii) the nature of the emitter material.
4. The stopping potential (V_0) depends on (i) the frequency of incident light, and (ii) the nature of the emitter material. For a given frequency of incident light, it is independent of its intensity. The stopping potential is directly related to the maximum kinetic energy of electrons emitted:

$$e V_0 = (1/2) m v_{max}^2 = K_{max}$$
5. Below a certain frequency (threshold frequency) ν_0 , characteristic of the metal, no photoelectric emission takes place, no matter how large the intensity may be.
6. The classical wave theory could not explain the main features of photoelectric effect. Its picture of continuous absorption of energy from radiation could not explain the independence of K_{max} on intensity, the existence of ν_0 and the instantaneous nature of the process. Einstein explained these features on the basis of photon picture of light. According to this, light is composed of discrete packets of energy called quanta or photons. Each photon carries an energy $E (= h \nu)$ and momentum $p (= h/\lambda)$, which depend on the frequency (ν) of incident light and not on its intensity. Photoelectric emission from the metal surface occurs due to absorption of a photon by an electron.
7. Einstein's photoelectric equation is in accordance with the energy conservation law as applied to the photon absorption by an electron in the metal. The maximum kinetic energy $(1/2)m v_{max}^2$ is equal to the photon energy ($h\nu$) minus the work function $\phi_0 (= h\nu_0)$ of the target metal:

$$\frac{1}{2} m v_{max}^2 = V_0 e = h\nu - \phi_0 = h(\nu - \nu_0)$$

This photoelectric equation explains all the features of the photoelectric effect. Millikan's first precise measurements confirmed the Einstein's photoelectric equation and obtained an accurate value of Planck's constant h . This led to the acceptance of particle or photon description (nature) of electromagnetic radiation, introduced by Einstein.

8. Radiation has dual nature: wave and particle. The nature of experiment determines whether a wave or particle description is best suited for understanding the experimental result. Reasoning that radiation and matter should be symmetrical in nature, Louis Victor de Broglie attributed a wave-like character to matter (material particles). The waves associated with the moving material particles are called matter waves or de Broglie waves.
9. The de Broglie wavelength (λ) associated with a moving particle is related to its momentum p as: $\lambda = h/p$. The dualism of matter is inherent in the de Broglie relation which contains a wave concept (λ) and a particle concept (p). The de Broglie wavelength is independent of the charge and nature of the material particle. It is significantly measurable (of the order of the atomic-planes spacing in crystals) only in case of sub-atomic particles like electrons, protons, etc. (due to smallness of their masses and hence, momenta). However, it is indeed very small, quite beyond measurement, in case of macroscopic objects, commonly encountered in everyday life.

Physical Quantity	Symbol	Dimensions	Unit	Remarks
Planck's constant	h	$[ML^2T^{-1}]$	J s	$E = h\nu$
Stopping potential	V_0	$[ML^2T^{-3}A^{-1}]$	V	$eV_0 = K_{\max}$
Work function	ϕ_0	$[ML^2T^{-2}]$	J; eV	$K_{\max} = E - \phi_0$
Threshold frequency	ν_0	$[T^{-1}]$	Hz	$\nu_0 = \phi_0 / h$
de Broglie wavelength	λ	[L]	m	$\lambda = h/p$

POINTS TO PONDER

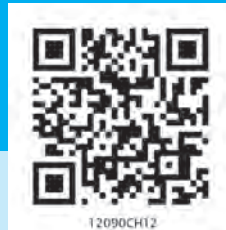
- Free electrons in a metal are free in the sense that they move inside the metal in a constant potential (This is only an approximation). They are not free to move out of the metal. They need additional energy to get out of the metal.
- Free electrons in a metal do not all have the same energy. Like molecules in a gas jar, the electrons have a certain energy distribution at a given temperature. This distribution is different from the usual Maxwell's distribution that you have learnt in the study of kinetic theory of gases. You will learn about it in later courses, but the difference has to do with the fact that electrons obey Pauli's exclusion principle.
- Because of the energy distribution of free electrons in a metal, the energy required by an electron to come out of the metal is different for different electrons. Electrons with higher energy require less additional energy to come out of the metal than those with lower energies. Work function is the least energy required by an electron to come out of the metal.
- Observations on photoelectric effect imply that in the event of matter-light interaction, *absorption of energy takes place in discrete units of $h\nu$* . This is not quite the same as saying that light consists of particles, each of energy $h\nu$.
- Observations on the stopping potential (its independence of intensity and dependence on frequency) are the crucial discriminator between the wave-picture and photon-picture of photoelectric effect.
- The wavelength of a matter wave given by $\lambda = \frac{h}{p}$ has physical significance; its phase velocity v_p has no physical significance. However, the group velocity of the matter wave is physically meaningful and equals the velocity of the particle.

EXERCISES

11.1 Find the

- maximum frequency, and
- minimum wavelength of X-rays produced by 30 kV electrons.

- 11.2** The work function of caesium metal is 2.14 eV. When light of frequency 6×10^{14} Hz is incident on the metal surface, photoemission of electrons occurs. What is the
- (a) maximum kinetic energy of the emitted electrons,
 - (b) Stopping potential, and
 - (c) maximum speed of the emitted photoelectrons?
- 11.3** The photoelectric cut-off voltage in a certain experiment is 1.5 V. What is the maximum kinetic energy of photoelectrons emitted?
- 11.4** Monochromatic light of wavelength 632.8 nm is produced by a helium-neon laser. The power emitted is 9.42 mW.
- (a) Find the energy and momentum of each photon in the light beam,
 - (b) How many photons per second, on the average, arrive at a target irradiated by this beam? (Assume the beam to have uniform cross-section which is less than the target area), and
 - (c) How fast does a hydrogen atom have to travel in order to have the same momentum as that of the photon?
- 11.5** In an experiment on photoelectric effect, the slope of the cut-off voltage versus frequency of incident light is found to be 4.12×10^{-15} V s. Calculate the value of Planck's constant.
- 11.6** The threshold frequency for a certain metal is 3.3×10^{14} Hz. If light of frequency 8.2×10^{14} Hz is incident on the metal, predict the cut-off voltage for the photoelectric emission.
- 11.7** The work function for a certain metal is 4.2 eV. Will this metal give photoelectric emission for incident radiation of wavelength 330 nm?
- 11.8** Light of frequency 7.21×10^{14} Hz is incident on a metal surface. Electrons with a maximum speed of 6.0×10^5 m/s are ejected from the surface. What is the threshold frequency for photoemission of electrons?
- 11.9** Light of wavelength 488 nm is produced by an argon laser which is used in the photoelectric effect. When light from this spectral line is incident on the emitter, the stopping (cut-off) potential of photoelectrons is 0.38 V. Find the work function of the material from which the emitter is made.
- 11.10** What is the de Broglie wavelength of
- (a) a bullet of mass 0.040 kg travelling at the speed of 1.0 km/s,
 - (b) a ball of mass 0.060 kg moving at a speed of 1.0 m/s, and
 - (c) a dust particle of mass 1.0×10^{-9} kg drifting with a speed of 2.2 m/s?
- 11.11** Show that the wavelength of electromagnetic radiation is equal to the de Broglie wavelength of its quantum (photon).



Chapter Twelve

ATOMS

12.1 INTRODUCTION

By the nineteenth century, enough evidence had accumulated in favour of atomic hypothesis of matter. In 1897, the experiments on electric discharge through gases carried out by the English physicist J. J. Thomson (1856 – 1940) revealed that atoms of different elements contain negatively charged constituents (electrons) that are identical for all atoms. However, atoms on a whole are electrically neutral. Therefore, an atom must also contain some positive charge to neutralise the negative charge of the electrons. But what is the arrangement of the positive charge and the electrons inside the atom? In other words, what is the structure of an atom?

The first model of atom was proposed by J. J. Thomson in 1898. According to this model, the positive charge of the atom is uniformly distributed throughout the volume of the atom and the negatively charged electrons are embedded in it like seeds in a watermelon. This model was picturesquely called *plum pudding model* of the atom. However subsequent studies on atoms, as described in this chapter, showed that the distribution of the electrons and positive charges are very different from that proposed in this model.

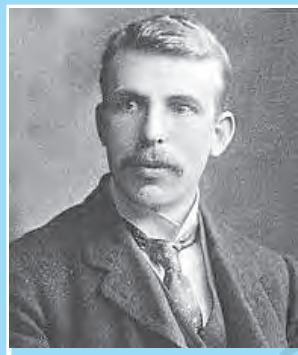
We know that condensed matter (solids and liquids) and dense gases at all temperatures emit electromagnetic radiation in which a continuous distribution of several wavelengths is present, though with different intensities. This radiation is considered to be due to oscillations of atoms

and molecules, governed by the interaction of each atom or molecule with its neighbours. *In contrast*, light emitted from rarefied gases heated in a flame, or excited electrically in a glow tube such as the familiar neon sign or mercury vapour light has only certain discrete wavelengths. The spectrum appears as a series of bright lines. In such gases, the average spacing between atoms is large. Hence, the radiation emitted can be considered due to individual atoms rather than because of interactions between atoms or molecules.

In the early nineteenth century it was also established that each element is associated with a characteristic spectrum of radiation, for example, hydrogen always gives a set of lines with fixed relative position between the lines. This fact suggested an intimate relationship between the internal structure of an atom and the spectrum of radiation emitted by it. In 1885, Johann Jakob Balmer (1825 – 1898) obtained a simple empirical formula which gave the wavelengths of a group of lines emitted by atomic hydrogen. Since hydrogen is simplest of the elements known, we shall consider its spectrum in detail in this chapter.

Ernst Rutherford (1871–1937), a former research student of J. J. Thomson, was engaged in experiments on α -particles emitted by some radioactive elements. In 1906, he proposed a classic experiment of scattering of these α -particles by atoms to investigate the atomic structure. This experiment was later performed around 1911 by Hans Geiger (1882–1945) and Ernst Marsden (1889–1970, who was 20 year-old student and had not yet earned his bachelor's degree). The details are discussed in Section 12.2. The explanation of the results led to the birth of Rutherford's planetary model of atom (also called the *nuclear model of the atom*). According to this the entire positive charge and most of the mass of the atom is concentrated in a small volume called the nucleus with electrons revolving around the nucleus just as planets revolve around the sun.

Rutherford's nuclear model was a major step towards how we see the atom today. However, it could not explain why atoms emit light of only discrete wavelengths. How could an atom as simple as hydrogen, consisting of a single electron and a single proton, emit a complex spectrum of specific wavelengths? In the classical picture of an atom, the electron revolves round the nucleus much like the way a planet revolves round the sun. However, we shall see that there are some serious difficulties in accepting such a model.



Ernst Rutherford (1871 – 1937) New Zealand born, British physicist who did pioneering work on radioactive radiation. He discovered alpha-rays and beta-rays. Along with Frederick Soddy, he created the modern theory of radioactivity. He studied the 'emanation' of thorium and discovered a new noble gas, an isotope of radon, now known as thoron. By scattering alpha-rays from the metal foils, he discovered the atomic nucleus and proposed the planetary model of the atom. He also estimated the approximate size of the nucleus.

ERNST RUTHERFORD (1871 – 1937)

12.2 ALPHA-PARTICLE SCATTERING AND RUTHERFORD'S NUCLEAR MODEL OF ATOM

At the suggestion of Ernst Rutherford, in 1911, H. Geiger and E. Marsden performed some experiments. In one of their experiments, as shown in

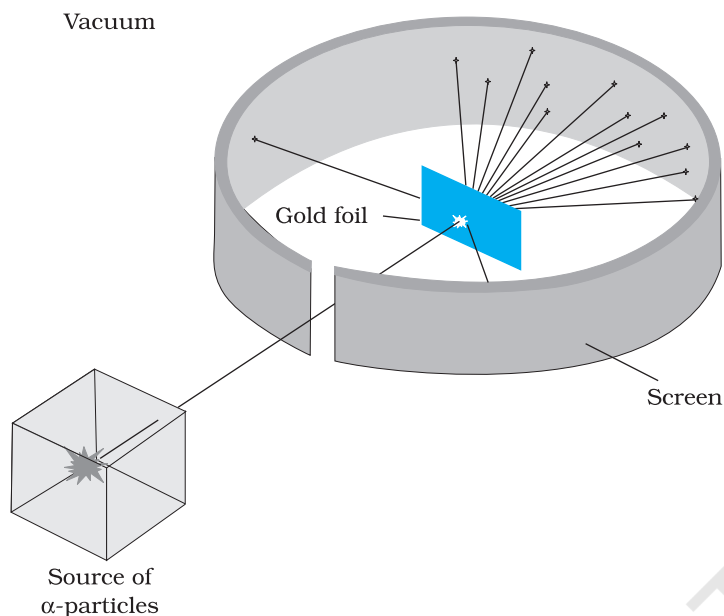


FIGURE 12.1 Geiger-Marsden scattering experiment. The entire apparatus is placed in a vacuum chamber (not shown in this figure).

Fig. 12.1, they directed a beam of 5.5 MeV α -particles emitted from a $^{214}_{83}\text{Bi}$ radioactive source at a thin metal foil made of gold. Figure 12.2 shows a schematic diagram of this experiment. Alpha-particles emitted by a $^{214}_{83}\text{Bi}$ radioactive source were collimated into a narrow beam by their passage through lead bricks. The beam was allowed to fall on a thin foil of gold of thickness 2.1×10^{-7} m. The scattered alpha-particles were observed through a rotatable detector consisting of zinc sulphide screen and a microscope. The scattered alpha-particles on striking the screen produced brief light flashes or scintillations. These flashes may be viewed through a microscope and the distribution of the number of scattered particles may be studied as a function of angle of scattering.

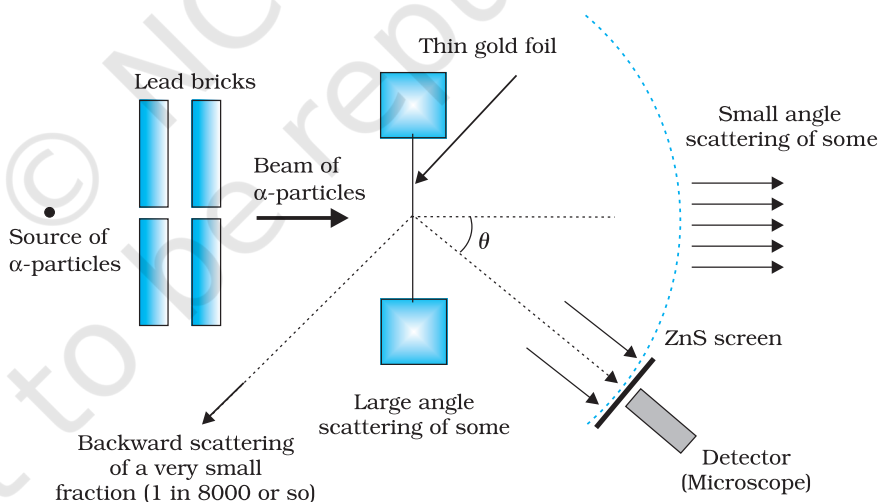


FIGURE 12.2 Schematic arrangement of the Geiger-Marsden experiment.

A typical graph of the total number of α -particles scattered at different angles, in a given interval of time, is shown in Fig. 12.3. The dots in this figure represent the data points and the solid curve is the theoretical prediction based on the assumption that the target atom has a small, dense, positively charged nucleus. Many of the α -particles pass through the foil. It means that they do not suffer any collisions. Only about 0.14% of the incident α -particles scatter by more than 1° ; and about 1 in 8000 deflect by more than 90° . Rutherford argued that, to deflect the α -particle backwards, it must experience a large repulsive force. This force could

be provided if the greater part of the mass of the atom and its positive charge were concentrated tightly at its centre. Then the incoming α -particle could get very close to the positive charge without penetrating it, and such a close encounter would result in a large deflection. This agreement supported the hypothesis of the nuclear atom. This is why Rutherford is credited with the *discovery* of the nucleus.

In Rutherford's nuclear model of the atom, the entire positive charge and most of the mass of the atom are concentrated in the nucleus with the electrons some distance away. The electrons would be moving in orbits about the nucleus just as the planets do around the sun. Rutherford's experiments suggested the size of the nucleus to be about 10^{-15} m to 10^{-14} m. From kinetic theory, the size of an atom was known to be 10^{-10} m, about 10,000 to 100,000 times larger than the size of the nucleus. Thus, the electrons would seem to be at a distance from the nucleus of about 10,000 to 100,000 times the size of the nucleus itself. Thus, most of an atom is empty space. With the atom being largely empty space, it is easy to see why most α -particles go right through a thin metal foil. However, when α -particle happens to come near a nucleus, the intense electric field there scatters it through a large angle. The atomic electrons, being so light, do not appreciably affect the α -particles.

The scattering data shown in Fig. 12.3 can be analysed by employing Rutherford's nuclear model of the atom. As the gold foil is very thin, it can be assumed that α -particles will suffer not more than one scattering during their passage through it. Therefore, computation of the trajectory of an alpha-particle scattered by a single nucleus is enough. Alpha-particles are nuclei of helium atoms and, therefore, carry two units, $2e$, of positive charge and have the mass of the helium atom. The charge of the gold nucleus is Ze , where Z is the atomic number of the atom; for gold $Z = 79$. Since the nucleus of gold is about 50 times heavier than an α -particle, it is reasonable to assume that it remains stationary throughout the scattering process. Under these assumptions, the trajectory of an alpha-particle can be computed employing Newton's second law of motion and the Coulomb's law for electrostatic force of repulsion between the alpha-particle and the positively charged nucleus.

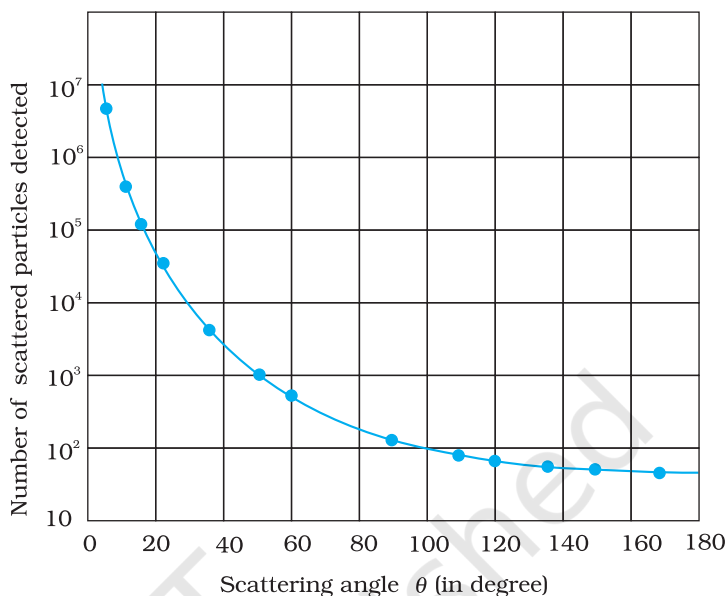


FIGURE 12.3 Experimental data points (shown by dots) on scattering of α -particles by a thin foil at different angles obtained by Geiger and Marsden using the setup shown in Figs. 12.1 and 12.2. Rutherford's nuclear model predicts the solid curve which is seen to be in good agreement with experiment.

The magnitude of this force is

$$F = \frac{1}{4\pi\epsilon_0} \frac{(2e)(Ze)}{r^2} \quad (12.1)$$

where r is the distance between the α -particle and the nucleus. The force is directed along the line joining the α -particle and the nucleus. The magnitude and direction of the force on an α -particle continuously changes as it approaches the nucleus and recedes away from it.

12.2.1 Alpha-particle trajectory

The trajectory traced by an α -particle depends on the impact parameter, b of collision. The *impact parameter* is the perpendicular distance of the initial velocity vector of the α -particle from the centre of the nucleus (Fig.

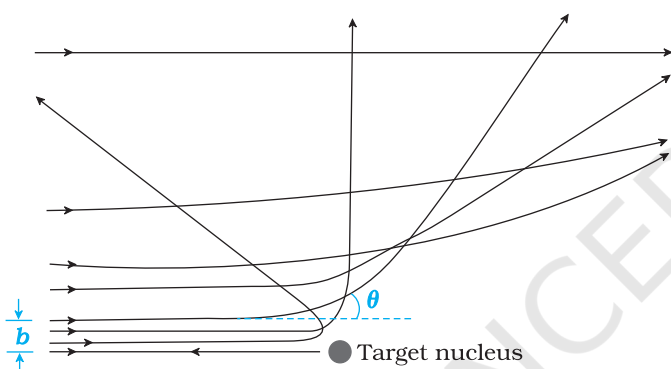


FIGURE 12.4 Trajectory of α -particles in the coulomb field of a target nucleus. The impact parameter, b and scattering angle θ are also depicted.

12.4). A given beam of α -particles has a distribution of impact parameters b , so that the beam is scattered in various directions with different probabilities (Fig. 12.4). (In a beam, all particles have nearly same kinetic energy.) It is seen that an α -particle close to the nucleus (small impact parameter) suffers large scattering. In case of head-on collision, the impact parameter is minimum and the α -particle rebounds back ($\theta \cong \pi$). For a large impact parameter, the α -particle goes nearly undeviated and has a small deflection ($\theta \cong 0$).

The fact that only a small fraction of the number of incident particles rebound back indicates that the number of α -particles undergoing head on collision is small. This,

in turn, implies that the mass and positive charge of the atom is concentrated in a small volume. Rutherford scattering therefore, is a powerful way to determine an upper limit to the size of the nucleus.

Example 12.1 In the Rutherford's nuclear model of the atom, the nucleus (radius about 10^{-15} m) is analogous to the sun about which the electron move in orbit (radius $\approx 10^{-10}$ m) like the earth orbits around the sun. If the dimensions of the solar system had the same proportions as those of the atom, would the earth be closer to or farther away from the sun than actually it is? The radius of earth's orbit is about 1.5×10^{11} m. The radius of sun is taken as 7×10^8 m.

Solution The ratio of the radius of electron's orbit to the radius of nucleus is $(10^{-10} \text{ m})/(10^{-15} \text{ m}) = 10^5$, that is, the radius of the electron's orbit is 10^5 times larger than the radius of nucleus. If the radius of the earth's orbit around the sun were 10^5 times larger than the radius of the sun, the radius of the earth's orbit would be $10^5 \times 7 \times 10^8 \text{ m} = 7 \times 10^{13} \text{ m}$. This is more than 100 times greater than the actual orbital radius of earth. Thus, the earth would be much farther away from the sun.

It implies that an atom contains a much greater fraction of empty space than our solar system does.

Example 12.2 In a Geiger-Marsden experiment, what is the distance of closest approach to the nucleus of a 7.7 MeV α -particle before it comes momentarily to rest and reverses its direction?

Solution The key idea here is that throughout the scattering process, the total mechanical energy of the system consisting of an α -particle and a gold nucleus is conserved. The system's initial mechanical energy is E_i , before the particle and nucleus interact, and it is equal to its mechanical energy E_f when the α -particle momentarily stops. The initial energy E_i is just the kinetic energy K of the incoming α -particle. The final energy E_f is just the electric potential energy U of the system. The potential energy U can be calculated from Eq. (12.1).

Let d be the centre-to-centre distance between the α -particle and the gold nucleus when the α -particle is at its stopping point. Then we can write the conservation of energy $E_i = E_f$ as

$$K = \frac{1}{4\pi\epsilon_0} \frac{(2e)(Ze)}{d} = \frac{2Ze^2}{4\pi\epsilon_0 d}$$

Thus the distance of closest approach d is given by

$$d = \frac{2Ze^2}{4\pi\epsilon_0 K}$$

The maximum kinetic energy found in α -particles of natural origin is 7.7 MeV or 1.2×10^{-12} J. Since $1/4\pi\epsilon_0 = 9.0 \times 10^9 \text{ N m}^2/\text{C}^2$. Therefore with $e = 1.6 \times 10^{-19} \text{ C}$, we have,

$$d = \frac{(2)(9.0 \times 10^9 \text{ Nm}^2/\text{C}^2)(1.6 \times 10^{-19} \text{ C})^2 Z}{1.2 \times 10^{-12} \text{ J}}$$

$$= 3.84 \times 10^{-16} Z \text{ m}$$

The atomic number of foil material gold is $Z = 79$, so that

$$d(\text{Au}) = 3.0 \times 10^{-14} \text{ m} = 30 \text{ fm. (1 fm (i.e. fermi) = } 10^{-15} \text{ m.)}$$

The radius of gold nucleus is, therefore, less than $3.0 \times 10^{-14} \text{ m}$. This is not in very good agreement with the observed result as the actual radius of gold nucleus is 6 fm. The cause of discrepancy is that the distance of closest approach is considerably larger than the sum of the radii of the gold nucleus and the α -particle. Thus, the α -particle reverses its motion without ever actually touching the gold nucleus.

EXAMPLE 12.2

12.2.2 Electron orbits

The Rutherford nuclear model of the atom which involves classical concepts, pictures the atom as an electrically neutral sphere consisting of a very small, massive and positively charged nucleus at the centre surrounded by the revolving electrons in their respective dynamically stable orbits. The electrostatic force of attraction, F_e between the revolving electrons and the nucleus provides the requisite centripetal force (F_c) to keep them in their orbits. Thus, for a dynamically stable orbit in a hydrogen atom

$$F_e = F_c$$

$$\frac{1}{4\pi\epsilon_0} \frac{e^2}{r^2} = \frac{mv^2}{r} \quad (12.2)$$

Thus the relation between the orbit radius and the electron velocity is

$$r = \frac{e^2}{4\pi\epsilon_0 mv^2} \quad (12.3)$$

The kinetic energy (K) and electrostatic potential energy (U) of the electron in hydrogen atom are

$$K = \frac{1}{2}mv^2 = \frac{e^2}{8\pi\epsilon_0 r} \text{ and } U = -\frac{e^2}{4\pi\epsilon_0 r}$$

(The negative sign in U signifies that the electrostatic force is in the $-r$ direction.) Thus the total energy E of the electron in a hydrogen atom is

$$\begin{aligned} E = K + U &= \frac{e^2}{8\pi\epsilon_0 r} - \frac{e^2}{4\pi\epsilon_0 r} \\ &= -\frac{e^2}{8\pi\epsilon_0 r} \end{aligned} \quad (12.4)$$

The total energy of the electron is negative. This implies the fact that the electron is bound to the nucleus. If E were positive, an electron will not follow a closed orbit around the nucleus.

EXAMPLE 12.3

Example 12.3 It is found experimentally that 13.6 eV energy is required to separate a hydrogen atom into a proton and an electron. Compute the orbital radius and the velocity of the electron in a hydrogen atom.

Solution Total energy of the electron in hydrogen atom is $-13.6 \text{ eV} = -13.6 \times 1.6 \times 10^{-19} \text{ J} = -2.2 \times 10^{-18} \text{ J}$. Thus from Eq. (12.4), we have

$$E = -\frac{e^2}{8\pi\epsilon_0 r} = -2.2 \times 10^{-18} \text{ J}$$

This gives the orbital radius

$$\begin{aligned} r &= -\frac{e^2}{8\pi\epsilon_0 E} = -\frac{(9 \times 10^9 \text{ N m}^2/\text{C}^2)(1.6 \times 10^{-19} \text{ C})^2}{(2)(-2.2 \times 10^{-18} \text{ J})} \\ &= 5.3 \times 10^{-11} \text{ m.} \end{aligned}$$

The velocity of the revolving electron can be computed from Eq. (12.3) with $m = 9.1 \times 10^{-31} \text{ kg}$,

$$v = \frac{e}{\sqrt{4\pi\epsilon_0 mr}} = 2.2 \times 10^6 \text{ m/s.}$$

12.3 ATOMIC SPECTRA

As mentioned in Section 12.1, each element has a characteristic spectrum of radiation, which it emits. When an atomic gas or vapour is excited at low pressure, usually by passing an electric current through it, the emitted radiation has a spectrum which contains certain specific wavelengths only. A spectrum of this kind is termed as emission line spectrum and it

consists of bright lines on a dark background. The spectrum emitted by atomic hydrogen is shown in Fig. 12.5. Study of emission line spectra of a material can therefore serve as a type of “fingerprint” for identification of the gas. When white light passes through a gas and we analyse the transmitted light using a spectrometer we find some dark lines in the spectrum. These dark lines correspond precisely to those wavelengths which were found in the emission line spectrum of the gas. This is called the *absorption spectrum* of the material of the gas.

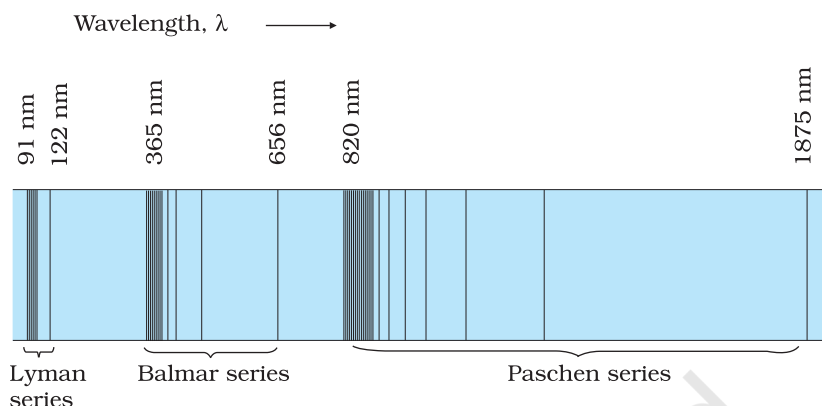
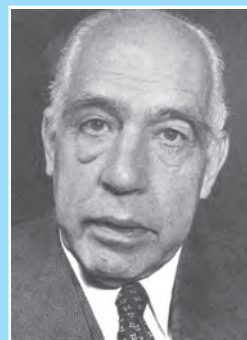


FIGURE 12.5 Emission lines in the spectrum of hydrogen.

12.4 BOHR MODEL OF THE HYDROGEN ATOM

The model of the atom proposed by Rutherford assumes that the atom, consisting of a central nucleus and revolving electron is stable much like sun-planet system which the model imitates. However, there are some fundamental differences between the two situations. While the planetary system is held by gravitational force, the nucleus-electron system being charged objects, interact by Coulomb's Law of force. We know that an object which moves in a circle is being constantly accelerated – the acceleration being centripetal in nature. According to classical electromagnetic theory, an accelerating charged particle emits radiation in the form of electromagnetic waves. The energy of an accelerating electron should therefore, continuously decrease. The electron would spiral inward and eventually fall into the nucleus (Fig. 12.6). Thus, such an atom can not be stable. Further, according to the classical electromagnetic theory, the frequency of the electromagnetic waves emitted by the revolving electrons is equal to the frequency of revolution. As the electrons spiral inwards, their angular velocities and hence their frequencies would change continuously, and so will the frequency of the light emitted. Thus, they would emit a continuous spectrum, in contradiction to the line spectrum actually observed. Clearly Rutherford model tells only a part of the story implying that the classical ideas are not sufficient to explain the atomic structure.



Niels Henrik David Bohr (1885 – 1962) Danish physicist who explained the spectrum of hydrogen atom based on quantum ideas. He gave a theory of nuclear fission based on the liquid-drop model of nucleus. Bohr contributed to the clarification of conceptual problems in quantum mechanics, in particular by proposing the complementary principle.

NIELS HENRIK DAVID BOHR (1885 – 1962)

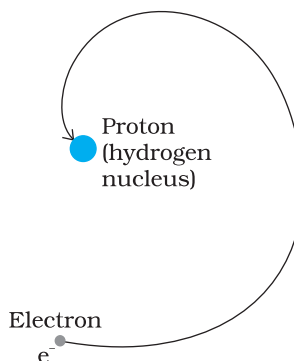


FIGURE 12.6 An accelerated atomic electron must spiral into the nucleus as it loses energy.

EXAMPLE 12.4

Example 12.4 According to the classical electromagnetic theory, calculate the initial frequency of the light emitted by the electron revolving around a proton in hydrogen atom.

Solution From Example 12.3 we know that velocity of electron moving around a proton in hydrogen atom in an orbit of radius 5.3×10^{-11} m is 2.2×10^6 m/s. Thus, the frequency of the electron moving around the proton is

$$\nu = \frac{v}{2\pi r} = \frac{2.2 \times 10^6 \text{ m s}^{-1}}{2\pi(5.3 \times 10^{-11} \text{ m})} \approx 6.6 \times 10^{15} \text{ Hz.}$$

According to the classical electromagnetic theory we know that the frequency of the electromagnetic waves emitted by the revolving electrons is equal to the frequency of its revolution around the nucleus. Thus the initial frequency of the light emitted is 6.6×10^{15} Hz.

It was Niels Bohr (1885 – 1962) who made certain modifications in this model by adding the ideas of the newly developing quantum hypothesis. Niels Bohr studied in Rutherford's laboratory for several months in 1912 and he was convinced about the validity of Rutherford nuclear model. Faced with the dilemma as discussed above, Bohr, in 1913, concluded that in spite of the success of electromagnetic theory in explaining large-scale phenomena, it could not be applied to the processes at the atomic scale. It became clear that a fairly radical departure from the established principles of classical mechanics and electromagnetism would be needed to understand the structure of atoms and the relation of atomic structure to atomic spectra. Bohr combined classical and early quantum concepts and gave his theory in the form of three postulates. These are :

- (i) Bohr's first postulate was that *an electron in an atom could revolve in certain stable orbits without the emission of radiant energy*, contrary to the predictions of electromagnetic theory. According to this postulate, each atom has certain definite stable states in which it

can exist, and each possible state has definite total energy. These are called the stationary states of the atom.

- (ii) Bohr's second postulate defines these stable orbits. This postulate states that the *electron* revolves around the nucleus *only in those orbits for which the angular momentum is some integral multiple of $h/2\pi$* where h is the Planck's constant ($= 6.6 \times 10^{-34}$ J s). Thus the angular momentum (L) of the orbiting electron is quantised. That is

$$L = nh/2\pi \quad (12.5)$$

- (iii) Bohr's third postulate incorporated into atomic theory the early quantum concepts that had been developed by Planck and Einstein. It states that *an electron might make a transition from one of its specified non-radiating orbits to another of lower energy. When it does so, a photon is emitted having energy equal to the energy difference between the initial and final states. The frequency of the emitted photon is then given by*

$$h\nu = E_i - E_f \quad (12.6)$$

where E_i and E_f are the energies of the initial and final states and $E_i > E_f$.

For a hydrogen atom, Eq. (12.4) gives the expression to determine the energies of different energy states. But then this equation requires the radius r of the electron orbit. To calculate r , Bohr's second postulate about the angular momentum of the electron—the quantisation condition – is used.

The radius of n^{th} possible orbit thus found is

$$r_n = \left(\frac{n^2}{m}\right) \left(\frac{h}{2\pi}\right)^2 \frac{4\pi\epsilon_0}{e^2} \quad (12.7)$$

The total energy of the electron in the stationary states of the hydrogen atom can be obtained by substituting the value of orbital radius in Eq. (12.4) as

$$E_n = -\left(\frac{e^2}{8\pi\epsilon_0}\right) \left(\frac{m}{n^2}\right) \left(\frac{2\pi}{h}\right)^2 \left(\frac{e^2}{4\pi\epsilon_0}\right)$$

$$\text{or } E_n = -\frac{me^4}{8n^2\epsilon_0^2h^2} \quad (12.8)$$

Substituting values, Eq. (12.8) yields

$$E_n = -\frac{2.18 \times 10^{-18}}{n^2} \text{ J} \quad (12.9)$$

Atomic energies are often expressed in electron volts (eV) rather than joules. Since $1 \text{ eV} = 1.6 \times 10^{-19} \text{ J}$, Eq. (12.9) can be rewritten as

$$E_n = -\frac{13.6}{n^2} \text{ eV} \quad (12.10)$$

The negative sign of the total energy of an electron moving in an orbit means that the electron is bound with the nucleus. Energy will thus be required to remove the electron from the hydrogen atom to a distance infinitely far away from its nucleus (or proton in hydrogen atom).

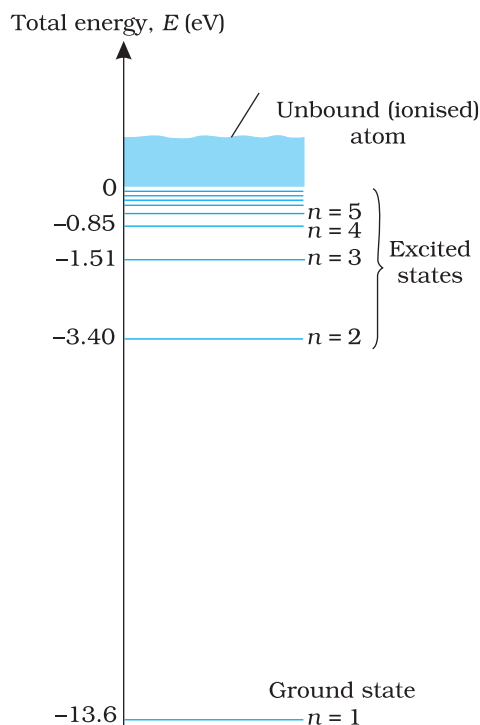


FIGURE 12.7 The energy level diagram for the hydrogen atom. The electron in a hydrogen atom at room temperature spends most of its time in the ground state. To ionise a hydrogen atom an electron from the ground state, 13.6 eV of energy must be supplied. (The horizontal lines specify the presence of allowed energy states.)

12.4.1 Energy levels

The energy of an atom is the *least* (largest negative value) when its electron is revolving in an orbit closest to the nucleus i.e., the one for which $n = 1$. For $n = 2, 3, \dots$ the absolute value of the energy E is smaller, hence the energy is progressively larger in the outer orbits. The *lowest* state of the atom, called the *ground state*, is that of the lowest energy, with the electron revolving in the orbit of smallest radius, the Bohr radius, a_0 . The energy of this state ($n = 1$), E_1 is -13.6 eV. Therefore, the minimum energy required to free the electron from the ground state of the hydrogen atom is 13.6 eV. It is called the *ionisation energy* of the hydrogen atom. This prediction of the Bohr's model is in excellent agreement with the experimental value of ionisation energy.

At room temperature, most of the hydrogen atoms are in *ground state*. When a hydrogen atom receives energy by processes such as electron collisions, the atom may acquire sufficient energy to raise the electron to higher energy states. The atom is then said to be in an *excited state*. From Eq. (12.10), for $n = 2$; the energy E_2 is -3.40 eV. It means that the energy required to excite an electron in hydrogen atom to its first excited state, is an energy equal to $E_2 - E_1 = -3.40 \text{ eV} - (-13.6) \text{ eV} = 10.2 \text{ eV}$. Similarly, $E_3 = -1.51$ eV and $E_3 - E_1 = 12.09$ eV, or to excite the hydrogen atom from its ground state ($n = 1$) to second excited state ($n = 3$), 12.09 eV energy is required, and so on. From these excited states the electron can then fall back to a state of lower energy, emitting a photon in the process. Thus, as the excitation of hydrogen atom increases (that is as n increases) the value of minimum energy required to free the electron from the excited atom decreases.

The energy level diagram* for the stationary states of a hydrogen atom, computed from Eq. (12.10), is given in Fig. 12.7. The principal quantum number n labels the stationary states in the ascending order of energy. In this diagram, the highest energy state corresponds to $n = \infty$ in Eq. (12.10) and has an energy of 0 eV. This is the energy of the atom when the electron is completely removed ($r = \infty$) from the nucleus and is at rest. Observe how the energies of the excited states come closer and closer together as n increases.

12.5 THE LINE SPECTRA OF THE HYDROGEN ATOM

According to the third postulate of Bohr's model, when an atom makes a transition from the higher energy state with quantum number n_i to the lower energy state with quantum number n_f ($n_f < n_i$), the difference of energy is carried away by a photon of frequency ν_{if} such that

* An electron can have any total energy above $E = 0$ eV. In such situations the electron is free. Thus there is a continuum of energy states above $E = 0$ eV, as shown in Fig. 12.7.

$$h\nu_{if} = E_{n_i} - E_{n_f} \quad (12.11)$$

Since both n_f and n_i are integers, this immediately shows that in transitions between different atomic levels, light is radiated in various discrete frequencies.

The various lines in the atomic spectra are produced when electrons jump from higher energy state to a lower energy state and photons are emitted. These spectral lines are called emission lines. But when an atom absorbs a photon that has precisely the same energy needed by the electron in a lower energy state to make transitions to a higher energy state, the process is called absorption. Thus if photons with a continuous range of frequencies pass through a rarefied gas and then are analysed with a spectrometer, a series of dark spectral absorption lines appear in the continuous spectrum. The dark lines indicate the frequencies that have been absorbed by the atoms of the gas.

The explanation of the hydrogen atom spectrum provided by Bohr's model was a brilliant achievement, which greatly stimulated progress towards the modern quantum theory. In 1922, Bohr was awarded Nobel Prize in Physics.

12.6 DE BROGLIE'S EXPLANATION OF BOHR'S SECOND POSTULATE OF QUANTISATION

Of all the postulates, Bohr made in his model of the atom, perhaps the most puzzling is his second postulate. It states that the angular momentum of the electron orbiting around the nucleus is quantised (that is, $L_n = nh/2\pi$; $n = 1, 2, 3 \dots$). Why should the angular momentum have only those values that are integral multiples of $h/2\pi$? The French physicist Louis de Broglie explained this puzzle in 1923, ten years after Bohr proposed his model.

We studied, in Chapter 11, about the de Broglie's hypothesis that material particles, such as electrons, also have a wave nature. C. J. Davisson and L. H. Germer later experimentally verified the wave nature of electrons in 1927. Louis de Broglie argued that the electron in its circular orbit, as proposed by Bohr, must be seen as a particle wave. In analogy to waves travelling on a string, particle waves too can lead to standing waves under resonant conditions. From Chapter 14 of Class XI Physics textbook, we know that when a string is plucked, a vast number of wavelengths are excited. However only those wavelengths survive which have nodes at the ends and form the standing wave in the string. It means that in a string, standing waves are formed when the total distance travelled by a wave down the string and back is one wavelength, two wavelengths, or any integral number of wavelengths. Waves with other wavelengths interfere with themselves upon reflection and their amplitudes quickly drop to zero. For an electron moving in n^{th} circular orbit of radius r_n , the total distance is the circumference of the orbit, $2\pi r_n$. Thus

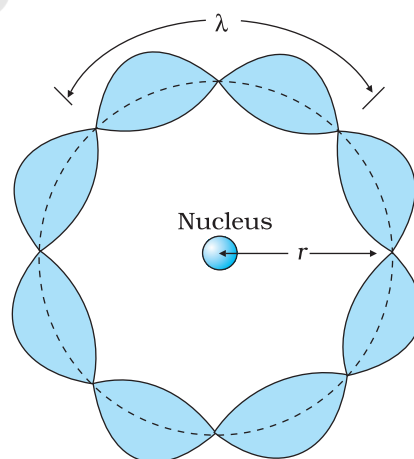


FIGURE 12.8 A standing wave is shown on a circular orbit where four de Broglie wavelengths fit into the circumference of the orbit.

$$2\pi r_n = n\lambda, \quad n = 1, 2, 3\ldots \quad (12.12)$$

Figure 12.8 illustrates a standing particle wave on a circular orbit for $n = 4$, i.e., $2\pi r_n = 4\lambda$, where λ is the de Broglie wavelength of the electron moving in n^{th} orbit. From Chapter 11, we have $\lambda = h/p$, where p is the magnitude of the electron's momentum. If the speed of the electron is much less than the speed of light, the momentum is mv_n . Thus, $\lambda = h/mv_n$. From Eq. (12.12), we have

$$2\pi r_n = n h/mv_n \quad \text{or} \quad m v_n r_n = nh/2\pi$$

This is the quantum condition proposed by Bohr for the angular momentum of the electron [Eq. (12.15)]. In Section 12.5, we saw that this equation is the basis of explaining the discrete orbits and energy levels in hydrogen atom. Thus de Broglie hypothesis provided an explanation for Bohr's second postulate for the quantisation of angular momentum of the orbiting electron. The quantised electron orbits and energy states are due to the wave nature of the electron and only resonant standing waves can persist.

Bohr's model, involving classical trajectory picture (planet-like electron orbiting the nucleus), correctly predicts the gross features of the hydrogenic atoms*, in particular, the frequencies of the radiation emitted or selectively absorbed. This model however has many limitations. Some are:

- (i) The Bohr model is applicable to hydrogenic atoms. It cannot be extended even to mere two electron atoms such as helium. The analysis of atoms with more than one electron was attempted on the lines of Bohr's model for hydrogenic atoms but did not meet with any success. Difficulty lies in the fact that each electron interacts not only with the positively charged nucleus but also with all other electrons. The formulation of Bohr model involves electrical force between positively charged nucleus and electron. It does not include the electrical forces between electrons which necessarily appear in multi-electron atoms.
- (ii) While the Bohr's model correctly predicts the frequencies of the light emitted by hydrogenic atoms, the model is unable to explain the relative intensities of the frequencies in the spectrum. In emission spectrum of hydrogen, some of the visible frequencies have weak intensity, others strong. Why? Experimental observations depict that some transitions are more favoured than others. Bohr's model is unable to account for the intensity variations.

Bohr's model presents an elegant picture of an atom and cannot be generalised to complex atoms. For complex atoms we have to use a new and radical theory based on Quantum Mechanics, which provides a more complete picture of the atomic structure.

* Hydrogenic atoms are the atoms consisting of a nucleus with positive charge $+Ze$ and a single electron, where Z is the proton number. Examples are hydrogen atom, singly ionised helium, doubly ionised lithium, and so forth. In these atoms more complex electron-electron interactions are nonexistent.

SUMMARY

1. Atom, as a whole, is electrically neutral and therefore contains equal amount of positive and negative charges.
2. In *Thomson's model*, an atom is a spherical cloud of positive charges with electrons embedded in it.
3. In *Rutherford's model*, most of the mass of the atom and all its positive charge are concentrated in a tiny nucleus (typically one by ten thousand the size of an atom), and the electrons revolve around it.
4. Rutherford nuclear model has two main difficulties in explaining the structure of atom: (a) It predicts that atoms are unstable because the accelerated electrons revolving around the nucleus must spiral into the nucleus. This contradicts the stability of matter. (b) It cannot explain the characteristic line spectra of atoms of different elements.
5. Atoms of most of the elements are stable and emit characteristic spectrum. The spectrum consists of a set of isolated parallel lines termed as line spectrum. It provides useful information about the atomic structure.
6. To explain the line spectra emitted by atoms, as well as the stability of atoms, Niel's Bohr proposed a model for hydrogenic (single electron) atoms. He introduced three postulates and laid the foundations of quantum mechanics:
 - (a) In a hydrogen atom, an electron revolves in certain stable orbits (called stationary orbits) without the emission of radiant energy.
 - (b) The stationary orbits are those for which the angular momentum is some integral multiple of $h/2\pi$. (Bohr's quantisation condition.) That is $L = nh/2\pi$, where n is an integer called the principal quantum number.
 - (c) The third postulate states that an electron might make a transition from one of its specified non-radiating orbits to another of lower energy. When it does so, a photon is emitted having energy equal to the energy difference between the initial and final states. The frequency (ν) of the emitted photon is then given by

$$h\nu = E_i - E_f$$

An atom absorbs radiation of the same frequency the atom emits, in which case the electron is transferred to an orbit with a higher value of n .

$$E_i + h\nu = E_f$$
7. As a result of the quantisation condition of angular momentum, the electron orbits the nucleus at only specific radii. For a hydrogen atom it is given by

$$r_n = \left(\frac{n^2}{m}\right)\left(\frac{h}{2\pi}\right)^2 \frac{4\pi\epsilon_0}{e^2}$$

The total energy is also quantised:

$$E_n = -\frac{me^4}{8n^2\epsilon_0^2h^2}$$

$$= -13.6 \text{ eV}/n^2$$

The $n = 1$ state is called ground state. In hydrogen atom the ground state energy is -13.6 eV . Higher values of n correspond to excited states ($n > 1$). Atoms are excited to these higher states by collisions with other atoms or electrons or by absorption of a photon of right frequency.

8. de Broglie's hypothesis that electrons have a wavelength $\lambda = h/mv$ gave an explanation for Bohr's quantised orbits by bringing in the wave-particle duality. The orbits correspond to circular standing waves in which the circumference of the orbit equals a whole number of wavelengths.
9. Bohr's model is applicable only to hydrogenic (single electron) atoms. It cannot be extended to even two electron atoms such as helium. This model is also unable to explain for the relative intensities of the frequencies emitted even by hydrogenic atoms.

POINTS TO PONDER

1. Both the Thomson's as well as the Rutherford's models constitute an unstable system. Thomson's model is unstable electrostatically, while Rutherford's model is unstable because of electromagnetic radiation of orbiting electrons.
2. What made Bohr quantise angular momentum (second postulate) and not some other quantity? Note, h has dimensions of angular momentum, and for circular orbits, angular momentum is a very relevant quantity. The second postulate is then so natural!
3. The orbital picture in Bohr's model of the hydrogen atom was inconsistent with the uncertainty principle. It was replaced by modern quantum mechanics in which Bohr's orbits are regions where the electron may be found with large probability.
4. Unlike the situation in the solar system, where planet-planet gravitational forces are very small as compared to the gravitational force of the sun on each planet (because the mass of the sun is so much greater than the mass of any of the planets), the electron-electron electric force interaction is comparable in magnitude to the electron-nucleus electrical force, because the charges and distances are of the same order of magnitude. This is the reason why the Bohr's model with its planet-like electron is not applicable to many electron atoms.
5. Bohr laid the foundation of the quantum theory by postulating specific orbits in which electrons do not radiate. Bohr's model include only one quantum number n . The new theory called quantum mechanics supports Bohr's postulate. However in quantum mechanics (more generally accepted), a given energy level may not correspond to just one quantum state. For example, a state is characterised by four quantum numbers (n , l , m , and s), but for a pure Coulomb potential (as in hydrogen atom) the energy depends only on n .
6. In Bohr model, contrary to ordinary classical expectation, the frequency of revolution of an electron in its orbit is not connected to the frequency of spectral line. The later is the difference between two orbital energies divided by h . For transitions between large quantum numbers (n to $n - 1$, n very large), however, the two coincide as expected.
7. Bohr's semiclassical model based on some aspects of classical physics and some aspects of modern physics also does not provide a true picture of the simplest hydrogenic atoms. The true picture is quantum mechanical affair which differs from Bohr model in a number of fundamental ways. But then if the Bohr model is not strictly correct, why do we bother about it? The reasons which make Bohr's model still useful are:

- (i) The model is based on just three postulates but accounts for almost all the general features of the hydrogen spectrum.
- (ii) The model incorporates many of the concepts we have learnt in classical physics.
- (iii) The model demonstrates how a theoretical physicist occasionally must quite literally ignore certain problems of approach in hopes of being able to make some predictions. If the predictions of the theory or model agree with experiment, a theoretician then must somehow hope to explain away or rationalise the problems that were ignored along the way.

EXERCISES

- 12.1** Choose the correct alternative from the clues given at the end of the each statement:
- (a) The size of the atom in Thomson's model is the atomic size in Rutherford's model. (much greater than/no different from/much less than.)
 - (b) In the ground state of electrons are in stable equilibrium, while in electrons always experience a net force. (Thomson's model/ Rutherford's model.)
 - (c) A *classical* atom based on is doomed to collapse. (Thomson's model/ Rutherford's model.)
 - (d) An atom has a nearly continuous mass distribution in a but has a highly non-uniform mass distribution in (Thomson's model/ Rutherford's model.)
 - (e) The positively charged part of the atom possesses most of the mass in (Rutherford's model/both the models.)
- 12.2** Suppose you are given a chance to repeat the alpha-particle scattering experiment using a thin sheet of solid hydrogen in place of the gold foil. (Hydrogen is a solid at temperatures below 14 K.) What results do you expect?
- 12.3** A difference of 2.3 eV separates two energy levels in an atom. What is the frequency of radiation emitted when the atom make a transition from the upper level to the lower level?
- 12.4** The ground state energy of hydrogen atom is -13.6 eV. What are the kinetic and potential energies of the electron in this state?
- 12.5** A hydrogen atom initially in the ground level absorbs a photon, which excites it to the $n = 4$ level. Determine the wavelength and frequency of photon.
- 12.6** (a) Using the Bohr's model calculate the speed of the electron in a hydrogen atom in the $n = 1, 2$, and 3 levels. (b) Calculate the orbital period in each of these levels.
- 12.7** The radius of the innermost electron orbit of a hydrogen atom is 5.3×10^{-11} m. What are the radii of the $n = 2$ and $n = 3$ orbits?
- 12.8** A 12.5 eV electron beam is used to bombard gaseous hydrogen at room temperature. What series of wavelengths will be emitted?
- 12.9** In accordance with the Bohr's model, find the quantum number that characterises the earth's revolution around the sun in an orbit of radius 1.5×10^{11} m with orbital speed 3×10^4 m/s. (Mass of earth = 6.0×10^{24} kg.)



Chapter Thirteen

NUCLEI

13.1 INTRODUCTION

In the previous chapter, we have learnt that in every atom, the positive charge and mass are densely concentrated at the centre of the atom forming its nucleus. The overall dimensions of a nucleus are much smaller than those of an atom. Experiments on scattering of α -particles demonstrated that the radius of a nucleus was smaller than the radius of an atom by a factor of about 10^4 . This means the volume of a nucleus is about 10^{-12} times the volume of the atom. In other words, an atom is almost empty. If an atom is enlarged to the size of a classroom, the nucleus would be of the size of pinhead. Nevertheless, the nucleus contains most (more than 99.9%) of the mass of an atom.

Does the nucleus have a structure, just as the atom does? If so, what are the constituents of the nucleus? How are these held together? In this chapter, we shall look for answers to such questions. We shall discuss various properties of nuclei such as their size, mass and stability, and also associated nuclear phenomena such as radioactivity, fission and fusion.

13.2 ATOMIC MASSES AND COMPOSITION OF NUCLEUS

The mass of an atom is very small, compared to a kilogram; for example, the mass of a carbon atom, ^{12}C , is 1.992647×10^{-26} kg. Kilogram is not a very convenient unit to measure such small quantities. Therefore, a

different mass unit is used for expressing atomic masses. This unit is the atomic mass unit (u), defined as $1/12^{\text{th}}$ of the mass of the carbon (^{12}C) atom. According to this definition

$$\begin{aligned} 1\text{u} &= \frac{\text{mass of one } ^{12}\text{C atom}}{12} \\ &= \frac{1.992647 \times 10^{-26} \text{ kg}}{12} \\ &= 1.660539 \times 10^{-27} \text{ kg} \end{aligned} \quad (13.1)$$

The atomic masses of various elements expressed in atomic mass unit (u) are close to being integral multiples of the mass of a hydrogen atom. There are, however, many striking exceptions to this rule. For example, the atomic mass of chlorine atom is 35.46 u.

Accurate measurement of atomic masses is carried out with a mass spectrometer. The measurement of atomic masses reveals the existence of different types of atoms of the same element, which exhibit the same chemical properties, but differ in mass. Such atomic species of the same element differing in mass are called *isotopes*. (In Greek, isotope means the same place, i.e. they occur in the same place in the periodic table of elements.) It was found that practically every element consists of a mixture of several isotopes. The relative abundance of different isotopes differs from element to element. Chlorine, for example, has two isotopes having masses 34.98 u and 36.98 u, which are nearly integral multiples of the mass of a hydrogen atom. The relative abundances of these isotopes are 75.4 and 24.6 per cent, respectively. Thus, the average mass of a chlorine atom is obtained by the weighted average of the masses of the two isotopes, which works out to be

$$\begin{aligned} &= \frac{75.4 \times 34.98 + 24.6 \times 36.98}{100} \\ &= 35.47 \text{ u} \end{aligned}$$

which agrees with the atomic mass of chlorine.

Even the lightest element, hydrogen has three isotopes having masses 1.0078 u, 2.0141 u, and 3.0160 u. The nucleus of the lightest atom of hydrogen, which has a relative abundance of 99.985%, is called the proton. The mass of a proton is

$$m_p = 1.00727 \text{ u} = 1.67262 \times 10^{-27} \text{ kg} \quad (13.2)$$

This is equal to the mass of the hydrogen atom ($= 1.00783\text{u}$), minus the mass of a single electron ($m_e = 0.00055 \text{ u}$). The other two isotopes of hydrogen are called deuterium and tritium. Tritium nuclei, being unstable, do not occur naturally and are produced artificially in laboratories.

The positive charge in the nucleus is that of the protons. A proton carries one unit of fundamental charge and is stable. It was earlier thought that the nucleus may contain electrons, but this was ruled out later using arguments based on quantum theory. All the electrons of an atom are outside the nucleus. We know that the number of these electrons outside the nucleus of the atom is Z , the atomic number. The total charge of the

atomic electrons is thus $(-Ze)$, and since the atom is neutral, the charge of the nucleus is $(+Ze)$. The number of protons in the nucleus of the atom is, therefore, exactly Z , the atomic number.

Discovery of Neutron

Since the nuclei of deuterium and tritium are isotopes of hydrogen, they must contain only one proton each. But the masses of the nuclei of hydrogen, deuterium and tritium are in the ratio of 1:2:3. Therefore, the nuclei of deuterium and tritium must contain, in addition to a proton, some neutral matter. The amount of neutral matter present in the nuclei of these isotopes, expressed in units of mass of a proton, is approximately equal to one and two, respectively. This fact indicates that the nuclei of atoms contain, in addition to protons, neutral matter in multiples of a basic unit. This hypothesis was verified in 1932 by James Chadwick who observed emission of neutral radiation when beryllium nuclei were bombarded with alpha-particles (α -particles are helium nuclei, to be discussed in a later section). It was found that this neutral radiation could knock out protons from light nuclei such as those of helium, carbon and nitrogen. The only neutral radiation known at that time was photons (electromagnetic radiation). Application of the principles of conservation of energy and momentum showed that if the neutral radiation consisted of photons, the energy of photons would have to be much higher than is available from the bombardment of beryllium nuclei with α -particles. The clue to this puzzle, which Chadwick satisfactorily solved, was to assume that the neutral radiation consists of a new type of neutral particles called *neutrons*. From conservation of energy and momentum, he was able to determine the mass of new particle 'as very nearly the same as mass of proton'.

The mass of a neutron is now known to a high degree of accuracy. It is

$$m_n = 1.00866 \text{ u} = 1.6749 \times 10^{-27} \text{ kg} \quad [13.3]$$

Chadwick was awarded the 1935 Nobel Prize in Physics for his discovery of the neutron.

A free neutron, unlike a free proton, is unstable. It decays into a proton, an electron and a antineutrino (another elementary particle), and has a mean life of about 1000s. It is, however, stable inside the nucleus.

The composition of a nucleus can now be described using the following terms and symbols:

Z - atomic number = number of protons [13.4(a)]

N - neutron number = number of neutrons [13.4(b)]

A - mass number = $Z + N$
= total number of protons and neutrons [13.4(c)]

One also uses the term nucleon for a proton or a neutron. Thus the number of nucleons in an atom is its mass number A .

Nuclear species or nuclides are shown by the notation ${}_Z^A\text{X}$ where X is the chemical symbol of the species. For example, the nucleus of gold is denoted by ${}_{79}^{197}\text{Au}$. It contains 197 nucleons, of which 79 are protons and the rest 118 are neutrons.

The composition of isotopes of an element can now be readily explained. The nuclei of isotopes of a given element contain the same number of protons, but differ from each other in their number of neutrons. Deuterium, ${}^2_1\text{H}$, which is an isotope of hydrogen, contains one proton and one neutron. Its other isotope tritium, ${}^3_1\text{H}$, contains one proton and two neutrons. The element gold has 32 isotopes, ranging from $A = 173$ to $A = 204$. We have already mentioned that chemical properties of elements depend on their electronic structure. As the atoms of isotopes have identical electronic structure they have identical chemical behaviour and are placed in the same location in the periodic table.

All nuclides with same mass number A are called *isobars*. For example, the nuclides ${}^3_1\text{H}$ and ${}^3_2\text{He}$ are isobars. Nuclides with same neutron number N but different atomic number Z , for example ${}^{198}_{80}\text{Hg}$ and ${}^{197}_{79}\text{Au}$, are called *isotones*.

13.3 SIZE OF THE NUCLEUS

As we have seen in Chapter 12, Rutherford was the pioneer who postulated and established the existence of the atomic nucleus. At Rutherford's suggestion, Geiger and Marsden performed their classic experiment: on the scattering of α -particles from thin gold foils. Their experiments revealed that the distance of closest approach to a gold nucleus of an α -particle of kinetic energy 5.5 MeV is about 4.0×10^{-14} m. The scattering of α -particle by the gold sheet could be understood by Rutherford by assuming that the coulomb repulsive force was solely responsible for scattering. Since the positive charge is confined to the nucleus, the actual size of the nucleus has to be less than 4.0×10^{-14} m.

If we use α -particles of higher energies than 5.5 MeV, the distance of closest approach to the gold nucleus will be smaller and at some point the scattering will begin to be affected by the short range nuclear forces, and differ from Rutherford's calculations. Rutherford's calculations are based on pure coulomb repulsion between the positive charges of the α -particle and the gold nucleus. From the distance at which deviations set in, nuclear sizes can be inferred.

By performing scattering experiments in which fast electrons, instead of α -particles, are projectiles that bombard targets made up of various elements, the sizes of nuclei of various elements have been accurately measured.

It has been found that a nucleus of mass number A has a radius

$$R = R_0 A^{1/3} \quad (13.5)$$

where $R_0 = 1.2 \times 10^{-15}$ m ($= 1.2$ fm; $1 \text{ fm} = 10^{-15}$ m). This means the volume of the nucleus, which is proportional to R^3 is proportional to A . Thus the density of nucleus is a constant, independent of A , for all nuclei. Different nuclei are like a drop of liquid of constant density. The density of nuclear matter is approximately $2.3 \times 10^{17} \text{ kg m}^{-3}$. This density is very large compared to ordinary matter, say water, which is 10^3 kg m^{-3} . This is understandable, as we have already seen that most of the atom is empty. Ordinary matter consisting of atoms has a large amount of empty space.

Example 13.1 Given the mass of iron nucleus as 55.85u and A=56, find the nuclear density?

Solution

$$m_{\text{Fe}} = 55.85, \quad u = 9.27 \times 10^{-26} \text{ kg}$$

$$\text{Nuclear density} = \frac{\text{mass}}{\text{volume}} = \frac{9.27 \times 10^{-26}}{(4\pi/3)(1.2 \times 10^{-15})^3} \times \frac{1}{56}$$

$$= 2.29 \times 10^{17} \text{ kg m}^{-3}$$

The density of matter in neutron stars (an astrophysical object) is comparable to this density. This shows that matter in these objects has been compressed to such an extent that they resemble a *big nucleus*.

13.4 MASS-ENERGY AND NUCLEAR BINDING ENERGY

13.4.1 Mass – Energy

Einstein showed from his theory of special relativity that it is necessary to treat mass as another form of energy. Before the advent of this theory of special relativity it was presumed that mass and energy were conserved separately in a reaction. However, Einstein showed that mass is another form of energy and one can convert mass-energy into other forms of energy, say kinetic energy and vice-versa.

Einstein gave the famous mass-energy equivalence relation

$$E = mc^2 \quad (13.6)$$

Here the energy equivalent of mass m is related by the above equation and c is the velocity of light in vacuum and is approximately equal to $3 \times 10^8 \text{ m s}^{-1}$.

Example 13.2 Calculate the energy equivalent of 1 g of substance.

Solution

$$\text{Energy, } E = 10^{-3} \times (3 \times 10^8)^2 \text{ J}$$

$$E = 10^{-3} \times 9 \times 10^{16} = 9 \times 10^{13} \text{ J}$$

Thus, if one gram of matter is converted to energy, there is a release of enormous amount of energy.

Experimental verification of the Einstein's mass-energy relation has been achieved in the study of nuclear reactions amongst nucleons, nuclei, electrons and other more recently discovered particles. In a reaction the conservation law of energy states that the initial energy and the final energy are equal provided the energy associated with mass is also included. This concept is important in understanding nuclear masses and the interaction of nuclei with one another. They form the subject matter of the next few sections.

13.4.2 Nuclear binding energy

In Section 13.2 we have seen that the nucleus is made up of neutrons and protons. Therefore it may be expected that the mass of the nucleus is equal to the total mass of its individual protons and neutrons. However,

the nuclear mass M is found to be always less than this. For example, let us consider $^{16}_8\text{O}$; a nucleus which has 8 neutrons and 8 protons. We have

$$\text{Mass of 8 neutrons} = 8 \times 1.00866 \text{ u}$$

$$\text{Mass of 8 protons} = 8 \times 1.00727 \text{ u}$$

$$\text{Mass of 8 electrons} = 8 \times 0.00055 \text{ u}$$

$$\begin{aligned} \text{Therefore the expected mass of } ^{16}_8\text{O nucleus} \\ = 8 \times 2.01593 \text{ u} = 16.12744 \text{ u.} \end{aligned}$$

The atomic mass of $^{16}_8\text{O}$ found from mass spectroscopy experiments is seen to be 15.99493 u. Subtracting the mass of 8 electrons ($8 \times 0.00055 \text{ u}$) from this, we get the experimental mass of $^{16}_8\text{O}$ nucleus to be 15.99053 u.

Thus, we find that the mass of the $^{16}_8\text{O}$ nucleus is less than the total mass of its constituents by 0.13691 u. The difference in mass of a nucleus and its constituents, ΔM , is called the *mass defect*, and is given by

$$\Delta M = [Zm_p + (A - Z)m_n] - M \quad (13.7)$$

What is the meaning of the mass defect? It is here that Einstein's equivalence of mass and energy plays a role. Since the mass of the oxygen nucleus is less than the sum of the masses of its constituents (8 protons and 8 neutrons, in the unbound state), the equivalent energy of the oxygen nucleus is less than that of the sum of the equivalent energies of its constituents. If one wants to break the oxygen nucleus into 8 protons and 8 neutrons, this extra energy $\Delta M c^2$, has to be supplied. This energy required E_b is related to the mass defect by

$$E_b = \Delta M c^2 \quad (13.8)$$

Example 13.3 Find the energy equivalent of one atomic mass unit, first in Joules and then in MeV. Using this, express the mass defect of $^{16}_8\text{O}$ in MeV/c^2 .

Solution

$$1 \text{ u} = 1.6605 \times 10^{-27} \text{ kg}$$

To convert it into energy units, we multiply it by c^2 and find that energy equivalent = $1.6605 \times 10^{-27} \times (2.9979 \times 10^8)^2 \text{ kg m}^2/\text{s}^2$

$$= 1.4924 \times 10^{-10} \text{ J}$$

$$= \frac{1.4924 \times 10^{-10}}{1.602 \times 10^{-19}} \text{ eV}$$

$$= 0.9315 \times 10^9 \text{ eV}$$

$$= 931.5 \text{ MeV}$$

$$\text{or, } 1 \text{ u} = 931.5 \text{ MeV}/c^2$$

$$\begin{aligned} \text{For } ^{16}_8\text{O}, \quad \Delta M = 0.13691 \text{ u} &= 0.13691 \times 931.5 \text{ MeV}/c^2 \\ &= 127.5 \text{ MeV}/c^2 \end{aligned}$$

The energy needed to separate $^{16}_8\text{O}$ into its constituents is thus 127.5 MeV/c^2 .

If a certain number of neutrons and protons are brought together to form a nucleus of a certain charge and mass, an energy E_b will be released

in the process. The energy E_b is called the *binding energy* of the nucleus. If we separate a nucleus into its nucleons, we would have to supply a total energy equal to E_b , to those particles. Although we cannot tear apart a nucleus in this way, the nuclear binding energy is still a convenient measure of how well a nucleus is held together. A more useful measure of the binding between the constituents of the nucleus is the *binding energy per nucleon*, E_{bn} , which is the ratio of the binding energy E_b of a nucleus to the number of the nucleons, A , in that nucleus:

$$E_{bn} = E_b / A \quad (13.9)$$

We can think of binding energy per nucleon as the average energy per nucleon needed to separate a nucleus into its individual nucleons.

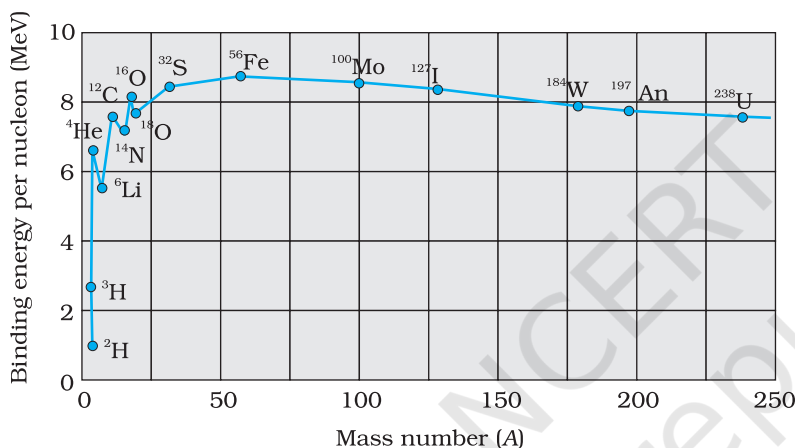


FIGURE 13.1 The binding energy per nucleon as a function of mass number.

Figure 13.1 is a plot of the binding energy per nucleon E_{bn} versus the mass number A for a large number of nuclei. We notice the following main features of the plot:

- (i) the binding energy per nucleon, E_{bn} , is practically constant, i.e. practically independent of the atomic number for nuclei of middle mass number ($30 < A < 170$). The curve has a maximum of about 8.75 MeV for $A = 56$ and has a value of 7.6 MeV for $A = 238$.
- (ii) E_{bn} is lower for both light nuclei ($A < 30$) and heavy nuclei ($A > 170$).

We can draw some conclusions from these two observations:

- (i) The force is attractive and sufficiently strong to produce a binding energy of a few MeV per nucleon.
- (ii) The constancy of the binding energy in the range $30 < A < 170$ is a consequence of the fact that the nuclear force is short-ranged. Consider a particular nucleon inside a sufficiently large nucleus. It will be under the influence of only some of its neighbours, which come within the range of the nuclear force. If any other nucleon is at a distance more than the range of the nuclear force from the particular nucleon it will have no influence on the binding energy of the nucleon under consideration. If a nucleon can have a maximum of p neighbours within the range of nuclear force, its binding energy would be proportional to p . Let the binding energy of the nucleus be pk , where k is a constant having the dimensions of energy. If we increase A by adding nucleons they will not change the binding energy of a nucleon inside. Since most of the nucleons in a large nucleus reside inside it and not on the surface, the change in binding energy per nucleon would be small. The binding energy per nucleon is a constant and is approximately equal to pk . The property that a given nucleon

influences only nucleons close to it is also referred to as saturation property of the nuclear force.

- (iii) A very heavy nucleus, say $A = 240$, has lower binding energy per nucleon compared to that of a nucleus with $A = 120$. Thus if a nucleus $A = 240$ breaks into two $A = 120$ nuclei, nucleons get more tightly bound. This implies energy would be released in the process. It has very important implications for energy production through *fission*, to be discussed later in Section 13.7.1.
- (iv) Consider two very light nuclei ($A \leq 10$) joining to form a heavier nucleus. The binding energy per nucleon of the fused heavier nuclei is more than the binding energy per nucleon of the lighter nuclei. This means that the final system is more tightly bound than the initial system. Again energy would be released in such a process of *fusion*. This is the energy source of sun, to be discussed later in Section 13.7.2.

13.5 NUCLEAR FORCE

The force that determines the motion of atomic electrons is the familiar Coulomb force. In Section 13.4, we have seen that for average mass nuclei the binding energy per nucleon is approximately 8 MeV, which is much larger than the binding energy in atoms. Therefore, to bind a nucleus together there must be a strong attractive force of a totally different kind. It must be strong enough to overcome the repulsion between the (positively charged) protons and to bind both protons and neutrons into the tiny nuclear volume. We have already seen that the constancy of binding energy per nucleon can be understood in terms of its short-range. Many features of the nuclear binding force are summarised below. These are obtained from a variety of experiments carried out during 1930 to 1950.

- (i) The nuclear force is much stronger than the Coulomb force acting between charges or the gravitational forces between masses. The nuclear binding force has to dominate over the Coulomb repulsive force between protons inside the nucleus. This happens only because the nuclear force is much stronger than the coulomb force. The gravitational force is much weaker than even Coulomb force.
- (ii) The nuclear force between two nucleons falls rapidly to zero as their distance is more than a few femtometres. This leads to *saturation of forces* in a medium or a large-sized nucleus, which is the reason for the constancy of the binding energy per nucleon.

A rough plot of the potential energy between two nucleons as a function of distance is shown in the Fig. 13.2. The potential energy is a minimum at a distance r_0 of about 0.8 fm. This means that the force is attractive for distances larger than 0.8 fm and repulsive if they are separated by distances less than 0.8 fm.

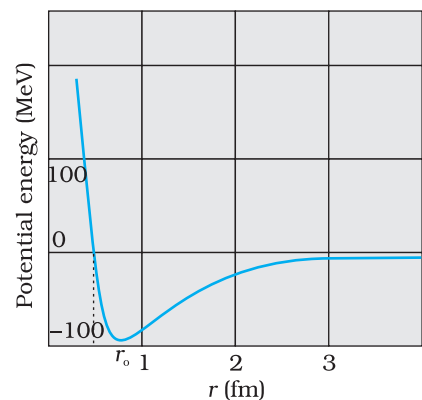


FIGURE 13.2 Potential energy of a pair of nucleons as a function of their separation. For a separation greater than r_0 , the force is attractive and for separations less than r_0 , the force is strongly repulsive.

(iii) The nuclear force between neutron-neutron, proton-neutron and proton-proton is approximately the same. The nuclear force does not depend on the electric charge.

Unlike Coulomb's law or the Newton's law of gravitation there is no simple mathematical form of the nuclear force.

13.6 RADIOACTIVITY

A. H. Becquerel discovered radioactivity in 1896 purely by accident. While studying the fluorescence and phosphorescence of compounds irradiated with visible light, Becquerel observed an interesting phenomenon. After illuminating some pieces of uranium-potassium sulphate with visible light, he wrapped them in black paper and separated the package from a photographic plate by a piece of silver. When, after several hours of exposure, the photographic plate was developed, it showed blackening due to something that must have been emitted by the compound and was able to penetrate both black paper and the silver.

Experiments performed subsequently showed that radioactivity was a nuclear phenomenon in which an unstable nucleus undergoes a decay. This is referred to as *radioactive decay*. Three types of radioactive decay occur in nature :

- (i) α -decay in which a helium nucleus ${}^4_2\text{He}$ is emitted;
- (ii) β -decay in which electrons or positrons (particles with the same mass as electrons, but with a charge exactly opposite to that of electron) are emitted;
- (iii) γ -decay in which high energy (hundreds of keV or more) photons are emitted.

13.7 NUCLEAR ENERGY

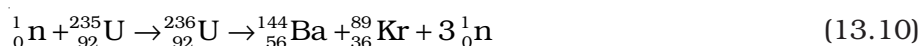
The curve of binding energy per nucleon E_{bn} , given in Fig. 13.1, has a long flat middle region between $A = 30$ and $A = 170$. In this region the binding energy per nucleon is nearly constant (8.0 MeV). For the lighter nuclei region, $A < 30$, and for the heavier nuclei region, $A > 170$, the binding energy per nucleon is less than 8.0 MeV, as we have noted earlier. Now, the greater the binding energy, the less is the total mass of a bound system, such as a nucleus. Consequently, if nuclei with less total binding energy transform to nuclei with greater binding energy, there will be a net energy release. This is what happens when a heavy nucleus decays into two or more intermediate mass fragments (*fission*) or when light nuclei fuse into a heavier nucleus (*fusion*.)

Exothermic chemical reactions underlie conventional energy sources such as coal or petroleum. Here the energies involved are in the range of electron volts. On the other hand, in a nuclear reaction, the energy release is of the order of MeV. Thus for the same quantity of matter, nuclear sources produce a million times more energy than a chemical source. Fission of 1 kg of uranium, for example, generates 10^{14} J of energy; compare it with burning of 1 kg of coal that gives 10^7 J.

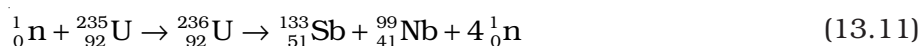
13.7.1 Fission

New possibilities emerge when we go beyond natural radioactive decays and study nuclear reactions by bombarding nuclei with other nuclear particles such as proton, neutron, α -particle, etc.

A most important neutron-induced nuclear reaction is fission. An example of fission is when a uranium isotope ${}_{92}^{235}\text{U}$ bombarded with a neutron breaks into two intermediate mass nuclear fragments



The same reaction can produce other pairs of intermediate mass fragments



Or, as another example,



The fragment products are radioactive nuclei; they emit β particles in succession to achieve stable end products.

The energy released (the Q value) in the fission reaction of nuclei like uranium is of the order of 200 MeV per fissioning nucleus. This is estimated as follows:

Let us take a nucleus with $A = 240$ breaking into two fragments each of $A = 120$. Then

E_{bn} for $A = 240$ nucleus is about 7.6 MeV,

E_{bn} for the two $A = 120$ fragment nuclei is about 8.5 MeV.

$\therefore$ Gain in binding energy for nucleon is about 0.9 MeV.

Hence the total gain in binding energy is 240×0.9 or 216 MeV.

The disintegration energy in fission events first appears as the kinetic energy of the fragments and neutrons. Eventually it is transferred to the surrounding matter appearing as heat. The source of energy in nuclear reactors, which produce electricity, is nuclear fission. The enormous energy released in an atom bomb comes from uncontrolled nuclear fission.

13.7.2 Nuclear fusion – energy generation in stars

When two light nuclei fuse to form a larger nucleus, energy is released, since the larger nucleus is more tightly bound, as seen from the binding energy curve in Fig. 13.1. Some examples of such energy liberating nuclear fusion reactions are :



In the first reaction, two protons combine to form a deuteron and a positron with a release of 0.42 MeV energy. In reaction [13.13(b)], two

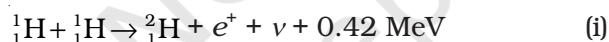
deuterons combine to form the light isotope of helium. In reaction (13.13c), two deuterons combine to form a triton and a proton. For fusion to take place, the two nuclei must come close enough so that attractive short-range nuclear force is able to affect them. However, since they are both positively charged particles, they experience coulomb repulsion. They, therefore, must have enough energy to overcome this coulomb barrier. The height of the barrier depends on the charges and radii of the two interacting nuclei. It can be shown, for example, that the barrier height for two protons is ~ 400 keV, and is higher for nuclei with higher charges. We can estimate the temperature at which two protons in a proton gas would (averagely) have enough energy to overcome the coulomb barrier:

$$(3/2)kT = K \simeq 400 \text{ keV, which gives } T \sim 3 \times 10^9 \text{ K.}$$

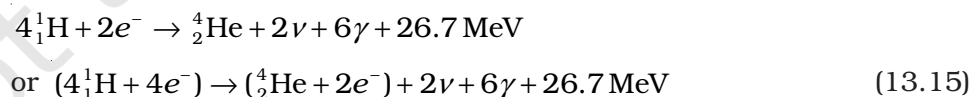
When fusion is achieved by raising the temperature of the system so that particles have enough kinetic energy to overcome the coulomb repulsive behaviour, it is called *thermonuclear fusion*.

Thermonuclear fusion is the source of energy output in the interior of stars. The interior of the sun has a temperature of 1.5×10^7 K, which is considerably less than the estimated temperature required for fusion of particles of average energy. Clearly, fusion in the sun involves protons whose energies are much above the average energy.

The fusion reaction in the sun is a multi-step process in which the hydrogen is burned into helium. Thus, the fuel in the sun is the hydrogen in its core. The *proton-proton (p, p) cycle* by which this occurs is represented by the following sets of reactions:



For the fourth reaction to occur, the first three reactions must occur twice, in which case two light helium nuclei unite to form ordinary helium nucleus. If we consider the combination $2(\text{i}) + 2(\text{ii}) + 2(\text{iii}) + (\text{iv})$, the net effect is



Thus, four hydrogen atoms combine to form an ${}_2^4\text{He}$ atom with a release of 26.7 MeV of energy.

Helium is not the only element that can be synthesized in the interior of a star. As the hydrogen in the core gets depleted and becomes helium, the core starts to cool. The star begins to collapse under its own gravity which increases the temperature of the core. If this temperature increases to about 10^8 K, fusion takes place again, this time of helium nuclei into carbon. This kind of process can generate through fusion higher and higher mass number elements. But elements more massive than those near the peak of the binding energy curve in Fig. 13.1 cannot be so produced.

The age of the sun is about 5×10^9 y and it is estimated that there is enough hydrogen in the sun to keep it going for another 5 billion years. After that, the hydrogen burning will stop and the sun will begin to cool and will start to collapse under gravity, which will raise the core temperature. The outer envelope of the sun will expand, turning it into the so called *red giant*.

13.7.3 Controlled thermonuclear fusion

The natural thermonuclear fusion process in a star is replicated in a thermonuclear fusion device. In controlled fusion reactors, the aim is to generate steady power by heating the nuclear fuel to a temperature in the range of 10^8 K. At these temperatures, the fuel is a mixture of positive ions and electrons (plasma). The challenge is to confine this plasma, since no container can stand such a high temperature. Several countries around the world including India are developing techniques in this connection. If successful, fusion reactors will hopefully supply almost unlimited power to humanity.

Example 13.4 Answer the following questions:

- Are the equations of nuclear reactions (such as those given in Section 13.7) 'balanced' in the sense a chemical equation (e.g., $2\text{H}_2 + \text{O}_2 \rightarrow 2\text{H}_2\text{O}$) is? If not, in what sense are they balanced on both sides?
- If both the number of protons and the number of neutrons are conserved in each nuclear reaction, in what way is mass converted into energy (or vice-versa) in a nuclear reaction?
- A general impression exists that mass-energy interconversion takes place only in nuclear reaction and never in chemical reaction. This is strictly speaking, incorrect. Explain.

Solution

- A chemical equation is balanced in the sense that the number of atoms of each element is the same on both sides of the equation. A chemical reaction merely alters the original combinations of atoms. In a nuclear reaction, elements may be transmuted. Thus, the number of atoms of each element is not necessarily conserved in a nuclear reaction. However, the number of protons and the number of neutrons are both separately conserved in a nuclear reaction. [Actually, even this is not strictly true in the realm of very high energies – what is strictly conserved is the total charge and total 'baryon number'. We need not pursue this matter here.] In nuclear reactions (e.g., Eq. 13.10), the number of protons and the number of neutrons are the same on the two sides of the equation.
- We know that the binding energy of a nucleus gives a negative contribution to the mass of the nucleus (mass defect). Now, since proton number and neutron number are conserved in a nuclear reaction, the total rest mass of neutrons and protons is the same on either side of a reaction. But the total binding energy of nuclei on the left side need not be the same as that on the right hand side. The difference in these binding energies appears as energy released or absorbed in a nuclear reaction. Since binding energy

contributes to mass, we say that the difference in the total mass of nuclei on the two sides get converted into energy or vice-versa. It is in these sense that a nuclear reaction is an example of mass-energy interconversion.

- (c) From the point of view of mass-energy interconversion, a chemical reaction is similar to a nuclear reaction *in principle*. The energy released or absorbed in a chemical reaction can be traced to the difference in chemical (not nuclear) binding energies of atoms and molecules on the two sides of a reaction. Since, strictly speaking, chemical binding energy also gives a negative contribution (mass defect) to the total mass of an atom or molecule, we can equally well say that the difference in the total mass of atoms or molecules, on the two sides of the chemical reaction gets converted into energy or vice-versa. However, the mass defects involved in a chemical reaction are almost a million times smaller than those in a nuclear reaction. This is the reason for the general impression, (which is *incorrect*) that mass-energy interconversion does not take place in a chemical reaction.

SUMMARY

1. An atom has a nucleus. The nucleus is positively charged. The radius of the nucleus is smaller than the radius of an atom by a factor of 10^4 . More than 99.9% mass of the atom is concentrated in the nucleus.
2. On the atomic scale, mass is measured in atomic mass units (u). By definition, 1 atomic mass unit (1u) is $1/12^{\text{th}}$ mass of one atom of ^{12}C ; $1\text{u} = 1.660563 \times 10^{-27} \text{ kg}$.
3. A nucleus contains a neutral particle called neutron. Its mass is almost the same as that of proton
4. The atomic number Z is the number of protons in the atomic nucleus of an element. The mass number A is the total number of protons and neutrons in the atomic nucleus; $A = Z + N$; Here N denotes the number of neutrons in the nucleus.

A nuclear species or a nuclide is represented as ${}_Z^A\text{X}$, where X is the chemical symbol of the species.

Nuclides with the same atomic number Z , but different neutron number N are called *isotopes*. Nuclides with the same A are *isobars* and those with the same N are *isotones*.

Most elements are mixtures of two or more isotopes. The atomic mass of an element is a weighted average of the masses of its isotopes and calculated in accordance to the relative abundances of the isotopes.

5. A nucleus can be considered to be spherical in shape and assigned a radius. Electron scattering experiments allow determination of the nuclear radius; it is found that radii of nuclei fit the formula

$$R = R_0 A^{1/3},$$

where $R_0 = \text{a constant} = 1.2 \text{ fm}$. This implies that the nuclear density is independent of A . It is of the order of 10^{17} kg/m^3 .

6. Neutrons and protons are bound in a nucleus by the short-range strong nuclear force. The nuclear force does not distinguish between neutron and proton.

7. The nuclear mass M is always less than the total mass, Σm , of its constituents. The difference in mass of a nucleus and its constituents is called the *mass defect*,

$$\Delta M = (Z m_p + (A - Z)m_n) - M$$

Using Einstein's mass energy relation, we express this mass difference in terms of energy as

$$\Delta E_b = \Delta M c^2$$

The energy ΔE_b represents the *binding energy* of the nucleus. In the mass number range $A = 30$ to 170 , the binding energy per nucleon is nearly constant, about 8 MeV/nucleon .

8. Energies associated with nuclear processes are about a million times larger than chemical process.

9. The Q -value of a nuclear process is

$$Q = \text{final kinetic energy} - \text{initial kinetic energy.}$$

Due to conservation of mass-energy, this is also,

$$Q = (\text{sum of initial masses} - \text{sum of final masses})c^2$$

10. Radioactivity is the phenomenon in which nuclei of a given species transform by giving out α or β or γ rays; α -rays are helium nuclei; β -rays are electrons. γ -rays are electromagnetic radiation of wavelengths shorter than X-rays.

11. Energy is released when less tightly bound nuclei are transmuted into more tightly bound nuclei. In fission, a heavy nucleus like ${}^{235}_{92}\text{U}$ breaks into two smaller fragments, e.g., ${}^{235}_{92}\text{U} + {}^1_0\text{n} \rightarrow {}^{133}_{51}\text{Sb} + {}^{99}_{41}\text{Nb} + 4 {}^1_0\text{n}$

12. In fusion, lighter nuclei combine to form a larger nucleus. Fusion of hydrogen nuclei into helium nuclei is the source of energy of all stars including our sun.

Physical Quantity	Symbol	Dimensions	Units	Remarks
Atomic mass unit		[M]	u	Unit of mass for expressing atomic or nuclear masses. One atomic mass unit equals $1/12^{\text{th}}$ of the mass of ${}^{12}\text{C}$ atom.
Disintegration or decay constant	λ	$[\text{T}^{-1}]$	s^{-1}	
Half-life	$T_{1/2}$	[T]	s	Time taken for the decay of one-half of the initial number of nuclei present in a radioactive sample.
Mean life	τ	[T]	s	Time at which number of nuclei has been reduced to e^{-1} of its initial value
Activity of a radioactive sample	R	$[\text{T}^{-1}]$	Bq	Measure of the activity of a radioactive source.

POINTS TO PONDER

1. The density of nuclear matter is independent of the size of the nucleus. The mass density of the atom does not follow this rule.
2. The radius of a nucleus determined by electron scattering is found to be slightly different from that determined by alpha-particle scattering. This is because electron scattering senses the charge distribution of the nucleus, whereas alpha and similar particles sense the nuclear matter.
3. After Einstein showed the equivalence of mass and energy, $E = mc^2$, we cannot any longer speak of separate laws of conservation of mass and conservation of energy, but we have to speak of a unified law of conservation of mass and energy. The most convincing evidence that this principle operates in nature comes from nuclear physics. It is central to our understanding of nuclear energy and harnessing it as a source of power. Using the principle, Q of a nuclear process (decay or reaction) can be expressed also in terms of initial and final masses.
4. The nature of the binding energy (per nucleon) curve shows that exothermic nuclear reactions are possible, when two light nuclei fuse or when a heavy nucleus undergoes fission into nuclei with intermediate mass.
5. For fusion, the light nuclei must have sufficient initial energy to overcome the coulomb potential barrier. That is why fusion requires very high temperatures.
6. Although the binding energy (per nucleon) curve is smooth and slowly varying, it shows peaks at nuclides like ${}^4\text{He}$, ${}^{16}\text{O}$ etc. This is considered as evidence of atom-like shell structure in nuclei.
7. Electron and positron are a particle-antiparticle pair. They are identical in mass; their charges are equal in magnitude and opposite. (It is found that when an electron and a positron come together, they annihilate each other giving energy in the form of gamma-ray photons.)
8. Radioactivity is an indication of the instability of nuclei. Stability requires the ratio of neutron to proton to be around 1:1 for light nuclei. This ratio increases to about 3:2 for heavy nuclei. (More neutrons are required to overcome the effect of repulsion among the protons.) Nuclei which are away from the stability ratio, i.e., nuclei which have an excess of neutrons or protons are unstable. In fact, only about 10% of known isotopes (of all elements), are stable. Others have been either artificially produced in the laboratory by bombarding α , p, d, n or other particles on targets of stable nuclear species or identified in astronomical observations of matter in the universe.

EXERCISES

You may find the following data useful in solving the exercises:

$$e = 1.6 \times 10^{-19} \text{ C} \quad N = 6.023 \times 10^{23} \text{ per mole}$$

$$1/(4\pi\epsilon_0) = 9 \times 10^9 \text{ N m}^2/\text{C}^2 \quad k = 1.381 \times 10^{-23} \text{ J K}^{-1}$$

$$1 \text{ MeV} = 1.6 \times 10^{-13} \text{ J} \quad 1 \text{ u} = 931.5 \text{ MeV}/c^2$$

$$1 \text{ year} = 3.154 \times 10^7 \text{ s}$$

$$m_{\text{H}} = 1.007825 \text{ u} \quad m_{\text{n}} = 1.008665 \text{ u}$$

$$m({}_2^4\text{He}) = 4.002603 \text{ u} \quad m_{\text{e}} = 0.000548 \text{ u}$$

13.1 Obtain the binding energy (in MeV) of a nitrogen nucleus (${}_{7}^{14}\text{N}$), given $m({}_{7}^{14}\text{N}) = 14.00307 \text{ u}$

13.2 Obtain the binding energy of the nuclei ${}_{26}^{56}\text{Fe}$ and ${}_{83}^{209}\text{Bi}$ in units of MeV from the following data:

$$m({}_{26}^{56}\text{Fe}) = 55.934939 \text{ u} \quad m({}_{83}^{209}\text{Bi}) = 208.980388 \text{ u}$$

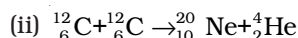
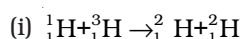
13.3 A given coin has a mass of 3.0 g. Calculate the nuclear energy that would be required to separate all the neutrons and protons from each other. For simplicity assume that the coin is entirely made of ${}_{29}^{63}\text{Cu}$ atoms (of mass 62.92960 u).

13.4 Obtain approximately the ratio of the nuclear radii of the gold isotope ${}_{79}^{197}\text{Au}$ and the silver isotope ${}_{47}^{107}\text{Ag}$.

13.5 The Q value of a nuclear reaction $A + b \rightarrow C + d$ is defined by

$$Q = [m_A + m_b - m_C - m_d]c^2$$

where the masses refer to the respective nuclei. Determine from the given data the Q -value of the following reactions and state whether the reactions are exothermic or endothermic.



Atomic masses are given to be

$$m({}_1^2\text{H}) = 2.014102 \text{ u}$$

$$m({}_1^3\text{H}) = 3.016049 \text{ u}$$

$$m({}_6^{12}\text{C}) = 12.000000 \text{ u}$$

$$m({}_{10}^{20}\text{Ne}) = 19.992439 \text{ u}$$

13.6 Suppose, we think of fission of a ${}_{26}^{56}\text{Fe}$ nucleus into two equal fragments, ${}_{13}^{28}\text{Al}$. Is the fission energetically possible? Argue by working out Q of the process. Given $m({}_{26}^{56}\text{Fe}) = 55.93494 \text{ u}$ and $m({}_{13}^{28}\text{Al}) = 27.98191 \text{ u}$.



- 13.7** The fission properties of ${}_{94}^{239}\text{Pu}$ are very similar to those of ${}_{92}^{235}\text{U}$. The average energy released per fission is 180 MeV. How much energy, in MeV, is released if all the atoms in 1 kg of pure ${}_{94}^{239}\text{Pu}$ undergo fission?
- 13.8** How long can an electric lamp of 100W be kept glowing by fusion of 2.0 kg of deuterium? Take the fusion reaction as
- $${}^2_1\text{H} + {}^2_1\text{H} \rightarrow {}^3_2\text{He} + n + 3.27 \text{ MeV}$$
- 13.9** Calculate the height of the potential barrier for a head on collision of two deuterons. (Hint: The height of the potential barrier is given by the Coulomb repulsion between the two deuterons when they just touch each other. Assume that they can be taken as hard spheres of radius 2.0 fm.)
- 13.10** From the relation $R = R_0 A^{1/3}$, where R_0 is a constant and A is the mass number of a nucleus, show that the nuclear matter density is nearly constant (i.e. independent of A).



Chapter Fourteen

SEMICONDUCTOR ELECTRONICS: MATERIALS, DEVICES AND SIMPLE CIRCUITS

14.1 INTRODUCTION

Devices in which a controlled flow of electrons can be obtained are the basic *building blocks* of all the electronic circuits. Before the discovery of transistor in 1948, such devices were mostly vacuum tubes (also called valves) like the vacuum diode which has two electrodes, viz., anode (often called plate) and cathode; triode which has three electrodes – cathode, plate and grid; tetrode and pentode (respectively with 4 and 5 electrodes). In a vacuum tube, the electrons are supplied by a heated cathode and the controlled flow of these electrons *in vacuum* is obtained by varying the voltage between its different electrodes. Vacuum is required in the inter-electrode space; otherwise the moving electrons may lose their energy on collision with the air molecules in their path. In these devices the electrons can flow only from the cathode to the anode (i.e., only in one direction). Therefore, such devices are generally referred to as *valves*. These vacuum tube devices are bulky, consume high power, operate generally at high voltages (~ 100 V) and have limited life and low reliability. The seed of the development of modern *solid-state semiconductor electronics* goes back to 1930's when it was realised that some solid-state semiconductors and their junctions offer the possibility of controlling the number and the direction of flow of charge carriers through them. Simple excitations like light, heat or small applied voltage can change the number of mobile charges in a semiconductor. Note that the supply

and flow of charge carriers in the semiconductor devices are *within the solid itself*, while in the earlier vacuum tubes/valves, the mobile electrons were obtained from a heated cathode and they were made to flow in an *evacuated* space or vacuum. No external heating or large evacuated space is required by the semiconductor devices. They are small in size, consume low power, operate at low voltages and have long life and high reliability. Even the Cathode Ray Tubes (CRT) used in television and computer monitors which work on the principle of vacuum tubes are being replaced by Liquid Crystal Display (LCD) monitors with supporting solid state electronics. Much before the full implications of the semiconductor devices was formally understood, a naturally occurring crystal of *galena* (Lead sulphide, PbS) with a metal point contact attached to it was used as *detector* of radio waves.

In the following sections, we will introduce the basic concepts of semiconductor physics and discuss some semiconductor devices like junction diodes (a 2-electrode device) and bipolar junction transistor (a 3-electrode device). A few circuits illustrating their applications will also be described.

14.2 CLASSIFICATION OF METALS, CONDUCTORS AND SEMICONDUCTORS

On the basis of conductivity

On the basis of the relative values of electrical conductivity (σ) or resistivity ($\rho = 1/\sigma$), the solids are broadly classified as:

(i) **Metals:** They possess very low resistivity (or high conductivity).

$$\rho \sim 10^{-2} - 10^{-8} \Omega \text{ m}$$

$$\sigma \sim 10^2 - 10^8 \text{ S m}^{-1}$$

(ii) **Semiconductors:** They have resistivity or conductivity intermediate to metals and insulators.

$$\rho \sim 10^{-5} - 10^6 \Omega \text{ m}$$

$$\sigma \sim 10^5 - 10^{-6} \text{ S m}^{-1}$$

(iii) **Insulators:** They have high resistivity (or low conductivity).

$$\rho \sim 10^{11} - 10^{19} \Omega \text{ m}$$

$$\sigma \sim 10^{-11} - 10^{-19} \text{ S m}^{-1}$$

The values of ρ and σ given above are indicative of magnitude and could well go outside the ranges as well. Relative values of the resistivity are not the only criteria for distinguishing metals, insulators and semiconductors from each other. There are some other differences, which will become clear as we go along in this chapter.

Our interest in this chapter is in the study of semiconductors which could be:

(i) *Elemental semiconductors:* Si and Ge

(ii) *Compound semiconductors:* Examples are:

- Inorganic: CdS, GaAs, CdSe, InP, etc.
- Organic: anthracene, doped phthalocyanines, etc.
- Organic polymers: polypyrrole, polyaniline, polythiophene, etc.

Most of the currently available semiconductor devices are based on elemental semiconductors Si or Ge and compound *inorganic* semiconductors. However, after 1990, a few semiconductor devices using

organic semiconductors and semiconducting polymers have been developed signalling the birth of a futuristic technology of polymer-electronics and molecular-electronics. In this chapter, we will restrict ourselves to the study of inorganic semiconductors, particularly elemental semiconductors Si and Ge. The general concepts introduced here for discussing the elemental semiconductors, by-and-large, apply to most of the compound semiconductors as well.

On the basis of energy bands

According to the Bohr atomic model, in an *isolated atom* the energy of any of its electrons is decided by the orbit in which it revolves. But when the atoms come together to form a solid they are close to each other. So the outer orbits of electrons from neighbouring atoms would come very close or could even overlap. This would make the nature of electron motion in a solid very different from that in an isolated atom.

Inside the crystal each electron has a unique position and no two electrons see exactly the same pattern of surrounding charges. Because of this, each electron will have a different *energy level*. These different energy levels with continuous energy variation form what are called *energy bands*. The energy band which includes the energy levels of the valence electrons is called the *valence band*. The energy band above the valence band is called the *conduction band*. With no external energy, all the valence electrons will reside in the valence band. If the lowest level in the conduction band happens to be lower than the highest level of the valence band, the electrons from the valence band can easily move into the conduction band. Normally the conduction band is empty. But when it overlaps on the valence band electrons can move freely into it. This is the case with metallic conductors.

If there is some gap between the conduction band and the valence band, electrons in the valence band all remain bound and no free electrons are available in the conduction band. This makes the material an insulator. But some of the electrons from the valence band may gain external energy to cross the gap between the conduction band and the valence band. Then these electrons will move into the conduction band. At the same time they will create vacant energy levels in the valence band where other valence electrons can move. Thus the process creates the possibility of conduction due to electrons in conduction band as well as due to vacancies in the valence band.

Let us consider what happens in the case of Si or Ge crystal containing N atoms. For Si, the outermost orbit is the third orbit ($n = 3$), while for Ge it is the fourth orbit ($n = 4$). The number of electrons in the outermost orbit is 4 ($2s$ and $2p$ electrons). Hence, the total number of outer electrons in the crystal is $4N$. The maximum possible number of electrons in the outer orbit is 8 ($2s + 6p$ electrons). So, for the $4N$ valence electrons there are $8N$ available energy states. These $8N$ discrete energy levels can either form a continuous band or they may be grouped in different bands depending upon the distance between the atoms in the crystal (see box on Band Theory of Solids).

At the distance between the atoms in the crystal lattices of Si and Ge, the energy band of these $8N$ states is split apart into two which are separated by an *energy gap* E_g (Fig. 14.1). The lower band which is completely occupied by the $4N$ valence electrons at temperature of absolute zero is the *valence band*. The other band consisting of $4N$ energy states, called the *conduction band*, is completely empty at absolute zero.

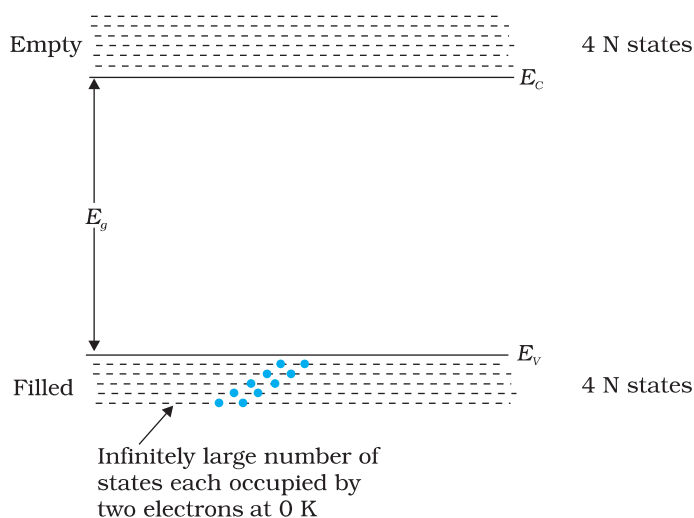


FIGURE 14.1 The energy band positions in a semiconductor at 0 K. The upper band, called the conduction band, consists of infinitely large number of closely spaced energy states. The lower band, called the valence band, consists of closely spaced completely filled energy states.

The lowest energy level in the conduction band is shown as E_C and highest energy level in the valence band is shown as E_V . Above E_C and below E_V there are a large number of closely spaced energy levels, as shown in Fig. 14.1.

The gap between the top of the valence band and bottom of the conduction band is called the *energy band gap* (Energy gap E_g). It may be large, small, or zero, depending upon the material. These different situations, are depicted in Fig. 14.2 and discussed below:

Case I: This refers to a situation, as shown in Fig. 14.2(a). One can have a metal either when the conduction band is partially filled and the balanced band is partially empty or when the conduction and valance bands overlap. When there is overlap electrons from valence band can easily move into the conduction band. This situation makes a large number of

electrons available for electrical conduction. When the valence band is partially empty, electrons from its lower level can move to higher level making conduction possible. Therefore, the resistance of such materials is low or the conductivity is high.

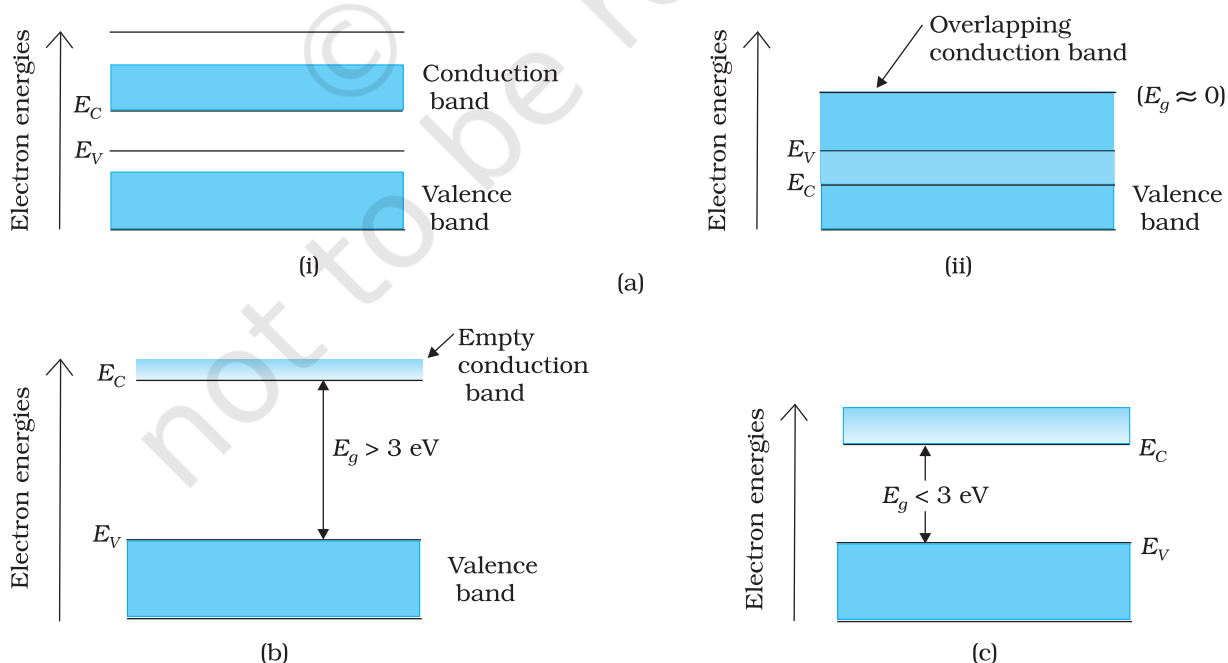


FIGURE 14.2 Difference between energy bands of (a) metals, (b) insulators and (c) semiconductors.

Case II: In this case, as shown in Fig. 14.2(b), a large band gap E_g exists ($E_g > 3$ eV). There are no electrons in the conduction band, and therefore no electrical conduction is possible. Note that the energy gap is so large that electrons cannot be excited from the valence band to the conduction band by thermal excitation. This is the case of *insulators*.

Case III: This situation is shown in Fig. 14.2(c). Here a finite but small band gap ($E_g < 3$ eV) exists. Because of the small band gap, at room temperature some electrons from valence band can acquire enough energy to cross the energy gap and enter the *conduction band*. These electrons (though small in numbers) can move in the conduction band. Hence, the resistance of *semiconductors* is not as high as that of the insulators.

In this section we have made a broad classification of metals, conductors and semiconductors. In the section which follows you will learn the conduction process in semiconductors.

14.3 INTRINSIC SEMICONDUCTOR

We shall take the most common case of Ge and Si whose lattice structure is shown in Fig. 14.3. These structures are called the *diamond-like structures*. Each atom is surrounded by four nearest neighbours. We know that Si and Ge have four valence electrons. In its crystalline structure, every Si or Ge atom tends to *share* one of its four valence electrons with each of its four nearest neighbour atoms, and also to *take share* of one electron from each such neighbour. These shared electron pairs are referred to as forming a *covalent bond* or simply a *valence bond*. The two shared electrons can be assumed to shuttle back-and-forth between the associated atoms holding them together strongly. Figure 14.4 schematically shows the 2-dimensional representation of Si or Ge structure shown in Fig. 14.3 which overemphasises the covalent bond. It shows an idealised picture in which no bonds are broken (all bonds are intact). Such a situation arises at low temperatures. As the temperature increases, more thermal energy becomes available to these electrons and some of these electrons may break-away (becoming *free* electrons contributing to conduction). The thermal energy effectively ionises only a few atoms in the crystalline lattice and creates a *vacancy* in the bond as shown in Fig. 14.5(a). The neighbourhood, from which the free electron (with charge $-q$) has come out leaves a vacancy with an effective charge $(+q)$. This *vacancy* with the effective positive electronic charge is called a *hole*. The hole behaves as an *apparent free particle* with effective positive charge.

In intrinsic semiconductors, the number of free electrons, n_e is equal to the number of holes, n_h . That is

$$n_e = n_h = n_i \quad (14.1)$$

where n_i is called intrinsic carrier concentration.

Semiconductors possess the unique property in which, apart from electrons, the holes also move. Suppose there is a hole at site 1 as shown

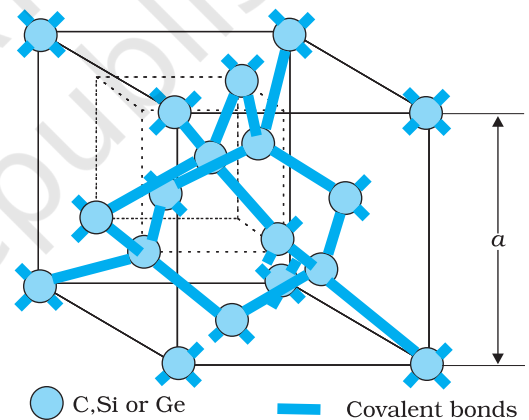


FIGURE 14.3 Three-dimensional diamond-like crystal structure for Carbon, Silicon or Germanium with respective lattice spacing a equal to 3.56, 5.43 and 5.66 Å.

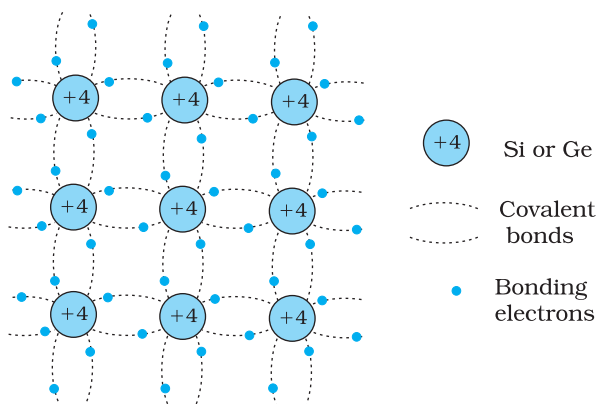


FIGURE 14.4 Schematic two-dimensional representation of Si or Ge structure showing covalent bonds at low temperature (all bonds intact). +4 symbol indicates inner cores of Si or Ge.

in Fig. 14.5(a). The movement of holes can be visualised as shown in Fig. 14.5(b). An electron from the covalent bond at site 2 may jump to the vacant site 1 (hole). Thus, after such a jump, the hole is at site 2 and the site 1 has now an electron. Therefore, apparently, the hole has moved from site 1 to site 2. Note that the electron originally set free [Fig. 14.5(a)] is not involved in this process of hole motion. The free electron moves completely independently as conduction electron and gives rise to an electron current, I_e under an applied electric field. Remember that the motion of hole is only a convenient way of describing the actual motion of *bound* electrons, whenever there is an empty bond anywhere in the crystal. Under the action of an electric field, these holes move towards negative potential giving the hole current, I_h . The total current, I is thus the sum of the electron current I_e and the hole current I_h :

$$I = I_e + I_h \quad (14.2)$$

It may be noted that apart from the *process of generation* of conduction electrons and holes, a simultaneous *process of recombination* occurs in which the electrons *recombine* with the holes. At equilibrium, the rate of generation is equal to the rate of recombination of charge carriers. The recombination occurs due to an electron colliding with a hole.

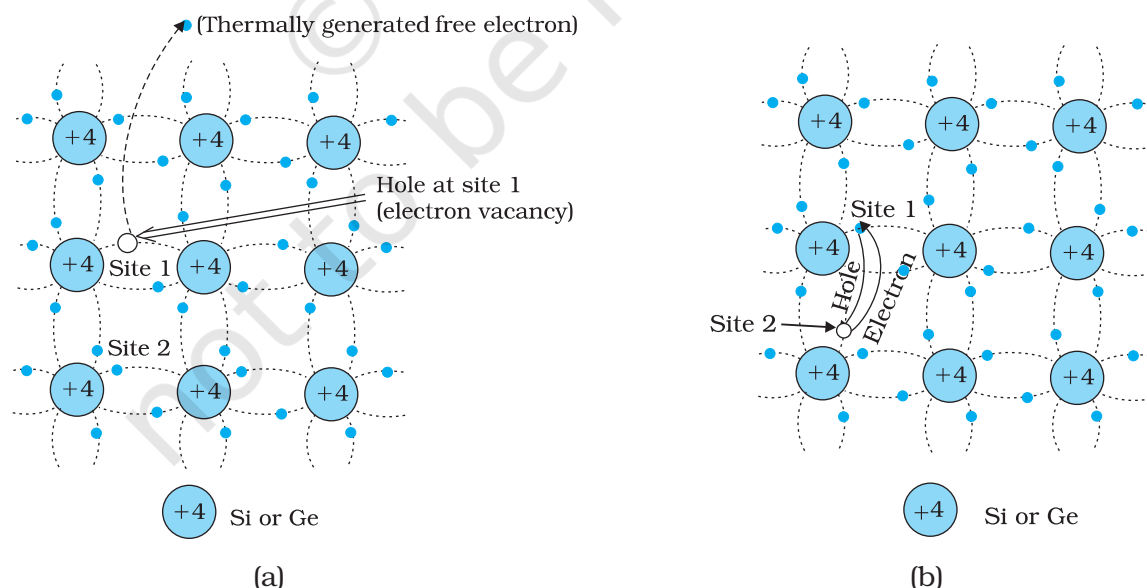


FIGURE 14.5 (a) Schematic model of generation of hole at site 1 and conduction electron due to thermal energy at moderate temperatures. (b) Simplified representation of possible thermal motion of a hole. The electron from the lower left hand covalent bond (site 2) goes to the earlier hole site 1, leaving a hole at its site indicating an apparent movement of the hole from site 1 to site 2.

An intrinsic semiconductor will behave like an insulator at $T = 0\text{ K}$ as shown in Fig. 14.6(a). It is the thermal energy at higher temperatures ($T > 0\text{ K}$), which excites some electrons from the valence band to the conduction band. These thermally excited electrons at $T > 0\text{ K}$, partially occupy the conduction band. Therefore, the energy-band diagram of an intrinsic semiconductor will be as shown in Fig. 14.6(b). Here, some electrons are shown in the conduction band. These have come from the valence band leaving equal number of holes there.

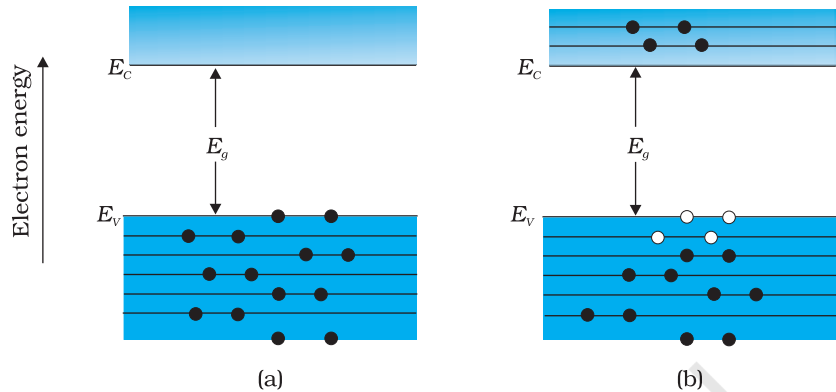


FIGURE 14.6 (a) An intrinsic semiconductor at $T = 0\text{ K}$ behaves like insulator. (b) At $T > 0\text{ K}$, four thermally generated electron-hole pairs. The filled circles (•) represent electrons and empty circles (○) represent holes.

Example 14.1 C, Si and Ge have same lattice structure. Why is C insulator while Si and Ge intrinsic semiconductors?

Solution The 4 bonding electrons of C, Si or Ge lie, respectively, in the second, third and fourth orbit. Hence, energy required to take out an electron from these atoms (i.e., ionisation energy E_g) will be least for Ge, followed by Si and highest for C. Hence, number of free electrons for conduction in Ge and Si are significant but negligibly small for C.

EXAMPLE 14.1

14.4 EXTRINSIC SEMICONDUCTOR

The conductivity of an intrinsic semiconductor depends on its temperature, but at room temperature its conductivity is very low. As such, no important electronic devices can be developed using these semiconductors. Hence there is a necessity of improving their conductivity. This can be done by making use of impurities.

When a small amount, say, a few parts per million (ppm), of a suitable impurity is added to the pure semiconductor, the conductivity of the semiconductor is increased manifold. Such materials are known as *extrinsic semiconductors* or *impurity semiconductors*. The deliberate addition of a desirable impurity is called *doping* and the impurity atoms are called *dopants*. Such a material is also called a *doped semiconductor*. The dopant has to be such that it does not distort the original pure semiconductor lattice. It occupies only a very few of the original semiconductor atom sites in the crystal. A necessary condition to attain this is that the sizes of the dopant and the semiconductor atoms should be nearly the same.

There are two types of dopants used in doping the tetravalent Si or Ge:

- (i) Pentavalent (valency 5); like Arsenic (As), Antimony (Sb), Phosphorous (P), etc.

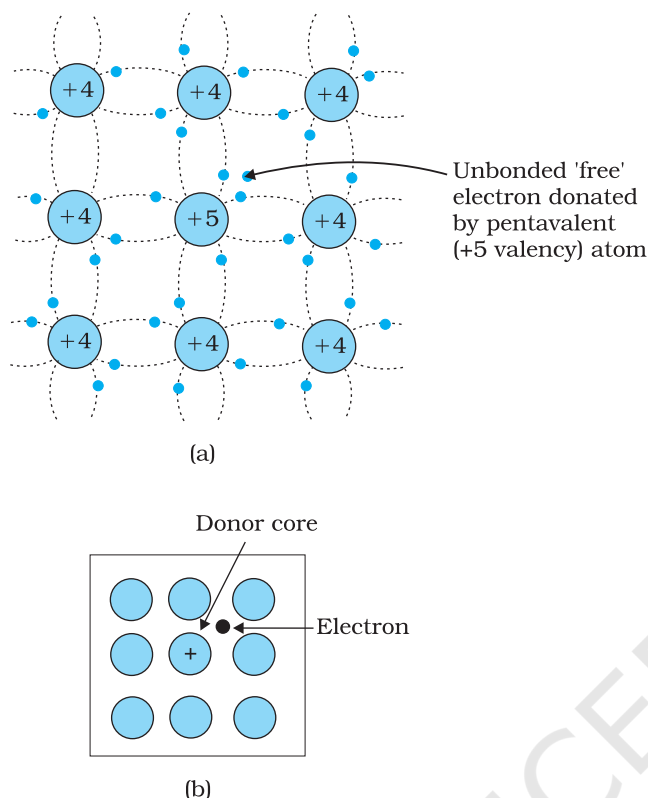


FIGURE 14.7 (a) Pentavalent donor atom (As, Sb, P, etc.) doped for tetravalent Si or Ge giving n-type semiconductor, and (b) Commonly used schematic representation of n-type material which shows only the fixed cores of the substituent donors with one additional effective positive charge and its associated extra electron.

(ii) Trivalent (valency 3); like Indium (In), Boron (B), Aluminium (Al), etc.

We shall now discuss how the doping changes the number of charge carriers (and hence the conductivity) of semiconductors. Si or Ge belongs to the fourth group in the Periodic table and, therefore, we choose the dopant element from nearby fifth or third group, expecting and taking care that the size of the dopant atom is nearly the same as that of Si or Ge. Interestingly, the pentavalent and trivalent dopants in Si or Ge give two entirely different types of semiconductors as discussed below.

(i) n-type semiconductor

Suppose we dope Si or Ge with a pentavalent element as shown in Fig. 14.7. When an atom of +5 valency element occupies the position of an atom in the crystal lattice of Si, four of its electrons bond with the four silicon neighbours while the fifth remains very weakly bound to its parent atom. This is because the four electrons participating in bonding are seen as part of the effective core of the atom by the fifth electron. As a result the ionisation energy required to set this electron free is very small and even at room temperature it will be free to move in the lattice of the semiconductor. For example, the energy required is ~ 0.01 eV for germanium, and 0.05 eV for silicon, to separate this

electron from its atom. This is in contrast to the energy required to jump the forbidden band (about 0.72 eV for germanium and about 1.1 eV for silicon) at room temperature in the intrinsic semiconductor. Thus, the pentavalent dopant is donating one extra electron for conduction and hence is known as *donor* impurity. The number of electrons made available for conduction by dopant atoms depends strongly upon the doping level and is independent of any increase in ambient temperature. On the other hand, the number of free electrons (with an equal number of holes) generated by Si atoms, increases weakly with temperature.

In a doped semiconductor the total number of conduction electrons n_e is due to the electrons contributed by donors and those generated intrinsically, while the total number of holes n_h is only due to the holes from the intrinsic source. But the rate of recombination of holes would increase due to the increase in the number of electrons. As a result, the number of holes would get reduced further.

Thus, with proper level of doping the number of conduction electrons can be made much larger than the number of holes. Hence in an extrinsic

semiconductor doped with pentavalent impurity, electrons become the *majority carriers* and holes the *minority carriers*. These semiconductors are, therefore, known as *n-type semiconductors*. For n-type semiconductors, we have,

$$n_e \gg n_h \quad (14.3)$$

(ii) p-type semiconductor

This is obtained when Si or Ge is doped with a trivalent impurity like Al, B, In, etc. The dopant has one valence electron less than Si or Ge and, therefore, this atom can form covalent bonds with neighbouring three Si atoms but does not have any electron to offer to the fourth Si atom. So the bond between the fourth neighbour and the trivalent atom has a vacancy or hole as shown in Fig. 14.8. Since the neighbouring Si atom in the lattice wants an electron in place of a hole, an electron in the outer orbit of an atom in the neighbourhood may jump to fill this vacancy, leaving a vacancy or hole at its own site. Thus the *hole* is available for conduction. Note that the trivalent foreign atom becomes effectively negatively charged when it shares fourth electron with neighbouring Si atom. Therefore, the dopant atom of p-type material can be treated as *core of one negative charge* along with its associated hole as shown in Fig. 14.8(b). It is obvious that one *acceptor* atom gives one *hole*. These holes are in addition to the intrinsically generated holes while the source of conduction electrons is only intrinsic generation. Thus, for such a material, the holes are the majority carriers and electrons are minority carriers. Therefore, extrinsic semiconductors doped with trivalent impurity are called *p-type semiconductors*. For p-type semiconductors, the recombination process will reduce the number (n_i) of intrinsically generated electrons to n_e . We have, for p-type semiconductors

$$n_h \gg n_e \quad (14.4)$$

Note that *the crystal maintains an overall charge neutrality as the charge of additional charge carriers is just equal and opposite to that of the ionised cores in the lattice.*

In extrinsic semiconductors, because of the abundance of majority current carriers, the minority carriers produced thermally have more chance of meeting majority carriers and thus getting destroyed. Hence, the dopant, by adding a large number of current carriers of one type, which become the majority carriers, indirectly helps to reduce the intrinsic concentration of minority carriers.

The semiconductor's energy band structure is affected by doping. In the case of extrinsic semiconductors, additional energy states due to donor impurities (E_D) and acceptor impurities (E_A) also exist. In the energy band diagram of n-type Si semiconductor, the donor energy level E_D is slightly below the bottom E_C of the conduction band and electrons from this level move into the conduction band with very small supply of energy. At room temperature, most of the donor atoms get ionised but very few ($\sim 10^{12}$) atoms of Si get ionised. So the conduction band will have most electrons coming from the donor impurities, as shown in Fig. 14.9(a). Similarly,

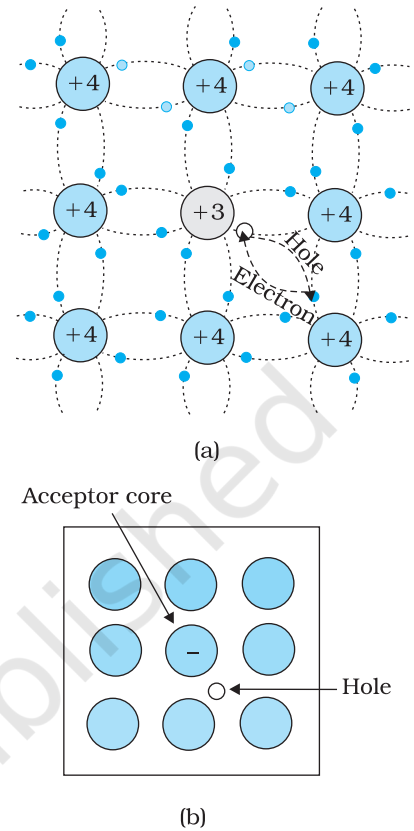


FIGURE 14.8 (a) Trivalent acceptor atom (In, Al, B etc.) doped in tetravalent Si or Ge lattice giving p-type semiconductor. (b) Commonly used schematic representation of p-type material which shows only the fixed core of the substituent acceptor with one effective additional negative charge and its associated hole.

for p-type semiconductor, the acceptor energy level E_A is slightly above the top E_V of the valence band as shown in Fig. 14.9(b). With very small supply of energy an electron from the valence band can jump to the level E_A and ionise the acceptor negatively. (Alternately, we can also say that with very small supply of energy the hole from level E_A sinks down into the valence band. Electrons rise up and holes fall down when they gain external energy.) At room temperature, most of the acceptor atoms get ionised leaving holes in the valence band. Thus at room temperature the density of holes in the valence band is predominantly due to impurity in the extrinsic semiconductor. The electron and hole concentration in a semiconductor in thermal equilibrium is given by

$$n_e n_h = n_i^2 \quad (14.5)$$

Though the above description is grossly approximate and hypothetical, it helps in understanding the difference between metals, insulators and semiconductors (extrinsic and intrinsic) in a simple manner. The difference in the resistivity of C, Si and Ge depends upon the energy gap between their conduction and valence bands. For C (diamond), Si and Ge, the energy gaps are 5.4 eV, 1.1 eV and 0.7 eV, respectively. Sn also is a group IV element but it is a metal because the energy gap in its case is 0 eV.

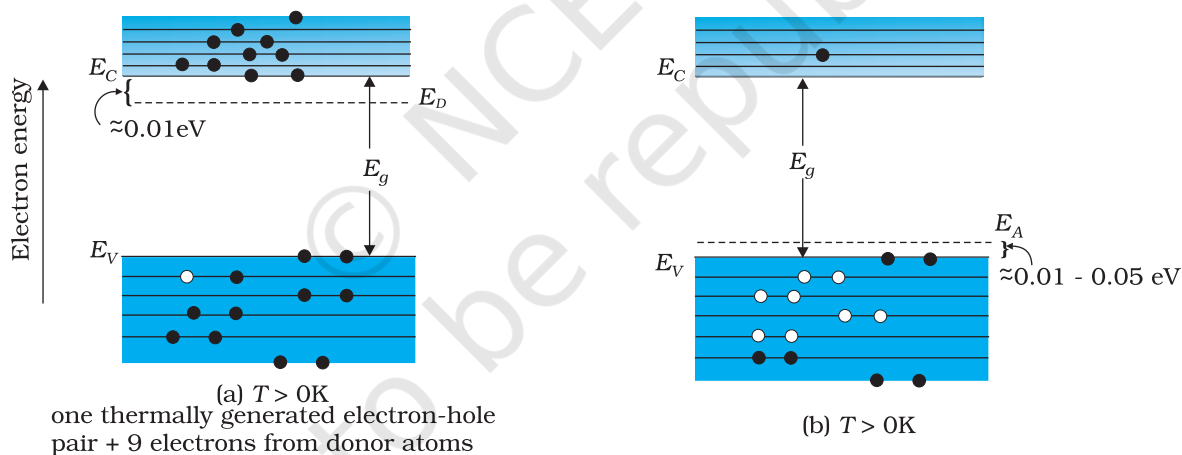


FIGURE 14.9 Energy bands of (a) n-type semiconductor at $T > 0K$, (b) p-type semiconductor at $T > 0K$.

EXAMPLE 14.2

Example 14.2 Suppose a pure Si crystal has 5×10^{28} atoms m^{-3} . It is doped by 1 ppm concentration of pentavalent As. Calculate the number of electrons and holes. Given that $n_i = 1.5 \times 10^{16} m^{-3}$.

Solution Note that thermally generated electrons ($n_i \sim 10^{16} m^{-3}$) are negligibly small as compared to those produced by doping.

Therefore, $n_e \approx N_D$.

Since $n_e n_h = n_i^2$, The number of holes

$$n_h = (2.25 \times 10^{32}) / (5 \times 10^{22})$$

$$\sim 4.5 \times 10^9 m^{-3}$$

14.5 p-n JUNCTION

A p-n junction is the basic building block of many semiconductor devices like diodes, transistor, etc. A clear understanding of the junction behaviour is important to analyse the working of other semiconductor devices. We will now try to understand how a junction is formed and how the junction behaves under the influence of external applied voltage (also called *bias*).

14.5.1 p-n junction formation

Consider a thin p-type silicon (p-Si) semiconductor wafer. By adding precisely a small quantity of pentavalent impurity, part of the p-Si wafer can be converted into n-Si. There are several processes by which a semiconductor can be formed. The wafer now contains p-region and n-region and a metallurgical junction between p-, and n- region.

Two important processes occur during the formation of a p-n junction: *diffusion* and *drift*. We know that in an n-type semiconductor, the concentration of electrons (number of electrons per unit volume) is more compared to the concentration of holes. Similarly, in a p-type semiconductor, the concentration of holes is more than the concentration of electrons. During the formation of p-n junction, and due to the concentration gradient across p-, and n- sides, holes diffuse from p-side to n-side ($p \rightarrow n$) and electrons diffuse from n-side to p-side ($n \rightarrow p$). This motion of charge carries gives rise to diffusion current across the junction.

When an electron diffuses from $n \rightarrow p$, it leaves behind an ionised donor on n-side. This ionised donor (positive charge) is immobile as it is bonded to the surrounding atoms. As the electrons continue to diffuse from $n \rightarrow p$, a layer of positive charge (or positive space-charge region) on n-side of the junction is developed.

Similarly, when a hole diffuses from $p \rightarrow n$ due to the concentration gradient, it leaves behind an ionised acceptor (negative charge) which is immobile. As the holes continue to diffuse, a layer of negative charge (or negative space-charge region) on the p-side of the junction is developed. This space-charge region on either side of the junction together is known as *depletion region* as the electrons and holes taking part in the initial movement across the junction *depleted* the region of its free charges (Fig. 14.10). The thickness of depletion region is of the order of one-tenth of a micrometre. Due to the positive space-charge region on n-side of the junction and negative space charge region on p-side of the junction, an electric field directed from positive charge towards negative charge develops. Due to this field, an electron on p-side of the junction moves to n-side and a hole on n-side of the junction moves to p-side. The motion of charge carriers due to the electric field is called drift. Thus a drift current, which is opposite in direction to the diffusion current (Fig. 14.10) starts.

Formation and working of p-n junction diode
<http://hyperphysics.phy-astr.gsu.edu/hbase/solids/pnjun.html>

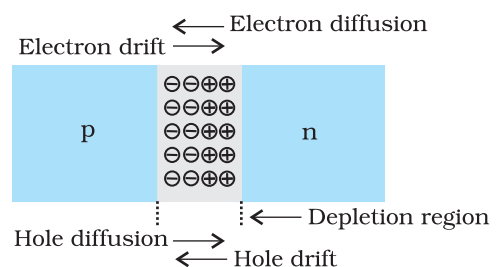
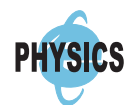


FIGURE 14.10 p-n junction formation process.

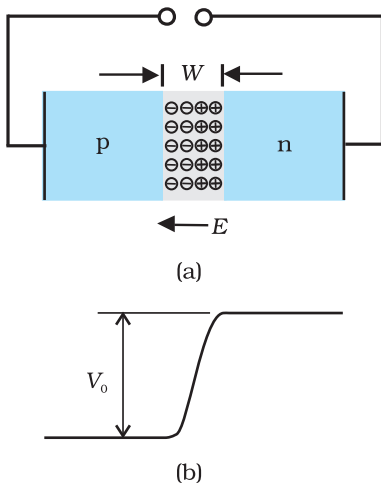


FIGURE 14.11 (a) Diode under equilibrium ($V = 0$), (b) Barrier potential under no bias.

Initially, diffusion current is large and drift current is small. As the diffusion process continues, the space-charge regions on either side of the junction extend, thus increasing the electric field strength and hence drift current. This process continues until the diffusion current equals the drift current. Thus a p-n junction is formed. In a p-n junction under equilibrium there is *no net current*.

The loss of electrons from the n-region and the gain of electron by the p-region causes a difference of potential across the junction of the two regions. The polarity of this potential is such as to oppose further flow of carriers so that a condition of equilibrium exists. Figure 14.11 shows the p-n junction at equilibrium and the potential across the junction. The n-material has lost electrons, and p material has acquired electrons. The n material is thus positive relative to the p material. Since this potential tends to prevent the movement of electron from the n region into the p region, it is often called a *barrier potential*.

EXAMPLE 14.3

Example 14.3 Can we take one slab of p-type semiconductor and physically join it to another n-type semiconductor to get p-n junction?

Solution No! Any slab, howsoever flat, will have roughness much larger than the inter-atomic crystal spacing (~ 2 to $3 \text{ \AA}$) and hence *continuous contact* at the atomic level will not be possible. The junction will behave as a *discontinuity* for the flowing charge carriers.

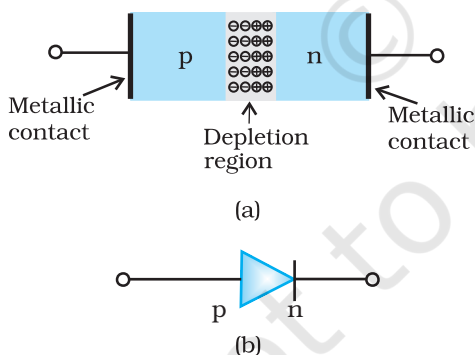


FIGURE 14.12 (a) Semiconductor diode, (b) Symbol for p-n junction diode.

14.6 SEMICONDUCTOR DIODE

A semiconductor diode [Fig. 14.12(a)] is basically a p-n junction with metallic contacts provided at the ends for the application of an external voltage. It is a two terminal device. A p-n junction diode is symbolically represented as shown in Fig. 14.12(b).

The direction of arrow indicates the conventional direction of current (when the diode is under forward bias). The equilibrium barrier potential can be altered by applying an external voltage V across the diode. The situation of p-n junction diode under equilibrium (without bias) is shown in Fig. 14.11(a) and (b).

14.6.1 p-n junction diode under forward bias

When an external voltage V is applied across a semiconductor diode such that p-side is connected to the positive terminal of the battery and n-side to the negative terminal [Fig. 14.13(a)], it is said to be *forward biased*.

The applied voltage mostly drops across the depletion region and the voltage drop across the p-side and n-side of the junction is negligible. (This is because the resistance of the depletion region – a region where there are no charges – is very high compared to the resistance of n-side and p-side.) The direction of the applied voltage (V) is opposite to the

built-in potential V_0 . As a result, the depletion layer width decreases and the barrier height is reduced [Fig. 14.13(b)]. The effective barrier height under forward bias is $(V_0 - V)$.

If the applied voltage is small, the barrier potential will be reduced only slightly below the equilibrium value, and only a small number of carriers in the material—those that happen to be in the uppermost energy levels—will possess enough energy to cross the junction. So the current will be small. If we increase the applied voltage significantly, the barrier height will be reduced and more number of carriers will have the required energy. Thus the current increases.

Due to the applied voltage, electrons from n-side cross the depletion region and reach p-side (where they are minority carries). Similarly, holes from p-side cross the junction and reach the n-side (where they are minority carries). This process under forward bias is known as minority carrier injection. At the junction boundary, on each side, the minority carrier concentration increases significantly compared to the locations far from the junction.

Due to this concentration gradient, the injected electrons on p-side diffuse from the junction edge of p-side to the other end of p-side. Likewise, the injected holes on n-side diffuse from the junction edge of n-side to the other end of n-side (Fig. 14.14). This motion of charged carriers on either side gives rise to current. The total diode forward current is sum of hole diffusion current and conventional current due to electron diffusion. The magnitude of this current is usually in mA.

14.6.2 p-n junction diode under reverse bias

When an external voltage (V) is applied across the diode such that n-side is positive and p-side is negative, it is said to be *reverse biased* [Fig.14.15(a)]. The applied voltage mostly drops across the depletion region. The direction of applied voltage is same as the direction of barrier potential. As a result, the barrier height increases and the depletion region widens due to the change in the electric field. The effective barrier height under reverse bias is $(V_0 + V)$, [Fig. 14.15(b)]. This suppresses the flow of electrons from $n \rightarrow p$ and holes from $p \rightarrow n$. Thus, diffusion current, decreases enormously compared to the diode under forward bias.

The electric field direction of the junction is such that if electrons on p-side or holes on n-side in their random motion come close to the junction, they will be swept to its majority zone. This drift of carriers gives rise to current. The drift current is of the order of a few μA . This is quite low because it is due to the motion of carriers from their minority side to their majority side across the junction. The drift current is also there under forward bias but it is negligible (μA) when compared with current due to injected carriers which is usually in mA.

The diode reverse current is not very much dependent on the applied voltage. Even a small voltage is sufficient to sweep the minority carriers from one side of the junction to the other side of the junction. The current

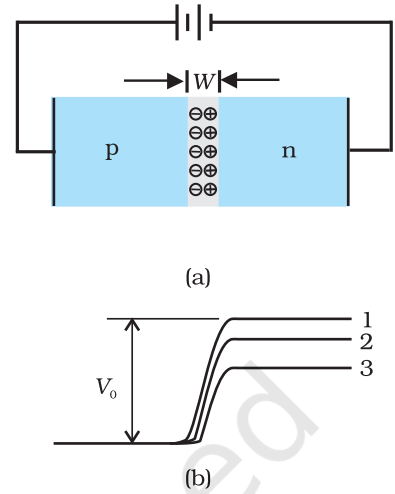


FIGURE 14.13 (a) p-n junction diode under forward bias, (b) Barrier potential (1) without battery, (2) Low battery voltage, and (3) High voltage battery.

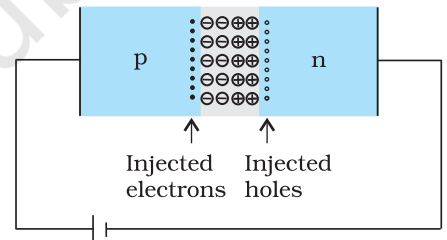


FIGURE 14.14 Forward bias minority carrier injection.

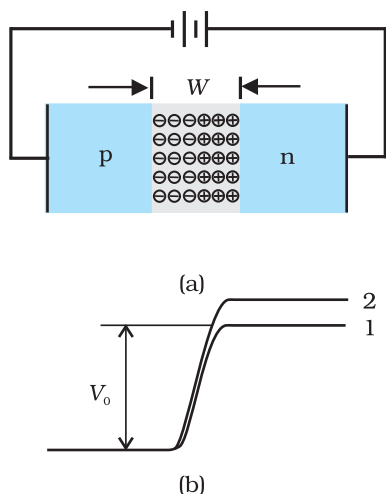


FIGURE 14.15 (a) Diode under reverse bias, (b) Barrier potential under reverse bias.

is not limited by the magnitude of the applied voltage but is limited due to the concentration of the minority carrier on either side of the junction.

The current under reverse bias is essentially voltage independent upto a critical reverse bias voltage, known as breakdown voltage (V_{br}). When $V = V_{br}$, the diode reverse current increases sharply. Even a slight increase in the bias voltage causes large change in the current. If the reverse current is not limited by an external circuit below the rated value (specified by the manufacturer) the p-n junction will get destroyed. Once it exceeds the rated value, the diode gets destroyed due to overheating. This can happen even for the diode under forward bias, if the forward current exceeds the rated value.

The circuit arrangement for studying the V - I characteristics of a diode, (i.e., the variation of current as a function of applied voltage) are shown in Fig. 14.16(a) and (b). The battery is connected to the diode through a potentiometer (or rheostat) so that the applied voltage to the diode can be changed. For different values of voltages, the value of the current is noted. A graph between V and I is obtained as in Fig. 14.16(c). Note that in forward bias measurement, we use a milliammeter since the expected current is large (as explained in the earlier section) while a micrometer is used in reverse bias to measure the current. You can see in Fig. 14.16(c) that in forward

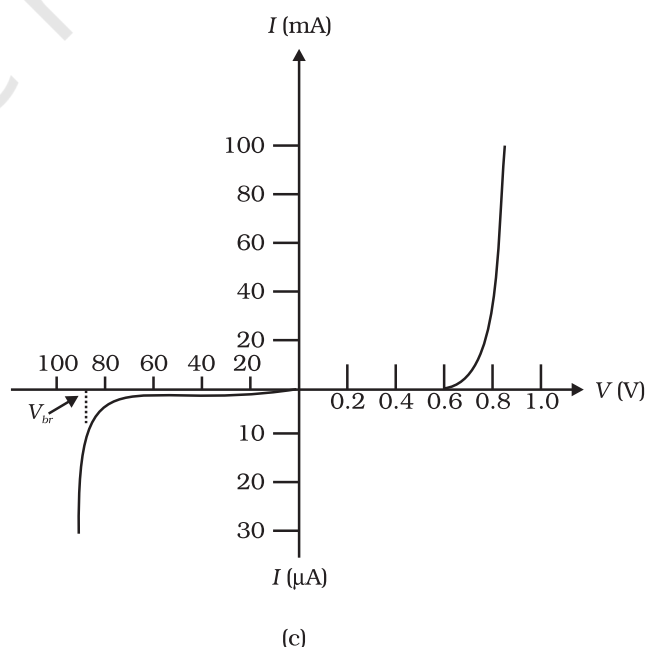
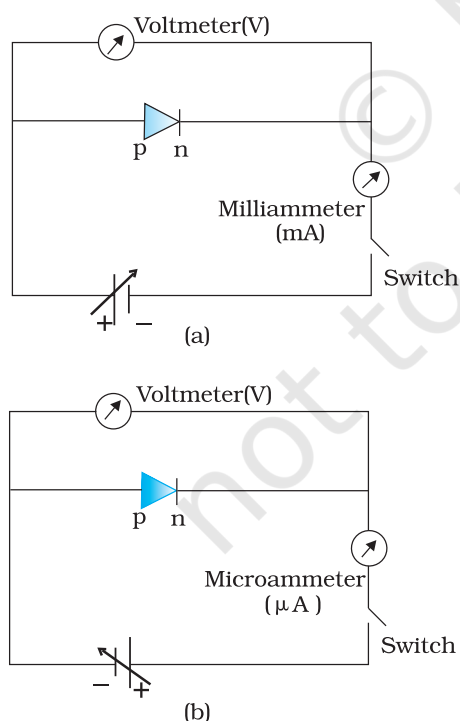


FIGURE 14.16 Experimental circuit arrangement for studying V - I characteristics of a p-n junction diode (a) in forward bias, (b) in reverse bias. (c) Typical V - I characteristics of a silicon diode.

bias, the current first increases very slowly, almost negligibly, till the voltage across the diode crosses a certain value. After the characteristic voltage, the diode current increases significantly (exponentially), even for a very small increase in the diode bias voltage. This voltage is called the *threshold voltage* or cut-in voltage ($\sim 0.2\text{V}$ for germanium diode and $\sim 0.7\text{V}$ for silicon diode).

For the diode in reverse bias, the current is very small ($\sim \mu\text{A}$) and almost remains constant with change in bias. It is called *reverse saturation current*. However, for special cases, at very high reverse bias (break down voltage), the current suddenly increases. This special action of the diode is discussed later in Section 14.8. The general purpose diode are not used beyond the reverse saturation current region.

The above discussion shows that the p-n junction diode primarily allows the flow of current only in one direction (forward bias). The forward bias resistance is low as compared to the reverse bias resistance. This property is used for rectification of ac voltages as discussed in the next section. For diodes, we define a quantity called *dynamic resistance* as the ratio of small change in voltage ΔV to a small change in current ΔI :

$$r_d = \frac{\Delta V}{\Delta I} \quad (14.6)$$

Example 14.4 The V-I characteristic of a silicon diode is shown in the Fig. 14.17. Calculate the resistance of the diode at (a) $I_D = 15\text{ mA}$ and (b) $V_D = -10\text{ V}$.

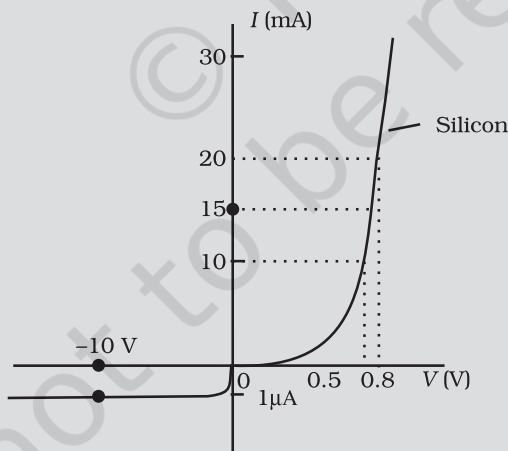


FIGURE 14.17

Solution Considering the diode characteristics as a straight line between $I = 10\text{ mA}$ to $I = 20\text{ mA}$ passing through the origin, we can calculate the resistance using Ohm's law.

(a) From the curve, at $I = 20\text{ mA}$, $V = 0.8\text{ V}$; $I = 10\text{ mA}$, $V = 0.7\text{ V}$

$$r_{fb} = \Delta V / \Delta I = 0.1\text{V} / 10\text{ mA} = 10\ \Omega$$

(b) From the curve at $V = -10\text{ V}$, $I = -1\ \mu\text{A}$,

Therefore,

$$r_{rb} = 10\text{ V} / 1\ \mu\text{A} = 1.0 \times 10^7\ \Omega$$

14.7 APPLICATION OF JUNCTION DIODE AS A RECTIFIER

From the V - I characteristic of a junction diode we see that it allows current to pass only when it is forward biased. So if an alternating voltage is applied across a diode the current flows only in that part of the cycle when the diode is forward biased. This property is used to *rectify* alternating voltages and the circuit used for this purpose is called a *rectifier*.

If an alternating voltage is applied across a diode in series with a load, a pulsating voltage will appear across the load only during the half cycles of the ac input during which the diode is forward biased. Such rectifier circuit, as shown in Fig. 14.18, is called a *half-wave rectifier*. The secondary of a transformer supplies the desired ac voltage across terminals A and B. When the voltage at A is positive, the diode is forward biased and it conducts. When A is negative, the diode is reverse-biased and it does not conduct. The reverse saturation current of a diode is negligible and can be considered equal to zero for practical purposes. (The reverse breakdown voltage of the diode must be sufficiently higher than the peak ac voltage at the secondary of the transformer to protect the diode from reverse breakdown.)

Therefore, in the positive *half-cycle* of ac there is a current through the load resistor R_L and we get an output voltage, as shown in Fig. 14.18(b), whereas there is no current in the negative half-cycle. In the next positive half-cycle, again we get the output voltage. Thus, the output voltage, though still varying, is restricted to *only one direction* and is said to be *rectified*. Since the rectified output of this circuit is only for half of the input ac wave it is called as *half-wave rectifier*.

The circuit using two diodes, shown in Fig. 14.19(a), gives output rectified voltage corresponding to both the positive as well as negative half of the ac cycle. Hence, it is known as *full-wave rectifier*. Here the p-side of the two diodes are connected to the ends of the secondary of the transformer. The n-side of the diodes are connected together and the output is taken between this common point of diodes and the midpoint of the secondary of the transformer. So for a full-wave rectifier the secondary of the transformer is provided with a centre tapping and so it is called *centre-tap transformer*. As can be seen from Fig. 14.19(c) the voltage rectified by each diode is only half the total secondary voltage. Each diode rectifies only for half the cycle, but the two do so for alternate cycles. Thus, the output between their common terminals and the centre-tap of the transformer becomes a full-wave rectifier output. (Note that there is another circuit of full wave rectifier which does not need a centre-tap transformer but needs four diodes.) Suppose the input voltage to A

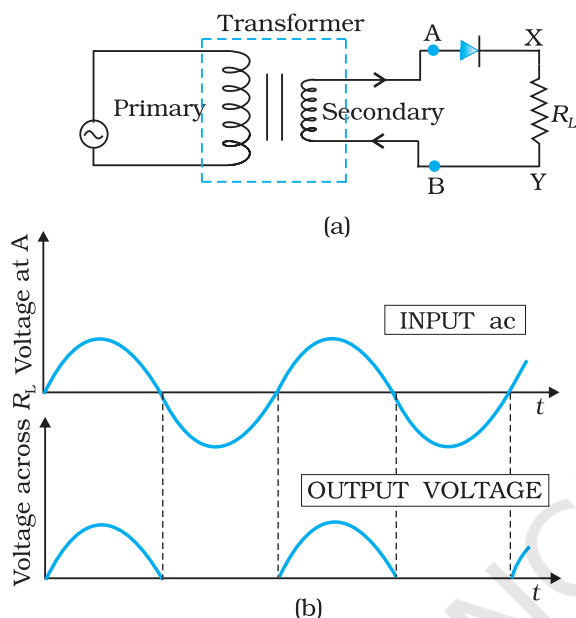


FIGURE 14.18 (a) Half-wave rectifier circuit, (b) Input ac voltage and output voltage waveforms from the rectifier circuit.

with respect to the centre tap at any instant is positive. It is clear that, at that instant, voltage at B being out of phase will be negative as shown in Fig. 14.19(b). So, diode D_1 gets forward biased and conducts (while D_2 being reverse biased is not conducting). Hence, during this positive half cycle we get an output current (and a output voltage across the load resistor R_L) as shown in Fig. 14.19(c). In the course of the ac cycle when the voltage at A becomes negative with respect to centre tap, the voltage at B would be positive. In this part of the cycle diode D_1 would not conduct but diode D_2 would, giving an output current and output voltage (across R_L) during the negative half cycle of the input ac. Thus, we get output voltage during both the positive as well as the negative half of the cycle. Obviously, this is a more efficient circuit for getting rectified voltage or current than the half-wave rectifier.

The rectified voltage is in the form of pulses of the shape of half sinusoids. Though it is unidirectional it does not have a steady value. To get steady dc output from the pulsating voltage normally a capacitor is connected across the output terminals (parallel to the load R_L). One can also use an inductor in series with R_L for the same purpose. Since these additional circuits appear to *filter* out the *ac ripple* and give a *pure dc* voltage, so they are called filters.

Now we shall discuss the role of capacitor in filtering. When the voltage across the capacitor is rising, it gets charged. If there is no external load, it remains charged to the peak voltage of the rectified output. When there is a load, it gets discharged through the load and the voltage across it begins to fall. In the next half-cycle of rectified output it again gets charged to the peak value (Fig. 14.20). The rate of fall of the voltage across the capacitor depends inversely upon the product of capacitance C and the effective resistance R_L used in the circuit and is called the *time constant*. To make the time constant large value of C should be large. So capacitor input filters use large capacitors. The *output voltage* obtained by using capacitor input filter is nearer to the *peak voltage* of the rectified voltage. This type of filter is most widely used in power supplies.

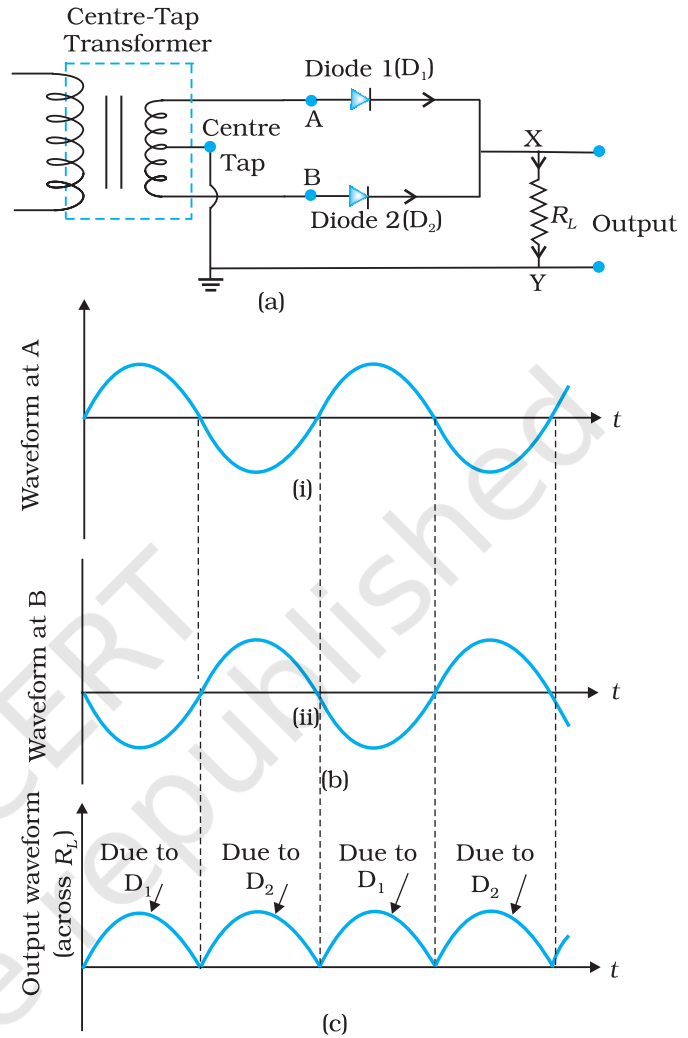


FIGURE 14.19 (a) A Full-wave rectifier circuit; (b) Input wave forms given to the diode D_1 at A and to the diode D_2 at B; (c) Output waveform across the load R_L connected in the full-wave rectifier circuit.

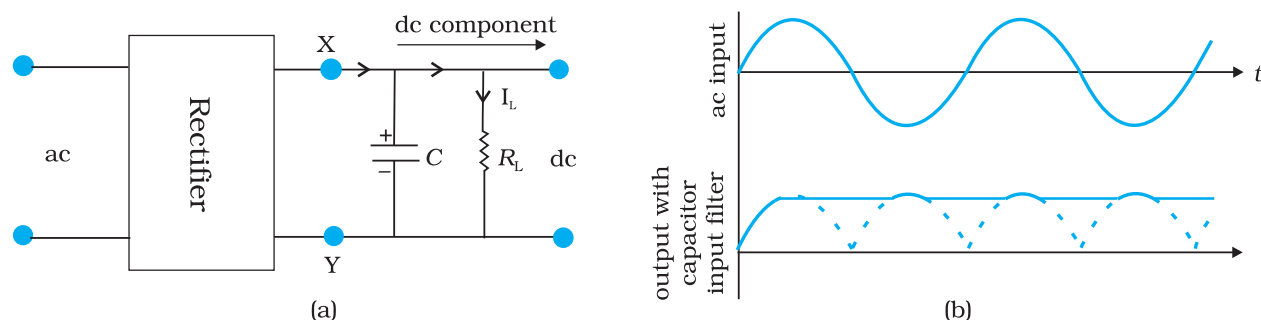


FIGURE 14.20 (a) A full-wave rectifier with capacitor filter, (b) Input and output voltage of rectifier in (a).

SUMMARY

1. Semiconductors are the basic materials used in the present solid state electronic devices like diode, transistor, ICs, etc.
2. Lattice structure and the atomic structure of constituent elements decide whether a particular material will be insulator, metal or semiconductor.
3. Metals have low resistivity (10^{-2} to $10^{-8} \Omega\text{m}$), insulators have very high resistivity ($>10^8 \Omega\text{m}^{-1}$), while semiconductors have intermediate values of resistivity.
4. Semiconductors are elemental (Si, Ge) as well as compound (GaAs, CdS, etc.).
5. Pure semiconductors are called 'intrinsic semiconductors'. The presence of charge carriers (electrons and holes) is an 'intrinsic' property of the material and these are obtained as a result of thermal excitation. The number of electrons (n_e) is equal to the number of holes (n_h) in intrinsic conductors. Holes are essentially electron vacancies with an effective positive charge.
6. The number of charge carriers can be changed by 'doping' of a suitable impurity in pure semiconductors. Such semiconductors are known as extrinsic semiconductors. These are of two types (n-type and p-type).
7. In n-type semiconductors, $n_e \gg n_h$ while in p-type semiconductors $n_h \gg n_e$.
8. n-type semiconducting Si or Ge is obtained by doping with pentavalent atoms (donors) like As, Sb, P, etc., while p-type Si or Ge can be obtained by doping with trivalent atom (acceptors) like B, Al, In etc.
9. $n_e n_h = n_i^2$ in all cases. Further, the material possesses an *overall charge neutrality*.
10. There are two distinct band of energies (called valence band and conduction band) in which the electrons in a material lie. Valence band energies are low as compared to conduction band energies. All energy levels in the valence band are filled while energy levels in the conduction band may be fully empty or partially filled. The electrons in the conduction band are free to move in a solid and are responsible for the conductivity. The extent of conductivity depends upon the energy gap (E_g) between the top of valence band (E_v) and the bottom of the conduction band E_c . The electrons from valence band can be excited by

heat, light or electrical energy to the conduction band and thus, produce a change in the current flowing in a semiconductor.

11. For insulators $E_g > 3$ eV, for semiconductors E_g is 0.2 eV to 3 eV, while for metals $E_g \approx 0$.
12. p-n junction is the 'key' to all semiconductor devices. When such a junction is made, a 'depletion layer' is formed consisting of immobile ion-cores devoid of their electrons or holes. This is responsible for a junction potential barrier.
13. By changing the external applied voltage, junction barriers can be changed. In forward bias (n-side is connected to negative terminal of the battery and p-side is connected to the positive), the barrier is decreased while the barrier increases in reverse bias. Hence, forward bias current is more (mA) while it is very small (μ A) in a p-n junction diode.
14. Diodes can be used for rectifying an ac voltage (restricting the ac voltage to one direction). With the help of a capacitor or a suitable filter, a dc voltage can be obtained.

POINTS TO PONDER

1. The energy bands (E_C or E_V) in the semiconductors are space delocalised which means that these are not located in any specific place inside the solid. The energies are the overall averages. When you see a picture in which E_C or E_V are drawn as straight lines, then they should be respectively taken simply as the *bottom* of conduction band energy levels and *top* of valence band energy levels.
2. In elemental semiconductors (Si or Ge), the n-type or p-type semiconductors are obtained by introducing 'dopants' as defects. In compound semiconductors, the change in relative stoichiometric ratio can also change the type of semiconductor. For example, in ideal GaAs the ratio of Ga:As is 1:1 but in Ga-rich or As-rich GaAs it could respectively be $\text{Ga}_{1.1}\text{As}_{0.9}$ or $\text{Ga}_{0.9}\text{As}_{1.1}$. In general, the presence of defects control the properties of semiconductors in many ways.

EXERCISES

- 14.1** In an n-type silicon, which of the following statement is true:
- (a) Electrons are majority carriers and trivalent atoms are the dopants.
 - (b) Electrons are minority carriers and pentavalent atoms are the dopants.
 - (c) Holes are minority carriers and pentavalent atoms are the dopants.
 - (d) Holes are majority carriers and trivalent atoms are the dopants.
- 14.2** Which of the statements given in Exercise 14.1 is true for p-type semiconductors.
- 14.3** Carbon, silicon and germanium have four valence electrons each. These are characterised by valence and conduction bands separated

by energy band gap respectively equal to $(E_g)_C$, $(E_g)_{Si}$ and $(E_g)_{Ge}$. Which of the following statements is true?

- (a) $(E_g)_{Si} < (E_g)_{Ge} < (E_g)_C$
- (b) $(E_g)_C < (E_g)_{Ge} > (E_g)_{Si}$
- (c) $(E_g)_C > (E_g)_{Si} > (E_g)_{Ge}$
- (d) $(E_g)_C = (E_g)_{Si} = (E_g)_{Ge}$

14.4 In an unbiased p-n junction, holes diffuse from the p-region to n-region because

- (a) free electrons in the n-region attract them.
- (b) they move across the junction by the potential difference.
- (c) hole concentration in p-region is more as compared to n-region.
- (d) All the above.

14.5 When a forward bias is applied to a p-n junction, it

- (a) raises the potential barrier.
- (b) reduces the majority carrier current to zero.
- (c) lowers the potential barrier.
- (d) None of the above.

14.6 In half-wave rectification, what is the output frequency if the input frequency is 50 Hz. What is the output frequency of a full-wave rectifier for the same input frequency.

Notes

© NCERT
not to be republished

APPENDICES

APPENDIX A 1 THE GREEK ALPHABET

Alpha	A	α	Iota	I	ι	Rho	P	ρ
Beta	B	β	Kappa	K	κ	Sigma	Σ	σ
Gamma	Γ	γ	Lambda	Λ	λ	Tau	T	τ
Delta	Δ	δ	Mu	M	μ	Upsilon	Y	υ
Epsilon	E	ε	Nu	N	ν	Phi	Φ	ϕ, φ
Zeta	Z	ζ	Xi	Ξ	ξ	Chi	X	χ
Eta	H	η	Omicron	O	o	Psi	Ψ	ψ
Theta	Θ	θ	Pi	Π	π	Omega	Ω	ω

APPENDIX A 2 COMMON SI PREFIXES AND SYMBOLS FOR MULTIPLES AND SUB-MULTIPLES

Multiple			Sub-Multiple		
Factor	Prefix	Symbol	Factor	Prefix	symbol
10^{18}	Exa	E	10^{-18}	atto	a
10^{15}	Peta	P	10^{-15}	femto	f
10^{12}	Tera	T	10^{-12}	pico	p
10^9	Giga	G	10^{-9}	nano	n
10^6	Mega	M	10^{-6}	micro	μ
10^3	kilo	k	10^{-3}	milli	m
10^2	Hecto	h	10^{-2}	centi	c
10^1	Deca	da	10^{-1}	deci	d

APPENDIX A 3 SOME IMPORTANT CONSTANTS

Name	Symbol	Value
Speed of light in vacuum	c	$2.9979 \times 10^8 \text{ m s}^{-1}$
Charge of electron	e	$1.602 \times 10^{-19} \text{ C}$
Gravitational constant	G	$6.673 \times 10^{-11} \text{ N m}^2 \text{ kg}^{-2}$
Planck constant	h	$6.626 \times 10^{-34} \text{ J s}$
Boltzmann constant	k	$1.381 \times 10^{-23} \text{ J K}^{-1}$
Avogadro number	N_A	$6.022 \times 10^{23} \text{ mol}^{-1}$
Universal gas constant	R	$8.314 \text{ J mol}^{-1} \text{ K}^{-1}$
Mass of electron	m_e	$9.110 \times 10^{-31} \text{ kg}$
Mass of neutron	m_n	$1.675 \times 10^{-27} \text{ kg}$
Mass of proton	m_p	$1.673 \times 10^{-27} \text{ kg}$
Electron-charge to mass ratio	e/m_e	$1.759 \times 10^{11} \text{ C/kg}$
Faraday constant	F	$9.648 \times 10^4 \text{ C/mol}$
Rydberg constant	R	$1.097 \times 10^7 \text{ m}^{-1}$
Bohr radius	a_0	$5.292 \times 10^{-11} \text{ m}$
Stefan-Boltzmann constant	σ	$5.670 \times 10^{-8} \text{ W m}^{-2} \text{ K}^{-4}$
Wien's Constant	b	$2.898 \times 10^{-3} \text{ m K}$
Permittivity of free space	ϵ_0 $1/4\pi \epsilon_0$	$8.854 \times 10^{-12} \text{ C}^2 \text{ N}^{-1} \text{ m}^{-2}$ $8.987 \times 10^9 \text{ N m}^2 \text{ C}^{-2}$
Permeability of free space	μ_0	$4\pi \times 10^{-7} \text{ T m A}^{-1}$ $\cong 1.257 \times 10^{-6} \text{ Wb A}^{-1} \text{ m}^{-1}$

OTHER USEFUL CONSTANTS

Name	Symbol	Value
Mechanical equivalent of heat	J	4.186 J cal^{-1}
Standard atmospheric pressure	1 atm	$1.013 \times 10^5 \text{ Pa}$
Absolute zero	0 K	$-273.15 \text{ }^\circ\text{C}$
Electron volt	1 eV	$1.602 \times 10^{-19} \text{ J}$
Unified Atomic mass unit	1 u	$1.661 \times 10^{-27} \text{ kg}$
Electron rest energy	mc^2	0.511 MeV
Energy equivalent of 1 u	1 u c^2	931.5 MeV
Volume of ideal gas (0 °C and 1atm)	V	22.4 L mol^{-1}
Acceleration due to gravity (sea level, at equator)	g	9.78049 m s^{-2}

ANSWERS

CHAPTER 9

- 9.1** $v = -54$ cm. The image is real, inverted and magnified. The size of the image is 5.0 cm. As $u \rightarrow f$, $v \rightarrow \infty$; for $u < f$, image is virtual.
- 9.2** $v = 6.7$ cm. Magnification = 5/9, i.e., the size of the image is 2.5 cm. As $u \rightarrow \infty$; $v \rightarrow f$ (but never beyond) while $m \rightarrow 0$.
- 9.3** 1.33; 1.7 cm
- 9.4** $n_{ga} = 1.51$; $n_{wa} = 1.32$; $n_{gw} = 1.144$; which gives $\sin r = 0.6181$ i.e., $r \simeq 38^\circ$.
- 9.5** $r = 0.8 \times \tan i_c$ and $\sin i_c = 1/1.33 \simeq 0.75$, where r is the radius (in m) of the largest circle from which light comes out and i_c is the critical angle for water-air interface, Area = 2.6 m^2
- 9.6** $n \simeq 1.53$ and D_m for prism in water $\simeq 10^\circ$
- 9.7** $R = 22$ cm
- 9.8** Here the object is virtual and the image is real. $u = +12$ cm (object on right; virtual)
- (a) $f = +20$ cm. Image is real and at 7.5 cm from the lens on its right side.
- (b) $f = -16$ cm. Image is real and at 48 cm from the lens on its right side.
- 9.9** $v = 8.4$ cm, image is erect and virtual. It is diminished to a size 1.8 cm. As $u \rightarrow \infty$, $v \rightarrow f$ (but never beyond f while $m \rightarrow 0$).
Note that when the object is placed at the focus of the concave lens (21 cm), the image is located at 10.5 cm (not at infinity as one might wrongly think).
- 9.10** A diverging lens of focal length 60 cm
- 9.11** (a) $v_e = -25$ cm and $f_e = 6.25$ cm give $u_e = -5$ cm; $v_o = (15 - 5) \text{ cm} = 10 \text{ cm}$,
 $f_o = u_o = -2.5$ cm; Magnifying power = 20
(b) $u_o = -2.59$ cm.
Magnifying power = 13.5.
- 9.12** Angular magnification of the eye-piece for image at 25 cm
 $= \frac{25}{2.5} + 1 = 11$; $|u_e| = \frac{25}{11} \text{ cm} = 2.27 \text{ cm}$; $v_o = 7.2 \text{ cm}$
Separation = 9.47 cm; Magnifying power = 88
- 9.13** 24; 150 cm
- 9.14** (a) Angular magnification = 1500
(b) Diameter of the image = 13.7 cm.

- 9.15** Apply mirror equation and the condition:
 (a) $f < 0$ (concave mirror); $u < 0$ (object on left)
 (b) $f > 0$; $u < 0$
 (c) $f > 0$ (convex mirror) and $u < 0$
 (d) $f < 0$ (concave mirror); $f < u < 0$
 to deduce the desired result.
- 9.16** The pin appears raised by 5.0 cm. It can be seen with an explicit ray diagram that the answer is independent of the location of the slab (for small angles of incidence).
- 9.17** (a) $\sin i'_c = 1.44/1.68$ which gives $i'_c = 59^\circ$. Total internal reflection takes place when $i > 59^\circ$ or when $r < r_{\max} = 31^\circ$. Now, $(\sin i_{\max}/\sin r_{\max}) = 1.68$, which gives $i_{\max} \simeq 60^\circ$. Thus, all incident rays of angles in the range $0 < i < 60^\circ$ will suffer total internal reflections in the pipe. (If the length of the pipe is finite, which it is in practice, there will be a lower limit on i determined by the ratio of the diameter to the length of the pipe.)
 (b) If there is no outer coating, $i'_c = \sin^{-1}(1/1.68) = 36.5^\circ$. Now, $i = 90^\circ$ will have $r = 36.5^\circ$ and $i' = 53.5^\circ$ which is greater than i'_c . Thus, all incident rays (in the range $53.5^\circ < i < 90^\circ$) will suffer total internal reflections.
- 9.18** For fixed distance s between object and screen, the lens equation does not give a real solution for u or v if f is greater than $s/4$. Therefore, $f_{\max} = 0.75$ m.
- 9.19** 21.4 cm
- 9.20** (a) (i) Let a parallel beam be the incident from the left on the convex lens first.
 $f_1 = 30$ cm and $u_1 = -\infty$, give $v_1 = +30$ cm. This image becomes a virtual object for the second lens.
 $f_2 = -20$ cm, $u_2 = + (30 - 8)$ cm = + 22 cm which gives, $v_2 = -220$ cm. The parallel incident beam appears to diverge from a point 216 cm from the centre of the two-lens system.
 (ii) Let the parallel beam be incident from the left on the concave lens first: $f_1 = -20$ cm, $u_1 = -\infty$, give $v_1 = -20$ cm. This image becomes a real object for the second lens: $f_2 = +30$ cm, $u_2 = - (20 + 8)$ cm = - 28 cm which gives, $v_2 = -420$ cm. The parallel incident beam appears to diverge from a point 416 cm on the left of the centre of the two-lens system.
- Clearly, the answer depends on which side of the lens system the parallel beam is incident. *Further we do not have a simple lens equation true for all u (and v) in terms of a definite constant of the system (the constant being determined by f_1 and f_2 , and the separation between the lenses).* The notion of effective focal length, therefore, does not seem to be meaningful for this system.
- (b) $u_1 = -40$ cm, $f_1 = 30$ cm, gives $v_1 = 120$ cm.
 Magnitude of magnification due to the first (convex) lens is 3.
 $u_2 = + (120 - 8)$ cm = +112 cm (object virtual);
 $f_2 = -20$ cm which gives $v_2 = -\frac{112 \times 20}{92}$ cm
 Magnitude of magnification due to the second (concave)

$$\text{lens} = 20/92.$$

$$\text{Net magnitude of magnification} = 0.652$$

$$\text{Size of the image} = 0.98 \text{ cm}$$

- 9.21** If the refracted ray in the prism is incident on the second face at the critical angle i_c , the angle of refraction r at the first face is $(60^\circ - i_c)$.

$$\text{Now, } i_c = \sin^{-1} (1/1.524) \simeq 41^\circ$$

$$\text{Therefore, } r = 19^\circ$$

$$\sin i = 0.4962; i \simeq 30^\circ$$

9.22 (a) $\frac{1}{v} + \frac{1}{9} = \frac{1}{10}$

$$\text{i.e., } v = -90 \text{ cm,}$$

$$\text{Magnitude of magnification} = 90/9 = 10.$$

$$\text{Each square in the virtual image has an area } 10 \times 10 \times 1 \text{ mm}^2 = 100 \text{ mm}^2 = 1 \text{ cm}^2$$

(b) Magnifying power = $25/9 = 2.8$

- (c) No, magnification of an image by a lens and angular magnification (or magnifying power) of an optical instrument are two separate things. The latter is the ratio of the angular size of the object (which is equal to the angular size of the image even if the image is magnified) to the angular size of the object if placed at the near point (25 cm). Thus, magnification magnitude is $|v/u|$ and magnifying power is $(25/|u|)$. Only when the image is located at the near point $|v| = 25 \text{ cm}$, are the two quantities equal.

- 9.23** (a) Maximum magnifying power is obtained when the image is at the near point (25 cm)

$$u = -7.14 \text{ cm.}$$

(b) Magnitude of magnification = $(25/|u|) = 3.5$.

(c) Magnifying power = 3.5

Yes, the magnifying power (when the image is produced at 25 cm) is equal to the magnitude of magnification.

9.24 Magnification = $\sqrt{6.25/1} = 2.5$

$$v = +2.5u$$

$$+\frac{1}{2.5u} - \frac{1}{u} = \frac{1}{10}$$

$$\text{i.e., } u = -6 \text{ cm}$$

$$|v| = 15 \text{ cm}$$

The virtual image is closer than the normal near point (25 cm) and cannot be seen by the eye distinctly.

- 9.25** (a) Even though the absolute image size is bigger than the object size, the angular size of the image is equal to the angular size of the object. The magnifier helps in the following way: without it object would be placed no closer than 25 cm; with it the object can be placed much closer. The closer object has larger angular size than the same object at 25 cm. It is in this sense that angular magnification is achieved.

- (b) Yes, it decreases a little because the angle subtended at the eye is then slightly less than the angle subtended at the lens. The

effect is negligible if the image is at a very large distance away.
[Note: When the eye is separated from the lens, the angles subtended at the eye by the first object and its image are not equal.]

- (c) First, grinding lens of very small focal length is not easy. More important, if you decrease focal length, aberrations (both spherical and chromatic) become more pronounced. So, in practice, you cannot get a magnifying power of more than 3 or so with a simple convex lens. However, using an aberration corrected lens system, one can increase this limit by a factor of 10 or so.
- (d) Angular magnification of eye-piece is $[(25/f_e) + 1]$ (f_e in cm) which increases if f_e is smaller. Further, magnification of the objective

$$\text{is given by } \frac{v_o}{|u_o|} = \frac{1}{(|u_o|/f_o) - 1}$$

which is large when $|u_o|$ is slightly greater than f_o . The microscope is used for viewing very close object. So $|u_o|$ is small, and so is f_o .

- (e) The image of the objective in the eye-piece is known as 'eye-ring'. All the rays from the object refracted by objective go through the eye-ring. Therefore, it is an ideal position for our eyes for viewing. If we place our eyes too close to the eye-piece, we shall not collect much of the light and also reduce our field of view. If we position our eyes on the eye-ring and the area of the pupil of our eye is greater or equal to the area of the eye-ring, our eyes will collect all the light refracted by the objective. The precise location of the eye-ring naturally depends on the separation between the objective and the eye-piece. When you view through a microscope by placing your eyes on one end, the ideal distance between the eyes and eye-piece is usually built-in the design of the instrument.

- 9.26** Assume microscope in normal use i.e., image at 25 cm. Angular magnification of the eye-piece

$$= \frac{25}{5} + 1 = 6$$

Magnification of the objective

$$= \frac{30}{6} = 5$$

$$\frac{1}{5u_o} - \frac{1}{u_o} = \frac{1}{1.25}$$

which gives $u_o = -1.5$ cm; $v_o = 7.5$ cm. $|u_e| = (25/6)$ cm = 4.17 cm. The separation between the objective and the eye-piece should be $(7.5 + 4.17)$ cm = 11.67 cm. Further the object should be placed 1.5 cm from the objective to obtain the desired magnification.

- 9.27** (a) $m = (f_o/f_e) = 28$

$$(b) \quad m = \frac{f_o}{f_e} \left[1 + \frac{f_o}{25} \right] = 33.6$$

- 9.28** (a) $f_o + f_e = 145 \text{ cm}$
 (b) Angle subtended by the tower = $(100/3000) = (1/30) \text{ rad}$.
 Angle subtended by the image produced by the objective

$$= \frac{h}{f_o} = \frac{h}{140}$$

 Equating the two, $h = 4.7 \text{ cm}$.
 (c) Magnification (magnitude) of the eye-piece = 6. Height of the final image (magnitude) = 28 cm.
- 9.29** The image formed by the larger (concave) mirror acts as virtual object for the smaller (convex) mirror. Parallel rays coming from the object at infinity will focus at a distance of 110 mm from the larger mirror. The distance of virtual object for the smaller mirror = $(110 - 20) = 90 \text{ mm}$. The focal length of smaller mirror is 70 mm. Using the mirror formula, image is formed at 315 mm from the smaller mirror.
- 9.30** The reflected rays get deflected by twice the angle of rotation of the mirror. Therefore, $d/1.5 = \tan 7^\circ$. Hence $d = 18.4 \text{ cm}$.
- 9.31** $n = 1.33$

CHAPTER 10

- 10.1** (a) Reflected light: (wavelength, frequency, speed same as incident light)
 $\lambda = 589 \text{ nm}$, $\nu = 5.09 \times 10^{14} \text{ Hz}$, $c = 3.00 \times 10^8 \text{ m s}^{-1}$
 (b) Refracted light: (frequency same as the incident frequency)
 $\nu = 5.09 \times 10^{14} \text{ Hz}$
 $\nu = (c/n) = 2.26 \times 10^8 \text{ m s}^{-1}$, $\lambda = (\nu/\nu) = 444 \text{ nm}$
- 10.2** (a) Spherical
 (b) Plane
 (c) Plane (a small area on the surface of a large sphere is nearly planar).
- 10.3** (a) $2.0 \times 10^8 \text{ m s}^{-1}$
 (b) No. The refractive index, and hence the speed of light in a medium, depends on wavelength. [When no particular wavelength or colour of light is specified, we may take the given refractive index to refer to yellow colour.] Now we know violet colour deviates more than red in a glass prism, i.e. $n_v > n_r$. Therefore, the violet component of white light travels slower than the red component.
- 10.4** $\lambda = \frac{1.2 \times 10^{-2} \times 0.28 \times 10^{-3}}{4 \times 1.4} \text{ m} = 600 \text{ nm}$
- 10.5** $K/4$
- 10.6** (a) 1.17 mm (b) 1.56 mm
- 10.7** 0.15°
- 10.8** $\tan^{-1}(1.5) \simeq 56.3^\circ$

10.9 $5000 \text{ \AA}$, $6 \times 10^{14} \text{ Hz}$; 45°

10.10 40m

CHAPTER 11

11.1 (a) $7.24 \times 10^{18} \text{ Hz}$ (b) 0.041 nm

11.2 (a) $0.34 \text{ eV} = 0.54 \times 10^{-19} \text{ J}$ (b) 0.34 V (c) 344 km/s

11.3 $1.5 \text{ eV} = 2.4 \times 10^{-19} \text{ J}$

11.4 (a) $3.14 \times 10^{-19} \text{ J}$, $1.05 \times 10^{-27} \text{ kg m/s}$ (b) $3 \times 10^{16} \text{ photons/s}$
(c) 0.63 m/s

11.5 $6.59 \times 10^{-34} \text{ Js}$

11.6 2.0 V

11.7 No, because $v < v_0$

11.8 $4.73 \times 10^{14} \text{ Hz}$

11.9 $2.16 \text{ eV} = 3.46 \times 10^{-19} \text{ J}$

11.10 (a) $1.7 \times 10^{-35} \text{ m}$ (b) $1.1 \times 10^{-32} \text{ m}$ (c) $3.0 \times 10^{-23} \text{ m}$

11.11 $\lambda = h/p = h/(hv/c) = c/v$

CHAPTER 12

12.1 (a) No different from
(b) Thomson's model; Rutherford's model
(c) Rutherford's model
(d) Thomson's model; Rutherford's model
(e) Both the models

12.2 The nucleus of a hydrogen atom is a proton. The mass of it is $1.67 \times 10^{-27} \text{ kg}$, whereas the mass of an incident α -particle is $6.64 \times 10^{-27} \text{ kg}$. Because the scattering particle is more massive than the target nuclei (proton), the α -particle won't bounce back in even in a head-on collision. It is similar to a football colliding with a tennis ball at rest. Thus, there would be no large-angle scattering.

12.3 $5.6 \times 10^{14} \text{ Hz}$

12.4 13.6 eV; -27.2 eV

12.5 $9.7 \times 10^{-8} \text{ m}$; $3.1 \times 10^{15} \text{ Hz}$.

12.6 (a) $2.18 \times 10^6 \text{ m/s}$; $1.09 \times 10^6 \text{ m/s}$; $7.27 \times 10^5 \text{ m/s}$
(b) $1.52 \times 10^{-16} \text{ s}$; $1.22 \times 10^{-15} \text{ s}$; $4.11 \times 10^{-15} \text{ s}$.

12.7 $2.12 \times 10^{-10} \text{ m}$; $4.77 \times 10^{-10} \text{ m}$

12.8 Lyman series: 103 nm and 122 nm; Balmer series: 656 nm.

12.9 2.6×10^{74}

CHAPTER 13

13.1 104.7 MeV

13.2 8.79 MeV, 7.84 MeV

13.3 $1.584 \times 10^{25} \text{ MeV}$ or $2.535 \times 10^{12} \text{ J}$

13.4 1.23

13.5 (i) $Q = -4.03$ MeV; endothermic

(ii) $Q = 4.62$ MeV; exothermic

13.6 $Q = m(^{56}_{26}\text{Fe}) - 2m(^{28}_{13}\text{Al}) = 26.90$ MeV; not possible.

13.7 4.536×10^{26} MeV

13.8 About 4.9×10^4 y

13.9 360 KeV

CHAPTER 14

14.1 (c)

14.2 (d)

14.3 (c)

14.4 (c)

14.5 (c)

14.6 50 Hz for half-wave, 100 Hz for full-wave

BIBLIOGRAPHY

TEXTBOOKS

For additional reading on the topics covered in this book, you may like to consult one or more of the following books. Some of these books however are more advanced and contain many more topics than this book.

- 1 **Ordinary Level Physics**, A.F. Abbott, Arnold-Heinemann (1984).
- 2 **Advanced Level Physics**, M. Nelkon and P. Parker, 6th Edition, Arnold-Heinemann (1987).
- 3 **Advanced Physics**, Tom Duncan, John Murray (2000).
- 4 **Fundamentals of Physics**, David Halliday, Robert Resnick and Jearl Walker, 7th Edition John Wily (2004).
- 5 **University Physics** (Sears and Zemansky's), H.D. Young and R.A. Freedman, 11th Edition, Addison—Wesley (2004).
- 6 **Problems in Elementary Physics**, B. Bukhovtso, V. Krivchenkov, G. Myakishev and V. Shalnov, MIR Publishers, (1971).
- 7 **Lectures on Physics** (3 volumes), R.P. Feynman, Addison – Wesley (1965).
- 8 **Berkeley Physics Course** (5 volumes) McGraw Hill (1965).
 - a. Vol. 1 – Mechanics: (Kittel, Knight and Ruderman)
 - b. Vol. 2 – Electricity and Magnetism (E.M. Purcell)
 - c. Vol. 3 – Waves and Oscillations (Frank S. Crawford)
 - d. Vol. 4 – Quantum Physics (Wichmann)
 - e. Vol. 5 – Statistical Physics (F. Reif)
- 9 **Fundamental University Physics**, M. Alonso and E. J. Finn, Addison – Wesley (1967).
- 10 **College Physics**, R.L. Weber, K.V. Manning, M.W. White and G.A. Weygand, Tata McGraw Hill (1977).
- 11 **Physics: Foundations and Frontiers**, G. Gamow and J.M. Cleveland, Tata McGraw Hill (1978).
- 12 **Physics for the Inquiring Mind**, E.M. Rogers, Princeton University Press (1960).
- 13 **PSSC Physics Course**, DC Heath and Co. (1965) Indian Edition, NCERT (1967).
- 14 **Physics Advanced Level**, Jim Breithampt, Stanley Thornes Publishers (2000).
- 15 **Physics**, Patrick Fullick, Heinemann (2000).
- 16 **Conceptual Physics**, Paul G. Hewitt, Addison—Wesley (1998).
- 17 **College Physics**, Raymond A. Serway and Jerry S. Faughn, Harcourt Brace and Co. (1999).
- 18 **University Physics**, Harris Benson, John Wiley (1996).
- 19 **University Physics**, William P. Crummet and Arthur B. Western, Wm.C. Brown (1994).
- 20 **General Physics**, Morton M. Sternheim and Joseph W. Kane, John Wiley (1988).
- 21 **Physics**, Hans C. Ohanian, W.W. Norton (1989).

- 22 Advanced Physics**, Keith Gibbs, Cambridge University Press (1996).
- 23 Understanding Basic Mechanics**, F. Reif, John Wiley (1995).
- 24 College Physics**, Jerry D. Wilson and Anthony J. Buffa, Prentice Hall (1997).
- 25 Senior Physics, Part – I**, I.K. Kikoin and A.K. Kikoin, MIR Publishers (1987).
- 26 Senior Physics, Part – II**, B. Bekhovtsev, MIR Publishers (1988).
- 27 Understanding Physics**, K. Cummings, Patrick J. Cooney, Priscilla W. Laws and Edward F. Redish, John Wiley (2005).
- 28 Essentials of Physics**, John D. Cutnell and Kenneth W. Johnson, John Wiley (2005).

GENERAL BOOKS

For instructive and entertaining general reading on science, you may like to read some of the following books. Remember however, that many of these books are written at a level far beyond the level of the present book.

- 1 Mr. Tompkins** in paperback, G. Gamow, Cambridge University Press (1967).
- 2 The Universe and Dr. Einstein**, C. Barnett, Time Inc. New York (1962).
- 3 Thirty years that Shook Physics**, G. Gamow, Double Day, New York (1966).
- 4 Surely You're Joking, Mr. Feynman**, R.P. Feynman, Bantam books (1986).
- 5 One, Two, Three... Infinity**, G. Gamow, Viking Inc. (1961).
- 6 The Meaning of Relativity**, A. Einstein, (Indian Edition) Oxford and IBH Pub. Co. (1965).
- 7 Atomic Theory and the Description of Nature**, Niels Bohr, Cambridge (1934).
- 8 The Physical Principles of Quantum Theory**, W. Heisenberg, University of Chicago Press (1930).
- 9 The Physics—Astronomy Frontier**, F. Hoyle and J.V. Narlikar, W.H. Freeman (1980).
- 10 The Flying Circus of Physics with Answer**, J. Walker, John Wiley and Sons (1977).
- 11 Physics for Everyone** (series), L.D. Landau and A.I. Kitaigorodski, MIR Publisher (1978).
 Book 1: Physical Bodies
 Book 2: Molecules
 Book 3: Electrons
 Book 4: Photons and Nuclei.
- 12 Physics can be Fun**, Y. Perelman, MIR Publishers (1986).
- 13 Power of Ten**, Philip Morrison and Eames, W.H. Freeman (1985).
- 14 Physics in your Kitchen Lab.**, I.K. Kikoin, MIR Publishers (1985).
- 15 How Things Work: The Physics of Everyday Life**, Louis A. Bloomfield, John Wiley (2005).
- 16 Physics Matters: An Introduction to Conceptual Physics**, James Trefil and Robert M. Hazen, John Wiley (2004).