

Sulabha K. Kulkarni

Nanotechnology: Principles and Practices

3rd Edition



Springer

Nanotechnology: Principles and Practices

Sulabha K. Kulkarni

Nanotechnology: Principles and Practices

Third Edition

 Springer



Sulabha K. Kulkarni
Physics Department
Indian Institute of Science
Education and Research
Pune, India

Co-published by Springer International Publishing, Cham, Switzerland, with Capital Publishing Company, New Delhi, India.

Sold and distributed in North, Central and South America by Springer, 233 Spring Street, New York 10013, USA.

In all other countries, except SAARC countries—Afghanistan, Bangladesh, Bhutan, India, Maldives, Nepal, Pakistan and Sri Lanka—sold and distributed by Springer, Haberstrasse 7, D-69126 Heidelberg, Germany.

In SAARC countries—Afghanistan, Bangladesh, Bhutan, India, Maldives, Nepal, Pakistan and Sri Lanka—printed book sold and distributed by Capital Publishing Company, 7/28, Mahaveer Street, Ansari Road, Daryaganj, New Delhi, 110 002, India.

ISBN 978-3-319-09170-9 ISBN 978-3-319-09171-6 (eBook)
DOI 10.1007/978-3-319-09171-6
Springer Cham Heidelberg New York Dordrecht London

Library of Congress Control Number: 2014953513

1st edition: © Capital Publishing Company 2007
2nd edition: © Capital Publishing Company 2011
© Capital Publishing Company 2015

This work is subject to copyright. All rights are reserved by the Publishers, whether the whole or part of the material is concerned, specifically the rights of translation, reprinting, reuse of illustrations, recitation, broadcasting, reproduction on microfilms or in any other physical way, and transmission or information storage and retrieval, electronic adaptation, computer software, or by similar or dissimilar methodology now known or hereafter developed. Exempted from this legal reservation are brief excerpts in connection with reviews or scholarly analysis or material supplied specifically for the purpose of being entered and executed on a computer system, for exclusive use by the purchaser of the work. Duplication of this publication or parts thereof is permitted only under the provisions of the Copyright Law of the Publishers' locations, in its current version, and permission for use must always be obtained from Springer. Permissions for use may be obtained through RightsLink at the Copyright Clearance Center. Violations are liable to prosecution under the respective Copyright Law.

The use of general descriptive names, registered names, trademarks, service marks, etc. in this publication does not imply, even in the absence of a specific statement, that such names are exempt from the relevant protective laws and regulations and therefore free for general use.

While the advice and information in this book are believed to be true and accurate at the date of publication, neither the authors nor the editors nor the publishers can accept any legal responsibility for any errors or omissions that may be made. The publishers make no warranty, express or implied, with respect to the material contained herein.

Printed on acid-free paper

Springer is part of Springer Science+Business Media (www.springer.com)

To my family

Preface to the Third Edition

Over the past 15 years, the scientific world has witnessed a dramatic interest in nanotechnology with research contributions from physicists, chemists, biologists as well as scientists working in the medical, agricultural, textile and environmental fields. The varied inputs from different disciplines and their complementary nature in the synthesis, characterization, understanding and applications from household items to space technologies have enriched the nanotechnology basics. This can be easily understood as about 80 journals have emerged on nanotechnology and hundreds of thousands of papers have been published in the last 12 years or so. Lecturers aware of this fact naturally want to introduce the new knowledge to the young generation and prepare their students for the great challenges of sustainable management of energy, water and environment faced by mankind.

The availability of a large number of books, review articles and journals on nanotechnology is essential in order to get an overall picture of such a vast field, which is a very difficult task. Particularly in the case of beginners in the field as well as students at the graduate level, the complex terminologies can prevent them from exploring the fascinating world of nanotechnology.

This book makes an attempt to introduce students, teachers and research scientists trying to enter the vast field of nanotechnology to nanosynthesis and various analysis methods and applications in a simple way, without overlooking the necessary principles underlying nanoscience. Continual attempts have been made right from the first edition (2006) to cover the contemporary ideas in the field. Due to the fast growth in the field by the launch of the second edition (2011), some more synthesis and analysis techniques which were becoming popular and newly discovered phenomena and effects were added to the basic contents of the first edition.

In the third edition, the book has been further reorganized keeping Chaps. 1, 2, 3, 4, 5, and 6 as in the second edition with some additional material on fuel cells and solar cells as well as including separate chapters on nanoelectronics and on nanotechnology and environment. The latter chapter, although briefly, gives an idea

of how nanotechnology can help in the detection of air and water pollution and in its remedial approaches. It also describes the negative aspects of nanotechnology that can be harmful to the environment or have adverse effects on human health. The chapter is rather brief because much research is awaited in this area. However, due to its importance, it is included as a separate chapter. In the chapter on nanoelectronics, single electron transistor and spin field-effect transistor are explained in more detail. Besides, two new appendices on the Kronig–Penney model and a data table on bulk semiconductors are included, which will provide some ready data for solving small problems like determination of band gaps in nanomaterials.

It is expected that this edition would give a good overview and provide a strong foundation of nanotechnology to its students. Most part of the book was prepared through my teaching notes, which were updated from time to time, while teaching nanotechnology at the University of Pune and the Indian Institute of Science Education and Research (IISER), Pune, for more than a decade. The book also contains some part of my own research (as well as that of my PhD and postgraduate students) in Nanotechnology and Condensed Matter Physics. Many of my students have contributed by way of preparing diagrams and sketches and providing photographs as well as carefully going through some parts of the book. For the third edition, I would like to thank Dr. Pavan G.V. Kumar, Dr. Smita Chaturvedi, Ms. Rashmi Runjhun, Mr. Prashant Bhaskar, Mr. Amey Apte, Mr. Arshad Nair and Ms. Supreet Singh for reading the drafts, making suggestions and drawing figures. I also thank the Director of IISER Pune Prof. K.N. Ganesh and Chair of the Physics Programme at IISER Pune Prof. Sunil Mukhi for their generous support while preparing this edition.

Last but not the least, I would like to thank Capital Publishing Company and Springer for their help, suggestions and careful proofreading of this edition.

Pune, India
May 2014

Sulabha K. Kulkarni

Preface to the Second Edition

Nanotechnology has now become a buzz word all over the world. Even the common man and school children hear the word “Nano”—may be through the cars rolling on the streets, housing projects and products like washing machines or mobiles. In some products there may not be any use of nanomaterials but the word indeed has become popular to indicate ‘small’. This indeed shows to some extent awareness, importance and acceptance of nanotechnology and nanoscience.

Nanotechnology is being introduced in many curricula at post-graduate and under-graduate levels although research in this field has not reached its saturation. It has therefore become indeed difficult to find any book which will be able to give all the new developments in the field with enough coverage keeping basic components intact.

In this edition attempt is made not to remove almost any of the original material of the first edition but add more material like ‘*Self Assembly*’ as a separate chapter (No. 6) and an Appendix (IV) on ‘*Vacuum Techniques*’. Although the word and some examples of ‘self assembly’ were part of the previous edition, progress in ‘Self Assembly’ and its importance is so overwhelming that it deserved a separate chapter in this book. However, I am quite aware that more material in this chapter also would have been in order. The appendix on ‘Vacuum Techniques’ would serve the readers to at least get an idea of this essential technology needed to synthesize nanomaterials like in physical and chemical vapour deposition systems. It also forms an integral part of numerous analysis instruments like ‘Electron Microscopes’ or spectroscopy techniques like ‘Photoelectron Spectrometer’ discussed in Chap. 7. In Chap. 7 the readers will also find some techniques like ‘*Dynamic Light Scattering*’ and ‘*Raman Spectroscopy*’ introduced in this edition. There are more illustrations introduced in this chapter. In Chap. 4 on Chemical Synthesis, additional techniques like ‘*Hydrothermal Synthesis*’, ‘*Sonochemical Synthesis*’, ‘*Microwave Synthesis*’ and ‘*Microreactor*’ or ‘*Lab-on-Chip*’ are introduced in the new edition. Similarly, in Chap. 9 on Nanolithography, the concepts of ‘*Nano Sphere Lithography*’ and ‘*Nanoimprint Lithography*’ are introduced.

Chapter 8 has a separate section on ‘*Clusters*’ and has been named as ‘Properties of Clusters and Nanomaterials’ as against the ‘Properties of Nanomaterials’ in the previous edition. In Chap. 10 on Some Special Nanomaterials one would find ‘*Graphene*’ and ‘*Metamaterials*’ or ‘*Negative Refractive Index*’ materials being introduced. Chapter 11 on applications is a chapter for which sky is the limit! Therefore attempt is made to introduce only the ‘*Solar Cells*’ and ‘*Lotus* (as well as *Gecko*) *Effect*’ which perhaps needed more attention. There is a considerable concern now on developing renewable, inexpensive, clean or pollution-free energy sources and solar cells based on nanomaterials open up a challenge for the scientists and technologists. Similarly with understanding of lotus effect, many new self cleaning products are around and deserve an introduction. One more experiment on making ordered pores in aluminum foil (*AAO templates*) is given in the new edition. This itself can be an interesting experiment as well as some may find it useful to deposit later materials in pores to make ordered nanorods for some applications like solar cells.

Thus it is hoped that the readers would find lot of new concepts, materials, techniques and application areas introduced in the second edition of the book.

Pune, India
March, 2011

Sulabha K. Kulkarni

Preface to the First Edition

Periodic table with 118 elements is quite limited in number. But by combination of different elements in certain proportions, nature has produced a number of gases, liquids, minerals and above all, the living world. Mankind too, ingeniously working, learnt to make a large number of materials and even cast or shape them for desired functions or operations. Starting with stone implements man learnt to separate metals and make alloys. He made wonderful organic and inorganic materials.

Now we use wood, metals, alloys, polymers etc. for our comforts. Some materials are directly taken from the nature and some are man-made. All the materials used today have a variety of functions, which is the fruit of skill and intellect of many generations of mankind. All this has led to wonderful electronic systems, communication tools, transport vehicles, textile, utensils, architectural materials, medicines etc.

Mankind has been constantly trying to overcome its physical limitations. Animals like horse can gallop at high speed and only the birds can fly. But by inventing the appropriate materials and understanding nature's aviation system, man has been able to develop vehicles that run faster than a horse and fly in the sky reaching a height impossible for a bird. He has developed the communication and navigation systems, which take him or his instruments to distant planets. He may not physically go everywhere but he has tried ingeniously to get some knowledge of different planets, stars or even the galaxies.

But in order to continue with such adventures, man needs more and more materials with controlled properties. He needs materials not known before. He needs materials, which are highly efficient, often small in size and novel. In the attempt of making lightweight and smaller and smaller electronic devices, scientists have reduced the size of materials to such an extent that it has reached nanometric dimensions. At such a scale new phenomena are observed for practically all the

materials known so far, leading to novel devices and potential applications in different fields from consumer goods to health related equipment. Perhaps we are now living in the age Nobel laureate Richard Feynman dreamt of in 1959. He, in his now very famous speech delivered to the American Physical Society, said that why not mimic nature and produce smaller and smaller functional materials, which will be highly efficient. At that time the computers were very big occupying large buildings yet only good enough to make calculations now done on a palm size calculator. Scientists used to be still proud of those computers. The radio sets used to be occupying large space and be power hungry. But now see the change of scenario after forty years. We have smaller and efficient Personal Computers (PC), Laptops, Compact Disc (CD) players, pocket transistors, mobiles with digital camera and what not. Although Feynman did not utter the words ‘Nanoscience’ or ‘Nanotechnology’ we see the advantages of making things small.

However, it may be remembered that nanomaterials are not really new. Michael Faraday synthesized stable gold colloidal particles of nano size in 1857 A.D. His gold samples are still in the British Museum in U.K. showing beautiful magenta-red colour (not golden!) solution. Decorative glass windows with beautiful designs in old churches and palaces indeed use nanoparticles of iron, cobalt, nickel, gold, silver etc. Photographic plates also have nanoparticles and whole branch of catalysis in chemistry has a variety of catalyst particles in nanometer range. However in all such examples there were systematic attempts to understand the size effects on properties—either lacking or insufficient. The lack of powerful microscopes, which would correlate the sizes to properties was the main barrier in early work. One can consider some of the milestones (or history) in Nanotechnology as given in Box 1.

Box 1: Milestones in Nanotechnology

- 1857—Michael Faraday synthesized gold colloids of nanosize
- 1915—W. Ostwald, a famous chemist, wrote a book ‘World of Neglected Dimensions’ in German
- 1931—E. Ruska and M. Knoll developed the first electron microscope
- 1951—E. Müller developed the Field Ion Microscope which enabled the imaging of atoms from the tip of metallic samples
- 1959—R. Feynman delivered his now very famous talk ‘There is Plenty of Room at the Bottom’ pointing out to the scientists that reduced dimensionality of materials would create fascinating materials
- 1968—A.Y. Cho and J. Arthur developed Molecular Beam Epitaxy technique for layer by layer growth of materials
- 1970—L. Esaki demonstrated the quantum size effect (QSE) in semiconductors

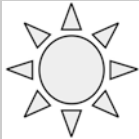
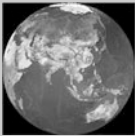
(continued)

Box 1 (continued)

- 1980—A.I. Akimov showed QSE in CdS and CdSe particles dispersed in glass, triggering the research on nanoparticles
- 1981—G. Binnig and H. Rohrer developed the scanning tunnelling microscope (STM) by which atomic resolution could be obtained. This was also followed by a family of scanning probe microscopes of various types
- 1985—R.F. Curl, H.W. Kroto and R.F. Smalley synthesized sixty atom carbon molecule, later named as ‘Fullerene’
- 1989—D.M. Eigler wrote letters ‘IBM’ using xenon atoms
- 1991—S. Iijima discovered ‘carbon nanotubes’
- 1999—C.A. Mirkin developed the ‘Dip Pen Lithography’
- 2000—D.M. Eigler devised ‘Quantum Mirage’ using Fe atoms on the copper substrate.

Beginning of the twenty first century has witnessed a tremendous upsurge of scientific activity in the field of ‘Nanoscience’ and ‘Nanotechnology’, whose seeds were sowed in the last century. Scientists, technocrats and even governments of many countries all over the world are convinced that nanotechnology based on nanoscience is the technology of twenty first century. The terms ‘Nanotechnology’ and ‘Nanoscience’ are often used synonymously. The literal meaning of ‘nano’ is ‘dwarf’ or an abnormally short person. However in scientific language it is a billionth (10^{-9}) part of some unit scale, e.g. nanometre or nanosecond mean 10^{-9} m [see Box 2] or 10^{-9} s respectively. Nanometre is so small that if you imagine only ten atoms of hydrogen placed in a line touching each other it will measure one nanometre. To get a qualitative idea, see Box 3. Nanotechnology is thus the technology of materials dealing with very small dimension materials usually in the range of 1–100 nm. When at least one of the dimensions of any type of material is reduced below 100 nm, its mechanical, thermal, optical, magnetic and other properties change at some size characteristic of that material. Thus within the same material one can get a range of properties. For example [See Box 4] consider a semiconductor like CdS, which is normally reddish in colour. If one brings down the particle size (i.e. diameter) of CdS to say 10 nm, its powder still has red colour. But below about 6 nm size, a dramatic change occurs in the optical properties of CdS. As illustrated in the photograph the colour of 4 nm size particles is orange, 3 nm size particles is yellow and that of 2 nm size particles is white. Not only the visual appearance but other properties also change dramatically. Melting point for pure bulk solids is very sharply defined. If we measure the melting point for CdS nanoparticles (i.e. particles with diameter in few nanometres range) then we will find that the melting point reduces with the particle size. Therefore by changing the particle size of a material one can achieve a range of properties.

Box 2					
Factor	Symbol	Prefix	Factor	Symbol	Prefix
10^{-18}	a	atto	10^1	da	deka
10^{-15}	f	femto	10^2	h	hecto
10^{-12}	p	pico	10^3	k	kilo
10^{-9}	n	nano	10^6	M	mega
10^{-6}	μ	micro	10^9	G	giga
10^{-3}	m	milli	10^{12}	T	tera
10^{-2}	c	centi	10^{15}	p	peta
10^{-1}	d	deci	10^{18}	E	exa

Box 3: Comparison of Different Objects		
	Diameter of the Sun	1,393,000 km
	Diameter of the Earth	12,715 km
	Height of Himalaya Mountain	8,848 m
	Height of a Man	1.65 m
	Virus	20–250 nm
	Cadmium Sulphide Nanoparticles	1–10 nm

Box 4: CdS Nanoparticles (Colour Change with Particle Size)

Technologists thus find tremendous potential of nanomaterials from consumer goods to space applications. In a span of few years there have emerged a number of manufacturers selling products like cosmetics, clothes, TVs, computers, medicines, toys, sports goods, automobile etc. which use some nanomaterial or the other.

Nanotechnology is an interdisciplinary science. It needs Physics, Chemistry, Engineering, Biology etc. so that its full potential can be exploited for the advantage of mankind. What has been achieved in nanotech so far is only the tip of the iceberg. However to fully explore the potential of nanotechnology it is essential to know what are nanomaterials, how and why do they differ from other materials, how to synthesize/analyze the nanomaterials organize them and understand some already proven application areas.

There are several books available on nanotechnology written by the leading scientists working in nanotechnology. Most of the books, however, are collection of research articles. This means the reader is expected to have considerable background in basic sciences. Often there is a jargon of technical terms, which is difficult to understand.

This book aims at developing sufficient background for students of graduate and post-graduate levels. Some boxes are given so that without leaving the main flow of the text additional material could be given. The book is meant to be self-sufficient.

Pune, India
September, 2006

Sulabha K. Kulkarni

Contents

1	Introduction to Quantum Mechanics	1
1.1	Introduction	1
1.2	Matter Waves	10
1.3	Heisenberg's Uncertainty Principle	12
1.4	Schrödinger Equation	15
1.5	Electron Confinement	19
1.5.1	Particle in a Box	20
1.5.2	Density of States	22
1.5.3	Particle in a Coulomb Potential	25
1.6	Tunnelling of a Particle Through Potential Barrier	27
	Further Reading	29
2	Structure and Bonding	31
2.1	Introduction	31
2.2	Arrangement of Atoms	33
2.3	Two Dimensional Crystal Structures	36
2.4	Three Dimensional Crystal Structures	36
2.5	Some Examples of Three Dimensional Crystals	38
2.5.1	Body Centred Cube (bcc)	38
2.6	Planes in the Crystals	38
2.7	Crystallographic Directions	38
2.8	Reciprocal Lattice	38
2.9	Quasi Crystals	41
2.10	Liquid Crystals	43
2.11	Bonding in Solids	44
2.11.1	Covalent Bond	45
2.11.2	Ionic Bond	46
2.11.3	Metallic Bond	47
2.11.4	Mixed Bonds	47
2.11.5	Secondary Bonds	48

2.12	Electronic Structure of Solids	49
2.12.1	Free Electron Motion	50
2.12.2	Bloch's Theorem.....	50
2.12.3	Origin of Band Structure	51
	Further Reading	53
3	Synthesis of Nanomaterials—I (Physical Methods)	55
3.1	Introduction	55
3.2	Mechanical Methods.....	55
3.2.1	High Energy Ball Milling	55
3.2.2	Melt Mixing.....	57
3.3	Methods Based on Evaporation	59
3.3.1	Physical Vapour Deposition with Consolidation	61
3.3.2	Ionized Cluster Beam Deposition	63
3.3.3	Laser Vapourization (Ablation)	64
3.3.4	Laser Pyrolysis.....	65
3.4	Sputter Deposition	65
3.4.1	DC Sputtering.....	67
3.4.2	RF Sputtering	69
3.4.3	Magnetron Sputtering.....	69
3.4.4	ECR Plasma Deposition	70
3.5	Chemical Vapour Deposition (CVD)	71
3.6	Electric Arc Deposition.....	73
3.7	Ion Beam Techniques (Ion Implantation).....	75
3.8	Molecular Beam Epitaxy (MBE).....	75
	Further Reading	76
4	Synthesis of Nanomaterials—II (Chemical Methods)	77
4.1	Introduction	77
4.2	Colloids and Colloids in Solutions	78
4.2.1	Interactions of Colloids and Medium	78
4.2.2	Colloids in Vacuum	83
4.2.3	Colloids in a Medium.....	84
4.2.4	Effect of Charges on Colloids	84
4.2.5	Stearic Repulsion	86
4.2.6	Synthesis of Colloids	87
4.3	Nucleation and Growth of Nanoparticles	87
4.4	Synthesis of Metal Nanoparticles by Colloidal Route	91
4.5	Synthesis of Semiconductor Nanoparticles by Colloidal Route ..	92
4.6	Langmuir-Blodgett (LB) Method	95
4.7	Microemulsions	98
4.8	Sol-Gel Method	103
4.9	Hydrothermal Synthesis	105
4.10	Sonochemical Synthesis	106

4.11	Microwave Synthesis	107
4.12	Synthesis Using Micro-reactor or Lab-On-Chip	107
	Further Reading	109
5	Synthesis of Nanomaterials—III (Biological Methods)	111
5.1	Introduction	111
5.2	Synthesis Using Microorganisms	116
5.3	Synthesis Using Plant Extracts	120
5.4	Use of Proteins, Templates Like DNA, S-Layers etc.	121
5.5	Synthesis of Nanoparticles Using DNA	123
	Further Reading	123
6	Self Assembly	125
6.1	Introduction	125
6.2	Mechanism of Self Assembly	127
6.3	Some Examples of Self Assembly	129
6.3.1	Self Assembly of Nanoparticles Using Organic Molecules	129
6.3.2	Self Assembly in Biological Systems	130
6.3.3	Self Assembly in Inorganic Materials	132
	Further Reading	133
7	Analysis Techniques	135
7.1	Introduction	135
7.2	Microscopes	135
7.2.1	Optical Microscopes	135
7.2.2	Confocal Microscope	140
7.3	Electron Microscopes	141
7.3.1	Scanning Electron Microscope	143
7.3.2	Transmission Electron Microscope (TEM)	146
7.4	Scanning Probe Microscopes (SPM)	148
7.4.1	Scanning Tunnelling Microscope	149
7.4.2	Atomic Force Microscope	152
7.4.3	Scanning Near-Field Optical Microscope (SNOM)	155
7.5	Diffraction Techniques	159
7.5.1	X-Ray Diffraction (XRD)	160
7.5.2	Atomic Scattering Factor	161
7.5.3	Bragg's Law of Diffraction	162
7.5.4	Diffraction from Different Types of Samples	165
7.5.5	Crystal Structure Factor	166
7.5.6	Diffraction from Nanoparticles	167
7.5.7	X-ray Diffractometer	170
7.5.8	Dynamic Light Scattering	171

7.6	Spectroscopies	173
7.6.1	Optical (Ultraviolet-Visible-Near Infra Red) Absorption Spectrometer	173
7.6.2	UV-Vis-NIR Spectrometer	175
7.6.3	Infra Red Spectrometers	176
7.6.4	Dispersive Infra Red Spectrometer	178
7.6.5	Fourier Transform Infra Red Spectrometer	179
7.6.6	Raman Spectroscopy	181
7.6.7	Luminescence	184
7.6.8	X-Ray and Ultra Violet Photoelectron Spectroscopies (XPS or ESCA and UPS)	186
7.6.9	Auger Electron Spectroscopy	190
7.7	Magnetic Measurements	192
7.7.1	Vibrating Sample Magnetometer (VSM)	192
7.8	Mechanical Measurements	194
7.8.1	Some Common Terminologies Related to Mechanical Properties	194
	Further Reading	197
8	Types of Nanomaterials and Their Properties	199
8.1	Introduction	199
8.2	Clusters	200
8.2.1	Types of Clusters	200
8.3	Semiconductor Nanoparticles	203
8.3.1	Excitons	204
8.3.2	Effective Mass Approximation	205
8.3.3	Optical Properties of Semiconductor Nanoparticles ...	208
8.4	Plasmonic Materials	214
8.4.1	Localized Surface Plasmon Resonance	215
8.4.2	Surface Plasmon Polariton	222
8.5	Nanomagnetism	225
8.5.1	Types of Magnetic Materials	227
8.6	Mechanical Properties of Nanomaterials	235
8.7	Structural Properties	237
8.8	Melting of Nanoparticles	238
	Further Reading	239
9	Nanolithography	241
9.1	Introduction	241
9.2	Lithography Using Photons (UV-VIS, Lasers and X-Rays).....	245
9.2.1	Lithography Using UV Light and Laser Beams	246
9.2.2	Use of X-rays in Lithography	246
9.3	Lithography Using Particle Beams	246
9.3.1	Electron Beam Lithography	247
9.3.2	Ion Beam Lithography	248

9.3.3	Neutral Beam Lithography	248
9.3.4	Nano Sphere Lithography	248
9.4	Scanning Probe Lithography	249
9.4.1	Mechanical Methods	249
9.4.2	Dip Pen Lithography	250
9.4.3	Optical Scanning Probe Lithography	251
9.4.4	Thermo-Mechanical Lithography	251
9.4.5	Electrical Scanning Probe Lithography	252
9.5	Soft Lithography	252
9.5.1	Microcontact Printing (μ CP)	253
9.5.2	Replica Molding (REM)	254
9.5.3	Microtransfer Molding (μ TM)	255
9.5.4	Micromolding in Capillaries (MIMIC)	255
9.5.5	Solvent-Assisted Micromolding (SAMIM)	255
	Further Reading	257
10	Nanoelectronics	259
10.1	Introduction	259
10.2	Coulomb Blockade	260
10.3	Single Electron Transistor (SET)	263
10.4	Spintronics	267
10.4.1	Giant Magneto Resistance	268
10.4.2	Spin Valve	270
10.4.3	Magnetic Tunnel Junction (MTJ)	271
10.4.4	Spin Field Effect Transistor (SFET)	271
10.5	Nanophotonics	272
	Further Reading	272
11	Some Special Nanomaterials	273
11.1	Introduction	273
11.2	Carbon Nanomaterials	273
11.2.1	Fullerenes	273
11.2.2	Carbon Nanotubes (CNTs)	274
11.2.3	Types of Carbon Nanotubes	279
11.2.4	Synthesis of Carbon Nanotubes	281
11.2.5	Growth Mechanism	283
11.2.6	Graphene	285
11.3	Porous Material	286
11.3.1	Porous Silicon	286
11.3.2	How to Make Silicon Porous?	288
11.3.3	Mechanism of Pores Formation	290
11.3.4	Properties of Porous Silicon Morphology	293
11.4	Aerogels	296
11.4.1	Types of Aerogels	298
11.4.2	Properties of Aerogels	302
11.4.3	Applications of Aerogels	302

11.5	Zeolites	303
11.5.1	Synthesis of Zeolites	304
11.5.2	Properties of Zeolites	305
11.6	Porosity Through Templates	306
11.6.1	Micelles as Templates	306
11.6.2	Metal Organic Frameworks (MOF)	307
11.7	Core-Shell Particles	308
11.7.1	Synthesis of Silica Cores	309
11.7.2	Core-Shell Assemblies	310
11.7.3	Properties of Core-Shell Particles	311
11.8	Metamaterials	311
11.9	Bioinspired Materials	313
11.9.1	Lotus Effect (Self Cleaning)	313
11.9.2	Gecko Effect (Adhesive Materials)	315
	Further Reading	315
12	Applications	317
12.1	Introduction	317
12.2	Energy	317
12.2.1	Dye Sensitized Photovoltaic Solar Cell (Grätzel Cell)	321
12.2.2	Organic (Polymer/Small Organic Molecules) Photovoltaic Cells	326
12.2.3	Fuel Cell	327
12.2.4	Hydrogen Generation and Storage	332
12.2.5	Hydrogen Storage (and Release)	334
12.2.6	Hybrid Energy Cells	335
12.3	Automobiles	336
12.4	Sports and Toys	338
12.5	Textiles	338
12.6	Cosmetics	339
12.7	Medical Field	339
12.7.1	Imaging	340
12.7.2	Drug Delivery	341
12.7.3	Cancer Therapy	343
12.7.4	Tissue Repair	344
12.8	Agriculture and Food	345
12.9	Domestic Appliances	346
12.10	Space, Defense and Engineering	347
	Further Reading	348
13	Nanotechnology and Environment	349
13.1	Introduction	349
13.2	Environmental Pollution and Role of Nanotechnology	350
13.3	Effect of Nanotechnology on Human Health	352
	Further Reading	354

14	Practicals	355
14.1	Introduction	355
14.2	Synthesis of Gold/Silver Nanoparticles	356
14.2.1	Chemicals	356
14.2.2	Equipments	357
14.2.3	Synthesis Procedure	357
14.2.4	Results	358
14.3	Synthesis of CdS Nanoparticles	359
14.3.1	Chemicals	359
14.3.2	Equipments	359
14.3.3	Synthesis Procedure	359
14.3.4	Results	360
14.4	Synthesis of ZnO Nanoparticles	360
14.4.1	Chemicals	361
14.4.2	Equipment	361
14.4.3	Synthesis Procedure	361
14.4.4	Results	361
14.5	Synthesis of TiO ₂ Nanoparticles	362
14.5.1	Chemicals	363
14.5.2	Equipment	363
14.5.3	Synthesis Procedure	363
14.5.4	Results	363
14.6	Synthesis of Fe ₂ O ₃ Nanoparticles	364
14.6.1	Chemicals	364
14.6.2	Equipment	365
14.6.3	Synthesis Procedure	365
14.6.4	Results	365
14.7	Synthesis of Porous Silicon	366
14.7.1	Materials	366
14.7.2	Equipment	366
14.7.3	Experimental Procedure	367
14.7.4	Results	368
14.8	Introductory Photolithography	368
14.8.1	Background	368
14.8.2	Chemicals	371
14.8.3	Equipment	371
14.8.4	Experimental Procedure	372
14.9	Introductory Nano (Soft) Lithography Using PDMS	372
14.9.1	Chemicals	373
14.9.2	Equipment	373
14.9.3	Synthesis Procedure	373
14.9.4	Results	374
14.9.5	Materials and Equipments	375
14.9.6	Experimental Procedure	375

14.9.7	Experimental Set-Up	376
14.9.8	Results	376
14.10	Fabrication of Porous Alumina or Anodized Alumina (AAO) Template	376
14.10.1	Chemicals	376
14.10.2	Equipment	376
14.10.3	Fabrication Procedure	377
14.10.4	Results	378
	Further Reading	379
Appendices		381
Appendix I		381
	Periodic Table Symbols of Elements and Their Atomic Numbers	381
Appendix II		382
	Electromagnetic Spectrum	382
Appendix III		383
	List of Fundamental Constants	383
Appendix IV		383
	Vacuum Techniques	383
	Vacuum System	384
	Vacuum Pumps	385
	Rotary Vane Pump	385
	Diffusion Pump	386
	Sorption Pump	387
	Ion Pump	388
	Vacuum Gauges	389
	U-tube Manometer	389
	McLeod Gauge	390
	Pirani Gauge	390
	Thermocouple Gauge	391
	Cold Cathode Gauge (Penning Gauge)	391
	Hot Cathode Ion Gauge	392
	Bayard-Alpert (B-A) Gauge	393
	Further Reading	393
Appendix V		394
	Properties of Some Semiconductors	394
Appendix VI		395
	Kronig Penney Model (1-D)	395
	Further Reading	398
Index		399

Chapter 1

Introduction to Quantum Mechanics

1.1 Introduction

In Nanotechnology we are concerned with natural and synthetic materials in the size range of $\sim 1\text{--}100$ nm. At such a small size, very familiar classical, Newtonian mechanics or thermodynamics are not able to explain the observed properties of materials. We have to use quantum mechanics sometimes directly and sometimes through subjects like solid state physics or chemistry which use it to explain the properties and phenomenon of different materials. Those of you who are familiar with quantum mechanics and solid state physics can skip this and the next chapter and directly go to the third chapter. For those who would like to start new, let us discuss first the need of quantum mechanics and how it got developed so that it can be used to understand atoms, molecules, solids and nanomaterials. Box 1.1 gives some historical milestones, which have led to quantum mechanics.

Box 1.1: Historical Milestones in the Development of Quantum Mechanics

Pre-quantum Era

- In 1669, Newton proposed that light had corpuscular or particle nature.
- Huygen claimed in 1690 that light had a wave nature.
- Kirchoff and others studied black body radiation around 1860.
- Maxwell proposed (1873) theory of electromagnetic waves.
- In 1803–04 Young performed double slit experiment, which showed that light had a wave nature.
- In 1887, Heinrich Hertz produced and detected electromagnetic radiation.

(continued)

Box 1.1 (continued)**Old Quantum Theory Period**

- In 1901, Max Planck showed that energy distribution in black body radiation could be explained properly only if one considered that the radiation was quantized or had a particle nature.
- In 1905, Einstein proposed a theory of photoelectric effect which decisively proved that quantum or particle nature was associated with electromagnetic waves.
- Compton effect (1920) could be explained only when particle nature of electromagnetic radiation was considered, supporting Max Planck's and Einstein's theories. Particles of electromagnetic waves were identified as 'photons'.
- De Broglie (1923) argued that if electromagnetic waves were particles (photons) then why not particles have waves associated with them?
- Bohr's atom model (1913) with stationary states (why electrons should have some fixed energies) could be explained with de Broglie hypothesis.

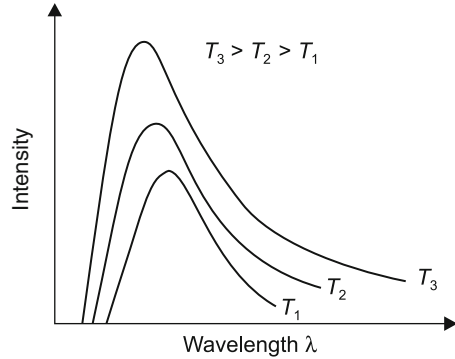
Modern Quantum Theory Begins

- Heisenberg introduced (1925) Matrix Mechanics.
- Schrödinger equation (1926) gave the firm foundation for de Broglie hypothesis and later explained the electronic structure of atoms, molecules and solids.
- Davisson and Germer showed in 1927 that electrons can be diffracted. Regularly spaced atoms constitute multi-slit analogue of Young's double slit experiment.
- Heisenberg proposed uncertainty principle in 1928.

This marks the beginning of Quantum Mechanics as it is practiced now!

In Kirchoff's time many scientists were interested in understanding the black body radiation (Box 1.2). When any radiation is incident on a body, what will it do? It will either reflect (r), absorb (a) or transmit (t), though partially, so that $r + a + t = 1$. A material is called a 'black body' (see Box 1.3) if it absorbs all radiation incident on it without reflecting or transmitting it. When a black body is heated, it gives out a spectrum as shown in Fig. 1.1. The black body spectrum spreads over a large range of wavelengths and has a maximum in the intensity. Experiments on black body radiation led Stefan (1879) and Boltzman (1884) to establish Stefan-Boltzman law. According to this law, the total radiation from a

Fig. 1.1 Spectra of Black body radiation. Note that as the temperature increases, the spectral intensity increases and maximum intensity shifts to shorter wavelength. T_1 , T_2 and T_3 are different temperatures in the increasing order



black body is proportional to the fourth power of absolute temperature. If E is the intensity of total radiation, and T the absolute temperature,

$$E = \sigma T^4 \quad (1.1)$$

where σ is known as Stefan's constant. Its value is $\sigma = 5.669 \times 10^{-8} \text{ W/m}^2$.

Later Wien showed that as the temperature of the black body increased, the maximum in the black body spectrum shifted to shorter wavelength or higher energy. To understand this, take a piece of charcoal or iron. At room temperature, they look black. If we heat these materials then what happens? First they look faint red, then they change to bright red and then become bright yellow or white. Wien's precise experiments on black body radiation resulted into a law which states that

$$\lambda_{\max} \times T = \text{Wien displacement constant} \cong 2.898 \times 10^{-3} \text{ mK} \quad (1.2)$$

where $\lambda_{\max}$ is the wavelength at which maximum intensity occurs for a black body held at temperature T . He also tried to write an expression, which could show how the intensity varied with wavelength as

$$E_{\lambda} = \frac{a}{\lambda^5} \exp\left(\frac{-b}{\lambda T}\right) \quad (1.3)$$

where ' a ' and ' b ' are constants. $E_{\lambda} d\lambda$ is the rate of energy emission per unit area in the wavelength range λ to $\lambda + d\lambda$. Therefore total energy radiated per unit time, per unit surface area is given by

$$E = \int_0^{\infty} E_{\lambda} d\lambda \quad (1.4)$$

Box 1.2: Sir Isaac Newton (1643–1727)

Isaac Newton was born on 4th January 1643 in Woolsthorpe, Lincolnshire, England. His father was a farmer who died two months before Isaac was born. His mother remarried and Newton had a difficult childhood. He graduated from Trinity College, Cambridge and studied Philosophy of Descartes, Gassendi, Hobbs and Boyle. He somewhere wrote, “Plato is my friend, Aristotle is my friend, but my best friend is truth”. He also studied Kepler’s optics and had a keen interest in mathematics. In 1665, he had to return home from college when an epidemic of plague broke.

In a period of 2 years at home, he did some revolutionary work in mathematics, physics and astronomy. He then laid the foundations of differential and integral calculus. He returned to college in 1667 and by 1669 he became well known due to his achievements in mathematics. In 1669, just at the age of 27, he was appointed as a Professor, Lucasian chair in Cambridge. His first lectures as a Professor were on Optics. All had believed since the time of Aristotle that white light was just a single component of radiation. But Newton’s work during 2 years at home during plague, made him think otherwise. Newton became the Fellow of Royal Society in 1672. He left Cambridge and took a Government position as a Warden of Royal Mint in 1696 and Master in 1699. Newton took his job very seriously and contributed a lot to the work of Mint. Officially he left his Cambridge position in 1701. Newton was a real genius and worked in almost all areas important at that time but he is best known for his work in



Laws of Gravity
Optics and
Co-foundation of Calculus

His famous books are

Philosophiae Naturalis Principa Mathematica (1687) and *Opticks* (1704). *Philosophiae Naturalis Principa Mathematica*, which is often referred to just as ‘Principa’, is considered to be one of the best scientific document ever written.

Newton was elected in 1703 as a President of Royal Society and remained so getting elected every year until his death on 31st March 1727 in London, England.

Box 1.3: Black Body Radiation

All materials absorb and emit energy. The intensity of energy radiated and absorbed by a body are equal if the object is in thermal equilibrium with its surrounding. However if the body is above the temperature of its surrounding then it emits radiation, which is known as black body radiation. Black body is thus an object that can absorb all radiation incident on it (no reflection or transmission!) or emits all radiation when above the temperature of the surrounding. Typical black body spectra are illustrated in Fig. 1.1.

Total intensity of radiation (area under the curve) and intensities at different emission wavelengths of a perfect black body irrespective of its material depend on the temperature. In practice a cavity with a small hole can act like a black body.

It was found that this formula holds good only for short wavelength side of the spectrum. Therefore scientists continued their efforts to find out a suitable equation which would be able to give or simulate the complete spectrum. Rayleigh and Jean arrived at an equation

$$E_{\lambda} = \frac{2c}{\lambda^4} kT \quad (1.5)$$

where ‘ c ’ is velocity of light and ‘ k ’ is Boltzman constant.

This equation could explain the long wavelength side of the black body spectrum. However it failed at short wavelengths. Thus none of the equations were satisfactory to explain the black body radiation over the entire range. The mystery continued until Planck proposed in 1901 a formula

$$E_{\lambda} = \frac{2\pi hc^2}{\lambda^5} \frac{1}{(e^{hc/\lambda KT} - 1)} \quad (1.6)$$

where h is Planck’s constant. Planck proposed that radiation cannot be absorbed or emitted continuously. One needs to consider the absorption or emission of radiation through some quantum of energy or ‘quanta’ like particles, later termed as photons. Energy of each quantum of radiation (absorbed or emitted) was assumed to be $h\nu$, where ν is the frequency of absorbed or emitted radiation. It can be shown that above equation is valid at short and long wavelengths equally well. The total radiated energy is given by

$$E = \frac{2\pi^5 k^4}{15c^2 h^3} T^4 \quad (1.7)$$

Thus Stefan's law follows from Planck's equation. Planck's equation can also be used to prove experimentally observed Wien's law. Planck's equation successfully explained all the regions of Black Body spectrum as well as previous findings and laws. Therefore Planck's idea that electromagnetic radiation should be considered as quantum of radiation turned out to be a milestone in the development of modern science (Box 1.4).

Box 1.4: Max Planck (1858–1947)

Max Planck was born in Kiel, Germany, in the year 1858. He studied physics in Munich as well as Berlin, Germany. He became a Professor of Theoretical Physics in 1892 in Berlin University. There he discovered in 1899 the fundamental constant ' h ' or Planck's constant named after him. Immediately in 1900 he discovered what is known now as 'Planck's law of black body radiation'.

This law became the foundation of quantum theory.

Max Planck was awarded the Nobel Prize for Physics in 1918. From 1930 to 1937, Planck was head of the Kaiser Wilhelm Gesellschaft zur Förderung der Wissenschaften (KWG, Emperor Wilhelm Society for the Advancement of Science) which was renamed after his death on 4th October 1947 in Göttingen, Germany as Max Planck Gesellschaft zur Förderung der Wissenschaften (MPG, Max Planck Society for the Advancement of Science). This institute continues to be one of the most important institutes for science in Germany.

**Box 1.5: Photoelectric Effect**

If two metal electrodes are placed in an evacuated tube, separated by a short distance, as illustrated in Fig. 1.2a then current flows in the circuit if cathode is irradiated with UV to visible light. This is known as photoelectric effect. It is easy to understand that the circuit is completed if electrons are emitted from cathode reaching the anode on illumination of the cathode. Einstein successfully explained in 1905 the observations by assuming that the incident beam of light behaved like photons or quanta of radiation proposed by Max Planck. His laws of photoelectricity are as follows:

(continued)

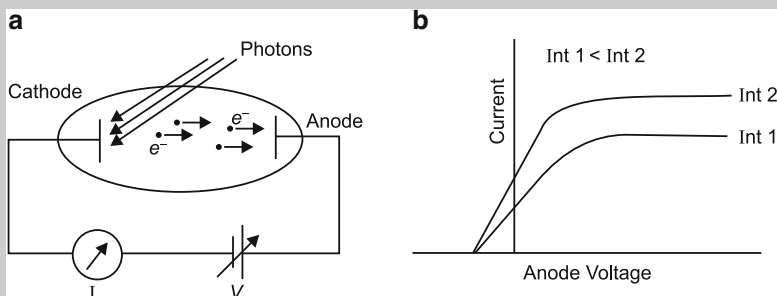
Box 1.5 (continued)

Fig. 1.2 (a) Circuit diagram to observe photoelectric effect. (b) Variation of current due to photoelectrons. Here 'Int' is the intensity of light

1. Photoelectric current (I) is proportional to the intensity of light (Int) falling on the cathode (see Fig. 1.2b).
2. There exists a threshold frequency dependent on the cathode material, which is necessary to emit photoelectrons. If light of lesser frequency is used then, even for any high intensity, no photoelectrons can be emitted.
3. This implies that there is a maximum kinetic energy which is necessary to produce photoelectric current.

$$k_{\max} = eV_0 = \frac{1}{2}mv^2 \quad (1.8)$$

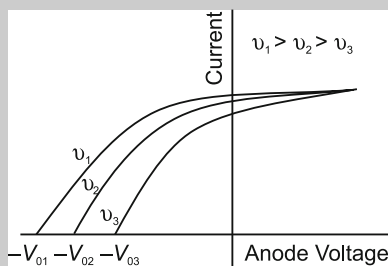
where V_0 is the stopping potential, e and m are electron charge and mass respectively. The velocity is v .

Even if anode is negative (V_0), electrons with maximum kinetic energy eV_0 can reach the anode.

Maximum energy of emitted photoelectrons does not depend upon the intensity of incident light but the frequency used as illustrated in Fig. 1.3.

4. Emitted photoelectrons have maximum energy (corresponding to $-V_0$) which depends upon the frequency of incident light.

Fig. 1.3 Negative anode voltage V_0 denotes the maximum kinetic energy the electrons can have in a photoemission process. It depends on the frequency ν of the incident light



(continued)

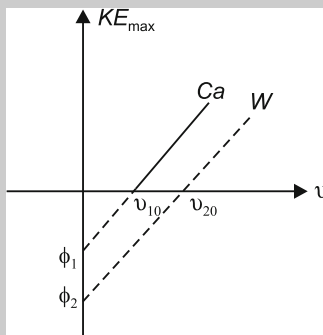
Box 1.5 (continued)

Fig. 1.4 Maximum kinetic energy depends upon the frequency of incident light and the material of cathode. Note that the lines have same slope (see Eq. 1.9). Here Ca and W are calcium and tungsten cathodes respectively. ϕ_1 and ϕ_2 are work functions of Ca and W respectively. Work function is a property of the material and denotes the amount of energy required to remove the electron from that material

Thus, number of ejected photoelectrons depends upon the intensity of light but the maximum kinetic energy of ejected photoelectrons depends upon the frequency of light. There is a minimum frequency ' ν_0 ' necessary to eject the photoelectrons which depends upon the material. This can be stated as

$$\text{Maximum kinetic energy} = h\nu - h\nu_0 \quad (1.9)$$

where $h\nu$ is the energy larger than minimum energy $h\nu_0$ required to eject the photoelectron (Fig. 1.4).

Box 1.6: Compton Scattering

X-rays with sufficiently high energy, when incident on a stationary electron as Compton imagined in 1920, can change their own direction as well as wavelength to longer side, reducing their energy. Reduced energy is imparted to the electron which gains a momentum $p = mv$ where m is the mass of the electron and v is the velocity gained by the electron.

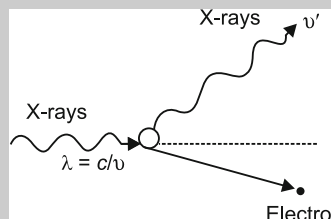
$$\text{K.E of electron} = h\nu - h\nu' \quad (1.10)$$

(continued)

Box 1.6 (continued)

where ν and ν' are the frequencies of incident X-rays before and after scattering (Fig. 1.5).

Fig. 1.5 Schematic of Compton effect $\nu' < \nu$



From conservation of momentum, momentum gained by the electron must be same as that lost by X-rays or we need to assume that X-rays which are electromagnetic waves have momentum $p = h/\lambda$ and X-rays behave as particles or photons.

Further, Einstein's theory of photoelectric effect (Box 1.5) and Compton effect (Box 1.6) could be explained only if one assumed the existence of photons or quanta of electromagnetic radiation.

However electromagnetic radiation also exhibits interference of light. This is quite evident from Young's diffraction experiments with single and double slits. It may, therefore, be concluded that it depends upon the type of experiment that electromagnetic radiation shows itself as waves or particles. Waves, of course, are continuous and particles are discrete in nature. The behaviour of electromagnetic waves sometimes as waves and sometimes as particles is termed as 'wave-particle duality' (Box 1.7).

Box 1.7: Albert Einstein (1879–1955)

Albert Einstein was born in Ulm, Germany, on 14th March 1879. At the age of four or five, his father showed him a compass and there begun his life long passion for physics. He wondered as a child how the mysterious force caused compass needle to move. He studied in a Catholic school in Munich, Germany. He was certainly not a top ranking student but he excelled in mathematics and physics. Even in school and college he was extremely original and independent in thinking.

He often challenged his teachers and the established scientific ideas. This perhaps resulted later into himself becoming a great scientist. At a very young age of 12, he studied geometry on his own and wrote his first scientific essay at the age of 16. He went to the University of Zurich for his college education and received Ph.D. there in 1905. In this year he wrote his four

(continued)

Box 1.7 (continued)

groundbreaking papers which were published in the journal ‘Annalen der Physik’. These are as follows:

1. “On a Heuristic Point of View Concerning the Production and Transformation of Light”

This deals with the photoelectric effect for which he received Nobel Prize in Physics in 1921.

2. “On the Movement of Small Particles Suspended in Stationary Liquids Required by the Molecular-Kinetic Theory of Heat”

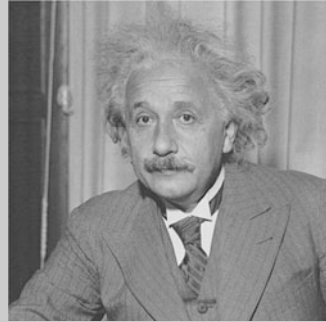
3. “On the Electrodynamics of Moving Bodies”

This paper is about his famous Theory of Relativity.

4. “Does the Inertia of a Body Depend upon its Energy Constant?”

$$\text{This paper has his famous equation } E = mc^2. \quad (1.11)$$

Interestingly, Einstein began his career as a Patent clerk at Swiss Federal Office for Intellectual Property in Bern. He spent 7 years there and was a guest lecturer for 1 year in the University of Bern. He then became the Professor of Theoretical Physics in the University of Zurich, Switzerland. Later he became a Swiss citizen. By 1909 he became very famous and visited many countries. In 1935 he went to U.S. and stayed there permanently in New Jersey. He was in the Institute of Advanced Studies in Princeton, New Jersey and made numerous scientific contributions. He also made until his death on 18th March 1955 numerous contributions to the world peace.



1.2 Matter Waves

It is quite logical to think that if electromagnetic waves sometimes behave like particles (photons) then why not particles behave like waves? This precisely was the point raised by de Broglie in 1923. He postulated that all the matter must have associated waves with wavelength given by

$$\lambda = \frac{h}{p} = \frac{h}{mv} \quad (1.12)$$

where h = Planck's constant, p = magnitude of the momentum, m = mass of the particle and v = velocity of the particle.

Equation 1.12 is known as the de Broglie relation and wavelength λ is known as de Broglie wavelength of the particle (Box 1.8).

Box 1.8: Louis de Broglie (1892–1987)

Louis de Broglie completed his school studies in Paris, France. He was interested in literary studies rather than science. He took a course in history at University of Paris as he wanted to make a career in diplomatic services. However he got interested in physics due to influence of his elder brother who was doing Ph.D. in experimental physics. So after getting a degree in Arts he enrolled himself for Ph.D. in physics.



However there broke the World War I and he had to serve in army. After 1920 he again returned to research in mathematical physics and was attracted by Planck's black body radiation and quantum theory concepts. In 1924 he put forth the particle-wave duality in his doctoral thesis for which he received the Nobel Prize for Physics in 1929. Some excerpt from his thesis which explains the background of his work is as follows.

Thirty years ago, physics was divided into two camps: ... the physics of matter, based on the concepts of particles and atoms which were supposed to obey the laws of classical Newtonian mechanics, and the physics of radiation, based on the idea of wave propagation in a hypothetical continuous medium, the luminous and electromagnetic ether. But these two systems of physics could not remain detached from each other: they had to be united by the formulation of a theory of exchanges of energy between matter and radiation. ... In the attempt to bring the two systems of physics together, conclusions were in fact reached which were neither correct nor even admissible. When applied to the energy equilibrium between matter and radiation ... Planck ... assumed ... that a light source ... emits its radiation in equal and finite quantities—in quanta. The success of Planck's ideas has been accompanied by serious consequences. If light is emitted in quanta, must it not, once emitted, possess a corpuscular structure? ... Jeans and Poincaré (showed) that if the motion of the material particles in a source of light took place according to the laws of classical mechanics, then the correct law of black-body radiation, Planck's law, could not be obtained.

He also once talked about his discovery as follows:

As in my conversations with my brother we always arrived at the conclusion that in the case of X-rays one had both waves and corpuscles, thus suddenly— ... it was certain in the course of summer 1923—I got the idea that one had to extend this duality to material particles, especially to electrons. And I realised that, on the one hand, the Hamilton-Jacobi theory pointed somewhat in that direction, for it can be applied to particles and, in addition, it represents a geometrical optics; on the other hand, in quantum phenomena one obtains quantum numbers, which are

(continued)

Box 1.8 (continued)

rarely found in mechanics but occur very frequently in wave phenomena and in all problems dealing with wave motion.

Thus I arrived at the following general idea which has guided my researches: for matter, just as much as for radiation, in particular light, we must introduce at one and the same time the corpuscle concept and the wave concept.

In other words, in both cases we must assume the existence of corpuscles accompanied by waves. But corpuscles and waves cannot be independent, since, according to Bohr, they are complementary to each other; consequently it must be possible to establish a certain parallelism between the motion of a corpuscle and the propagation of the wave which is associated with it.

De Broglie wrote many popular science books and was awarded first Kalinga Prize for the same in 1952 instituted by UNESCO.

This concept was initially hard to accept, as how can something be both wave and particle! However, later on, it turned out to be a very effective way of explaining various properties of atoms and subatomic particles like electrons, protons and neutrons. This concept of matter waves along with others discussed earlier laid down the foundation of quantum mechanics. ‘Matter Waves’ were not considered in classical mechanics.

1.3 Heisenberg’s Uncertainty Principle

Additionally it is also necessary to consider ‘uncertainty principle’ put forth by W. Heisenberg in 1928. This also has no classical analogue but is one of the basic concepts in quantum mechanics.

In order to understand uncertainty principle, consider an experiment as follows. If we have a source of monochromatic light (single wavelength) at the back of two slits as shown in Fig. 1.6, we would obtain a diffraction pattern i.e. bands of dark and light areas. Dark bands correspond to the regions where light does not strike and light areas correspond to areas where light is able to reach. Replace now the source of light with source of mono energetic electrons from an electron gun. We would be able to see an interference pattern similar to that with light. We would get the diffraction pattern of light and dark regions by exposure of photographic plate to electrons. If we make few more experiments by reducing the electron beam intensity, we can see some interesting effect (Box 1.9).

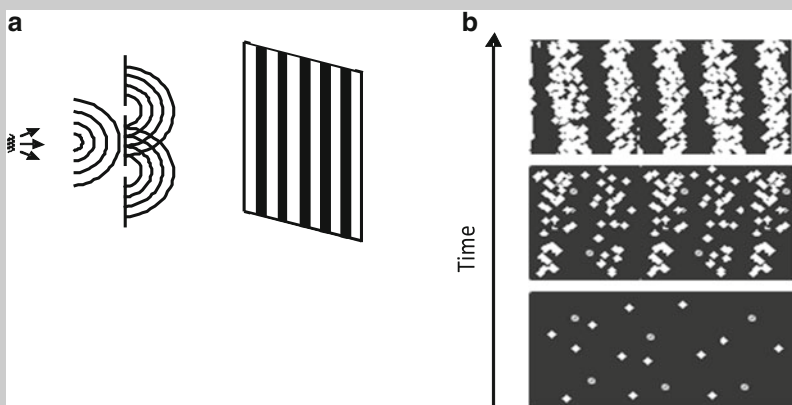
Box 1.9: Interference of Light and Electrons

Fig. 1.6 (a) Schematic diagram to obtain diffraction pattern (slits S1 and S2) on the photographic plate P (one may use some counter also to detect the intensity). One can use photons or electrons as the source. (b) Electron diffraction pattern for change in exposure (counting) time

For much reduced intensity we would see blurred interference pattern and with further reduction of electron beam intensity we may get only a couple of spots and no bands of light and dark areas. However electrons appear to have gone to places where intensity maxima appear in diffraction pattern. The question is how do the electrons know to reach the places where intensity maxima occur? Through which slits the electrons go? If we reduce the intensity of beam such that within the exposure time of photographic plate only a single electron went through the slit, should it go to the region with maximum intensity when number of electrons were large and produced diffraction pattern similar to that of light? How should a single electron know what kind of pattern would be produced if more electrons were there and it should not go to certain regions? How does a single electron know about the presence of the other slit? Or did it exist at two slits simultaneously? This is perhaps hard to accept as we think of any particle or body having precise location at a precise time. This means that it would appear at one of the slits at a given time (classical mechanics concept).

Can we now think of an experiment to determine through which slit the electrons go to produce diffraction pattern? For this, we can perform an experiment in which we shine the slits with light. Suppose we illuminate the slits with X-rays of short wavelength for accurate measurement. By measuring the scattered radiation we should know through which slit the electrons passed. This would give us a result showing that 50 % electrons went through one slit and 50 % through the other. However we would see that there is no interference or diffraction pattern produced at all. Our new experiment has destroyed the interference pattern. This is because, as

we know, due to Compton effect electron changes its energy and direction. Without X-rays striking the electrons, they went to produce interference pattern or there were only certain momenta allowed for the electrons. After Compton scattering the electrons got other momenta from the photons so that they could reach the previously forbidden regions. On the other hand we can use very long wavelength. Location of electron can be determined with a precision of $\lambda/2$. If wavelength is very large, precision would be lost and we would not know through which slit the electron came out. In other words we cannot precisely know position and momentum of the electron simultaneously with arbitrary accuracy. If we reduce the momentum uncertainty (using long wavelength), position becomes uncertain and if we measure the position with certainty using short wavelength then momentum becomes uncertain destroying the diffraction pattern. Indeed it is not possible to keep both position and momentum measurement precise, simultaneously. Precise measurement of one disturbs the other measurement. This is known as Heisenberg's uncertainty principle (Box 1.10). It is expressed as

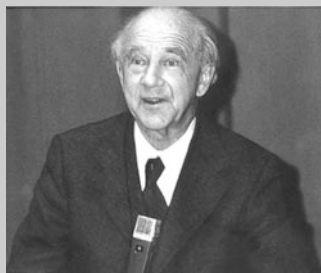
$$\Delta x \cdot \Delta p \geq \frac{h}{2\pi} \quad (1.13)$$

Box 1.10: W. Heisenberg (1901–1976)

Werner Heisenberg was born on 5th December 1901 in Würzburg, Germany. His father became a Professor of Greek language in University of Munich and W. Heisenberg went to Munich. He completed his school studies in Munich but studied under eminent physicists like Sommerfeld, Wien, Rosenthal, Max Born, Hilbert etc. at various places.

In 1926 he was appointed as a lecturer under Niels Bohr and within a year as a Professor of theoretical physics at University of Leipzig in Germany. In 1929, he went to U.S., Japan and India to deliver a lecture series.

Heisenberg published theory of quantum mechanics at the age of just 23. He received the Nobel Prize in Physics for the theory he developed and specially its applications which resulted into the discovery of allotropes of hydrogen. Later he was interested in plasma physics. He also worked on the unified field theory of elementary particles. In 1941 he was appointed as a professor of physics at the University of Berlin and Director of Kaiser Wilhelm Institute of Physics. In later years he was in Munich and died in 1976.



where Δx is the dispersion $\left(= \sqrt{\langle (x - \langle x \rangle)^2} \right)$ in the measurement of position in x direction and Δp is the dispersion in the simultaneous measurement of x component of momentum of a particle. It may be emphasized here that if we try to reduce the uncertainty in x or try to make a very precise measurement of position, the uncertainty in measurement of p increases and vice versa. The Eq. (1.13) is always valid irrespective of method of measurement. In other words, uncertainty principle gives the maximum possible accuracy in a simultaneous measurement of two quantities viz. position and momentum of a particle, in this case.

In general, Heisenberg's uncertainty principle states that "it is impossible to determine precisely and simultaneously the values of both the members of a particular pair of physical variables such as position components in a particular direction and momentum in that direction, energy of a bound state and its decay time etc."

1.4 Schrödinger Equation

In classical mechanics, in order to determine the position of a given particle under some conditions like its speed, initial velocity etc. we use Newton's equations. Similarly in quantum mechanics we make use of Schrödinger equation to understand the behaviour of subatomic particles (Box 1.11).

Box 1.11: Erwin Schrödinger (1887–1961)

Erwin Schrödinger was born on 12th August 1887, in Vienna, Austria. During his childhood he developed broad interests in various subjects from poetry to science. He studied from 1906 to 1910 in the University of Vienna. During World War I, he was an officer in artillery.

From 1920 he went on academic positions to various universities like in Stuttgart, Breslau and Zurich.

He worked on many problems in theoretical physics and produced significant work on atomic spectra, specific heats of solids, and physiological studies of colour. His greatest work, however, is the Schrödinger's wave equation which appeared in 1926 and he shared 1933 Nobel Prize for the same with Dirac. He however did not like the statistical interpretation of waves and generally accepted description in terms of waves and particles. In his own words:



(continued)

Box 1.11 (continued)

What we observe as material bodies and forces are nothing but shapes and variations in the structure of space. Particles are just schaumkommen (appearances). The world is given to me only once, not one existing and one perceived. Subject and object are only one. The barrier between them cannot be said to have broken down as a result of recent experience in the physical sciences, for this barrier does not exist. Let me say at the outset, that in this discourse, I am opposing not a few special statements of quantum mechanics held today (1950s). I am opposing as it were the whole of it. I am opposing its basic views that have been shaped 25 years ago, when Max Born put forward his probability interpretation, which was accepted by almost everybody. (Schrödinger Erwin, The Interpretation of Quantum Mechanics. Ox Bow Press, Woodbridge, CN, 1995).

I don't like it, and I'm sorry I ever had anything to do with it (Erwin Schrödinger talking about Born's Probability Wave Interpretation of Quantum Mechanics.)

In 1927 he went to Berlin but left it in 1933 when Hitler came into power. During World War II he moved to various places and settled in Dublin as director of School for Theoretical Physics. He returned to Vienna after he retired from there in 1955. He died on 4th January 1961.

For a free particle, in one dimension, time dependent Schrödinger equation is as follows:

$$\frac{-\hbar^2}{2m} \frac{\partial^2 \psi(x, t)}{\partial x^2} = i\hbar \frac{\partial \psi(x, t)}{\partial t} \quad (1.14)$$

where $\psi(x, t)$ is the wave function of particle of mass m . We neglect the potential energy V and assume that

$$E = \frac{\hbar^2 k^2}{2m} = \hbar\omega$$

which is the kinetic energy of the particle.

If we want to generalize the equation by writing $E = \text{kinetic energy} + \text{potential energy}$

i.e.
$$E = \frac{\hbar^2 k^2}{2m} + V$$

the Schrödinger equation becomes

$$\frac{-\hbar^2}{2m} \frac{\partial^2 \psi(x, t)}{\partial x^2} + V(x, t) \psi(x, t) = i\hbar \frac{\partial \psi(x, t)}{\partial t} \quad (1.15)$$

Equation (1.15) is one dimensional form of Schrödinger equation. A three dimensional form can be written as

$$\begin{aligned} & \frac{-\hbar^2}{2m} \left(\frac{\partial^2}{\partial x^2} + \frac{\partial^2}{\partial y^2} + \frac{\partial^2}{\partial z^2} \right) \psi(x, y, z, t) + V(x, y, z, t) \psi(x, y, z, t) \\ &= i\hbar \frac{\partial \psi(x, y, z, t)}{\partial t} \end{aligned} \quad (1.16)$$

We can use Laplacian notation viz.

$$\frac{\partial^2}{\partial x^2} + \frac{\partial^2}{\partial y^2} + \frac{\partial^2}{\partial z^2} = \nabla^2 \quad (1.17)$$

and $\mathbf{r}$ for (x, y, z) .

Equation (1.16) takes then a compact form as

$$\frac{-\hbar^2}{2m} \nabla^2 \psi(\mathbf{r}, t) + V(\mathbf{r}, t) \psi(\mathbf{r}, t) = i\hbar \frac{\partial \psi(\mathbf{r}, t)}{\partial t} \quad (1.18)$$

This is a time-dependent Schrödinger equation in three dimensions.

We can also get a time independent Schrödinger equation as follows. In case the potential energy of a particle does not change with time (as in case of energy states of electrons in an atom), we can write for one dimension

$$V(x, t) = V(x) \quad (1.19)$$

and Schrödinger equation as

$$\frac{-\hbar^2}{2m} \frac{\partial^2 \psi(x, t)}{\partial x^2} + V(x) \psi(x, t) = i\hbar \frac{\partial \psi(x, t)}{\partial t} \quad (1.20)$$

Assuming that wave function $\psi(x, t)$ can be separated into two parts viz. one varying with position and the other with time, it can be written as

$$\psi(x, t) = \Phi(x)\theta(t)$$

Substituting it in Eq. (1.20), we get

$$\frac{-\hbar^2}{2m} \theta(t) \frac{d^2 \Phi(x)}{dx^2} + V(x) \Phi(x) \theta(t) = i\hbar \Phi(x) \frac{d\theta(t)}{dt} \quad (1.21)$$

$$\text{or} \quad \frac{-\hbar^2}{2m \Phi(x)} \frac{d^2 \Phi(x)}{dx^2} + V(x) = \frac{i\hbar}{\theta(t)} \frac{d\theta(t)}{dt} \quad (1.22)$$

Note that in the above equation the left hand side is only position dependent and the right hand side is only time dependent. Therefore each side of the equation must be a constant, say A . We can therefore write

$$\frac{-\hbar^2}{2m \Phi(x)} \frac{d^2 \Phi(x)}{dx^2} + V(x) = A \quad (1.23)$$

and

$$\frac{i\hbar}{\theta(t)} \frac{d\theta(t)}{dt} = A \quad (1.24)$$

Solving Eq. (1.24), we obtain

$$\theta(t) = \theta(0)e^{\frac{-iAt}{\hbar}} \quad (1.25)$$

where $\theta(0)$ is $\theta(t)$ at time $t = 0$.

Value of constant A can be found out as follows.

Compare equation

$$i\hbar \frac{\partial \psi(x, t)}{\partial t} = E \psi(x, t)$$

and

$$\frac{i\hbar}{\theta(t)} \frac{d\theta(t)}{dt} = A$$

They are similar at $x = 0$. This implies that $A = E$. Therefore Eq. (1.24) can be rewritten as

$$\frac{d\theta(t)}{\theta(t)} = -iE \frac{dt}{\hbar} \quad (1.26)$$

Integrating it, we obtain

$$\theta(t) = \exp\left(\frac{-iEt}{\hbar}\right) \quad (1.27)$$

Substituting A in Eq. (1.23) with E now, we get

$$\frac{-\hbar^2}{2m\Phi(x)} \frac{d^2\Phi(x)}{dx^2} + V(x) = E \quad (1.28)$$

or

$$\frac{-\hbar^2}{2m} \frac{d^2\Phi(x)}{dx^2} + V(x)\Phi(x) = E\Phi(x) \quad (1.29)$$

This is one dimensional, time independent Schrödinger equation. For three dimensions we can write

$$\psi(\mathbf{r}, t) = \psi(\mathbf{r}) \Phi(t) \quad (1.30)$$

and time independent three dimensional Schrödinger equation becomes

$$\frac{-\hbar^2}{2m} \nabla^2 \psi(\mathbf{r}) + V(\mathbf{r}) \psi(\mathbf{r}) = E \psi(\mathbf{r}) \quad (1.31)$$

In order to determine the energy states of a particle it is necessary to know the form of potential V and how particle amplitude varies in space or $\psi(\mathbf{r})$. There may be a number of wave functions satisfying Schrödinger equation, but all may not be able to describe a given situation. The acceptable wave function should satisfy physical boundary conditions. It is necessary that ψ should satisfy following conditions so that using Schrödinger equation one can obtain the realistic values of energy.

1. $\psi(\mathbf{x}, \mathbf{y}, \mathbf{z})$ must be finite and single valued at all points in coordinate space.
2. It is necessary that $\psi(x, y, z)$ is a continuous function at all points in coordinate space.
3. Derivatives of ψ viz. $\frac{\partial \psi}{\partial x}$, $\frac{\partial \psi}{\partial y}$ and $\frac{\partial \psi}{\partial z}$ also need to be finite, single valued and continuous functions in coordinate space.
4. $\int |\psi|^2 d\tau$ must be finite, where $d\tau$ is the volume.

It is convenient to use complex wave form of ψ .

$$\psi = u + i v \text{ with } u \text{ and } v \text{ as real functions.} \quad (1.32)$$

$$\psi^* = u - i v \text{ is the complex conjugate of } \psi. \quad (1.33)$$

$$\psi \psi^* = u^2 + v^2 \quad (1.34)$$

Although ψ and ψ^* do not have physical significance as such, $\psi \psi^*$ or $|\psi|^2$ is interpreted as the probability density of particle being there in integration limits of $\int |\psi|^2 d\tau$ (Box 1.12).

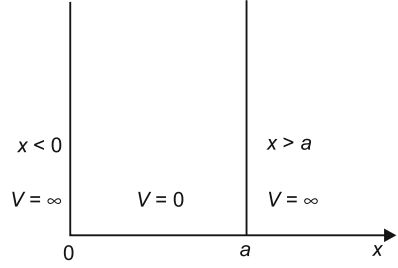
Box 1.12: Physical Interpretation of de Broglie Wave

1. The wave function $\psi(\mathbf{r}, t)$ is useful to determine the probability of finding a moving particle.
2. $\psi(\mathbf{r}, t)$ may be complex.
3. $|\psi(\mathbf{r}, t)|^2$ is the probability of finding the particle at time t at a point given by $\mathbf{r}$ in real space.
4. As the particle should be found somewhere in space it can be obtained using the normalization condition $\int |\psi(\mathbf{r}, t)|^2 d\tau = 1$.

1.5 Electron Confinement

As briefly mentioned in the introduction of this chapter, nano structured materials have at least one of the dimensions in the range of 1–100 nm. In such cases properties of materials can be understood using quantum mechanics rather than classical mechanics. Even the properties of atoms, molecules or extended solids can be understood using quantum mechanics because in order to understand various

Fig. 1.7 One dimensional potential box



properties even in case of bulk solids we need to understand the electron properties which are subatomic particles. For detailed descriptions it is necessary to refer to books devoted to atoms and molecules and solid state physics. Here we shall outline the essential features by which one can understand the methodology.

We shall discuss in this section confinement of a particle in one dimension (freedom in 2D), confinement in two dimensions (freedom in 1D) and confinement in all three dimensions (no freedom in any direction or 0D material). These confinements when referred to practical cases are often known as essentially 2D quantum wells (thin film), 1D wire (wire) and 0D quantum dot (nanoparticles).

1.5.1 Particle in a Box

Consider a box of length ‘ a ’ such that

$$\begin{aligned} \text{Potential} \quad & V = 0 \quad \text{if } 0 < x < a \\ \text{and} \quad & V = \infty \quad \text{if } x < 0 \text{ or } x > a \end{aligned}$$

as illustrated in Fig. 1.7.

Energy states of the particle of mass m can be obtained using time independent Schrödinger equation for one dimension as

$$-\frac{\hbar^2}{2m} \frac{\partial^2}{\partial x^2} \psi(x) + V(x)\psi(x) = E\psi(x) \quad (1.35)$$

Let $\psi(x)$ have a general form as

$$\psi(x) = A \sin \left(\frac{2mE}{\hbar^2} \right)^{\frac{1}{2}} x + B \cos \left(\frac{2mE}{\hbar^2} \right)^{\frac{1}{2}} x \quad (1.36)$$

As the particle exists only inside the box, wavefunction should not exist outside the box and should be zero at the boundaries.

At $x = 0$, boundary condition $\psi(x) = 0$ leads to $B = 0$. Therefore,

$$\psi(x) = A \sin \left(\frac{2mE}{\hbar^2} \right)^{\frac{1}{2}} x \quad (1.37)$$

At $x = a$, boundary condition $\psi(a) = 0$ leads to

$$\psi(a) = 0 = A \sin\left(\frac{2mE}{\hbar^2}\right)^{\frac{1}{2}} a \quad (1.38)$$

But ' a ' is not zero, therefore

$$\sin\left(\frac{2mE}{\hbar^2}\right)^{\frac{1}{2}} a = 0 \quad \text{or,} \quad \left(\frac{2mE}{\hbar^2}\right)^{\frac{1}{2}} a = n\pi$$

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

Therefore

$$E_n = \frac{n^2 \hbar^2 \pi^2}{2ma^2} = \frac{n^2 h^2}{8ma^2} \quad (1.39)$$

$$\text{i.e.,} \quad E_n \propto n^2 \quad (1.40)$$

Putting E_n in Eq. (1.37) we get

$$\psi_n = A \sin\left(\frac{n\pi}{a}\right) x \quad (1.41)$$

Although n can take any integer value according to Eq. (1.39), in practice $n = 0$ and $\psi_n = 0$ are not allowed inside the box, because, if allowed, $|\psi|^2$ would be 0 and probability of finding the particle inside the box would be zero. For the same reason ψ cannot be zero inside the box. Therefore n takes the values $n = 1, 2, 3, \dots$. This shows that energies of particle in a one dimensional potential box are quantized and can be illustrated as in Fig. 1.8a.

Corresponding wave functions and probabilities of different states of particle in the box would look like those in Fig. 1.8b, c respectively.

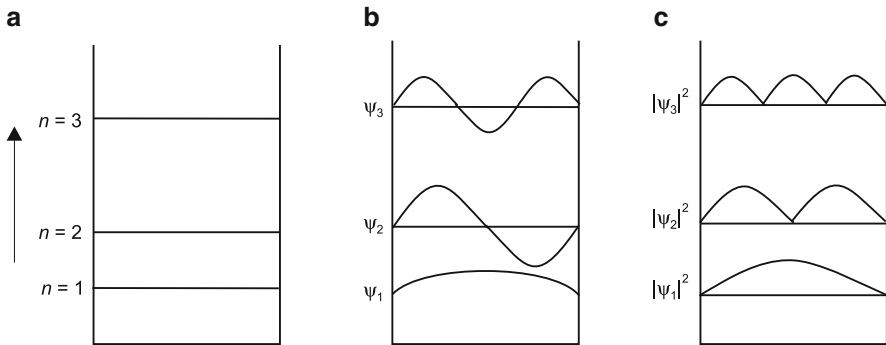


Fig. 1.8 (a) Quantized energy states, (b) corresponding wavefunctions and (c) probability of finding the particle at different locations between '0' and 'a' in the box

1.5.2 Density of States

Density of states $D(E)$ is an important quantity. It enables to gain understanding of various spectroscopic and transport properties of materials. It is defined as the number of states per unit energy range and in general can be obtained as follows. Consider energy as is given for a particle in a 1D box.

$$E_n = n^2 = \frac{h^2}{8ma^2} \quad (\text{Henceforth we drop suffix of } E)$$

$$dE = \frac{h^2}{8ma^2} \cdot 2n \cdot dn \quad (1.42)$$

$$\frac{dn}{dE} = \frac{8ma^2}{h^2} \cdot \frac{1}{2n} = \frac{a}{h} \sqrt{\frac{2m}{E}} \quad (1.43)$$

$$D(E) = \frac{dn}{dE} \propto E^{-1/2} \quad (1.44)$$

The Eq. (1.44) shows how the density of states will vary with change of energy. Depending upon the energy state's function, density of states would take form as briefly stated below.

1.5.2.1 Density of States for a Zero Dimensional (0D) Solid

Neglecting the periodic potential existing in solids, we can imagine a zero dimensional solid in which electron is confined in a three dimensional potential box with extremely small (<100 nm) length, breadth and height as a 0D solid. This will have the discrete energy levels as discussed above with density of states given as

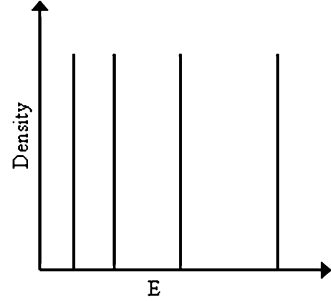
$$D(E) = \frac{dN}{dE} = \sum_{\varepsilon_i} \delta(E - \varepsilon_i) \quad (1.45)$$

where ε_i are discrete energy levels and δ is Dirac function. The density of states as a function of energy would appear as illustrated in Fig. 1.9.

1.5.2.2 Density of States in a One Dimensional (1D) Potential Box, Wire

A particle confined in one dimension is like a particle in a one dimensional potential well as in the previous case. However the potential in two directions is infinitely large but length is not very small. This gives rise to density of states as follows:

Fig. 1.9 Density of states for a particle in a zero dimensional solid



$$D(E) = \frac{dN}{dE} = \sum_{\varepsilon_i < E} \delta(E - \varepsilon_i)^{-1/2} \quad (1.46)$$

where ε_i are discrete energy levels. Figure 1.10 graphically illustrates nature of density of states for a one dimensional solid.

1.5.2.3 Density of States in a Two Dimensional (2D) Potential Box, Thin Film

It can be shown that the density of states ($D(E)$) in the two dimensional solid, which is nothing but the case of thin films, is given as

$$D(E) = \frac{dN}{dE} = \sum_{\varepsilon_i < E} 1 \quad (1.47)$$

Thus it can be shown that the density of states in two dimensional case is constant (see Fig. 1.11). In this case $\frac{dN}{dE}$ would correspond to the states available in an area.

1.5.2.4 Density of States for a Particle in a Three Dimensional Box

Following the previous cases it is a straightforward task to show that in a box of length ' a ', width ' b ', and height ' c ' with potential $V = 0$ inside the box and $V = \infty$ outside the box, the energy states can be obtained as

$$E_{n_x, n_y, n_z} = \frac{\hbar^2}{2m} (n_x^2 + n_y^2 + n_z^2) \quad (1.48)$$

Fig. 1.10 Density of states for a particle in a one dimensional solid

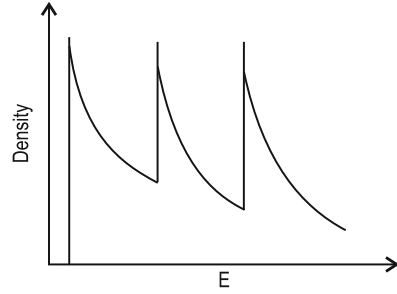


Fig. 1.11 Density of states for a 2D potential box

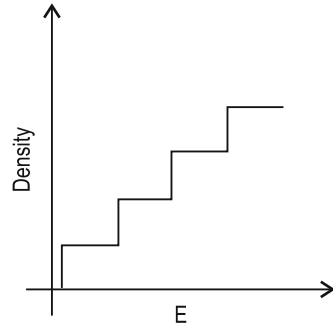
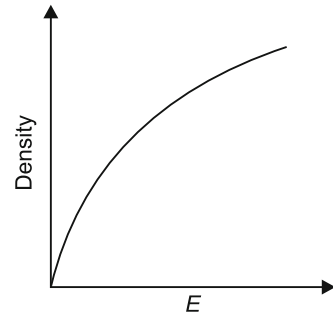


Fig. 1.12 Density of states for a particle in a 3D potential box



The wave function has the form

$$\psi_n(x, y, z) = A \sin\left(\frac{\pi n_x x}{a}\right) \sin\left(\frac{\pi n_y y}{b}\right) \sin\left(\frac{\pi n_z z}{c}\right) \quad (1.49)$$

Hence to consider the density of states we shall require the volume. It can be shown that

$$D(E) \propto E^{1/2} \quad (1.50)$$

This is graphically illustrated in Fig. 1.12.

1.5.3 Particle in a Coulomb Potential

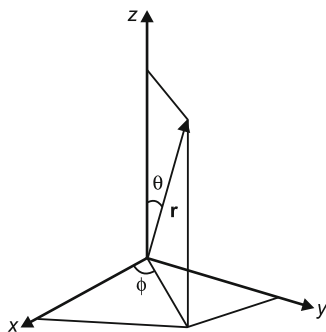
So far we have considered the particle confined to a potential box in which $V = 0$ inside the box and abruptly outside the box $V = \infty$. We shall consider now another type of potential viz. Coulomb potential experienced by a particle. The potential can be written as

$$V(r) = \frac{-Ze^2}{r} \quad (1.51)$$

This type of potential can arise when a particle moving in three dimensions is attracted towards the centre. For example an electron with negative charge attracted towards a positively charged nucleus at the centre of an atom.

In Eq. (1.51) ' Ze ' is the charge due to Z number of electrons and r is the distance between nucleus and the charge. The negative sign arises by convention to show attraction between oppositely charged particles.

It is convenient to solve this problem in spherical coordinates rather than in Cartesian coordinates by writing ψ , the wave function in spherical coordinates $\psi(r, \theta, \phi)$.



$$x = r \sin \theta \cos \phi \quad (1.52)$$

$$y = r \sin \theta \sin \phi \quad (1.53)$$

$$z = r \cos \theta \text{ and } r = \sqrt{x^2 + y^2 + z^2} \quad (1.54)$$

The time independent Schrödinger equation in spherical coordinates can be written as

$$\nabla^2 \psi(\mathbf{r}, \theta, \phi) + \frac{2m}{\hbar^2} [E - V(\mathbf{r})] \psi(\mathbf{r}, \theta, \phi) = 0 \quad (1.55)$$

We shall not go here in the details of solving this equation but state the major steps and result. It is the spherically symmetric form of the potential which enables us to simplify the procedure.

Wavefunction $\psi(\mathbf{r}, \theta, \phi)$ is separable into functions of $\mathbf{r}$, θ and ϕ as

$$\psi = R(r)Y(\theta, \phi) \quad (1.56)$$

It can be written as

$$\psi_{n\ell m}(r, \theta, \phi) = \frac{U_{n\ell}(r)}{r} Y_{\ell m}(\theta, \phi) \quad (1.57)$$

$Y_{\ell m}$ are spherical functions.

The energy states can be obtained by considering one dimensional form of Eq. (1.55) as follows

$$\frac{-\hbar^2}{2m} \frac{d^2 U}{dr^2} + \left[V(r) + \frac{\hbar^2}{2mr^2} \ell(\ell+1) \right] U = EU \quad (1.58)$$

Energy state of the system is described with three quantum numbers: n (principal quantum number), ℓ (orbital quantum number) and m (magnetic quantum number).

The angular momentum of the state is given by L , as

$$L^2 = \ell(\ell+1)\hbar^2 \text{ where } \ell = 0, 1, 2, \dots (n-1) \quad (1.59)$$

The magnetic quantum number m essentially gives component of angular momentum L parallel to z axis as

$$L_z = m\hbar \text{ where } m = 0, \pm 1, \pm 2, \dots \pm \ell \quad (1.60)$$

The states are often denoted by $s, p, d, f \dots$ to mean $\ell = 0, 1, 2$ etc. respectively. If we consider now a Coulomb potential experienced by an electron of the form

$$V(r) = \frac{-Ze^2}{r}$$

where Z is atomic number, a situation in hydrogen atom (H), helium ion (He^+), lithium ion (Li^{2+}) etc.

Corresponding Schrödinger equation can be written as

$$\frac{\partial^2 \psi}{\partial x^2} + \frac{\partial^2 \psi}{\partial y^2} + \frac{\partial^2 \psi}{\partial z^2} + \frac{8\pi^2 m}{h^2} \left(E + \frac{Ze^2}{r} \right) \psi = 0 \quad (1.61)$$

The energy states by solving the equation in spherical coordinates gives

$$E_n = \frac{-2\pi^2\mu Z^2 e^4}{n^2 h^2} \quad (1.62)$$

where μ is the electron mass corrected for motion of nucleus

$$\frac{1}{\mu} = \frac{1}{m} + \frac{1}{M} \quad (1.63)$$

Note that n is an integer which takes the values,

$$\begin{aligned} n &= \ell + 1, \ell + 2, \ell + 3 \dots \\ \ell &= 0, 1, 2, 3 \dots (n - 1) \end{aligned} \quad (1.64)$$

One can see here that E_n is proportional to $1/n^2$, which means that energies will come closer as n increases, as opposed to potential in a box with potential inside as zero.

1.6 Tunnelling of a Particle Through Potential Barrier

Consider a situation in which a matter wave is incident on a potential barrier as illustrated in Fig. 1.13 from negative x direction. For the sake of simplicity, we shall consider one dimensional case. Energy of the particle is purely kinetic in region I and potential $V = 0$. However, $E < V_0$ in region I, where V_0 is the potential at $x = 0$. It can be shown using Schrödinger equation that for particle with less energy than the potential energy as in region I, there certainly exists a possibility that particle can not only enter the region II but also get transmitted in region III and propagate as long as V_0 is not infinite. This is impossible in classical mechanics. To understand

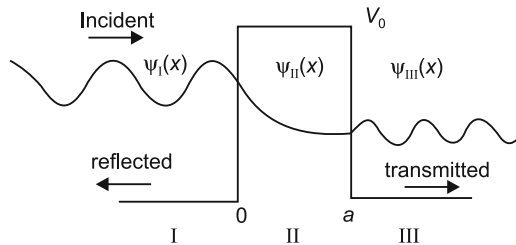


Fig. 1.13 Particle tunnelling through a potential barrier of height V_0 . The kinetic energy of the particle is $E < V_0$. The particle is incident from left side in region I and by tunnelling through region II, escapes to region III with some probability of transmission, given by Eq. (1.73)

this phenomenon of ‘tunnelling’ in quantum mechanics, consider that we have three different wave functions in regions I, II and III as follows

$$\psi_I(x) = A_1 e^{ikx} + A_2 e^{-ikx} \quad (1.65)$$

$$\psi_{II}(x) = B_1 e^{\alpha x} + B_2 e^{-\alpha x} \quad (1.66)$$

$$\psi_{III}(x) = C_1 e^{ikx} + C_2 e^{-ikx} \quad (1.67)$$

In all these wave functions we have considered that wave function is composed of two parts viz. a wave in positive direction (denoted by e^{ikx} and $e^{\alpha x}$) and negative direction (denoted by e^{-ikx} and $e^{-\alpha x}$). Waves in regions I and III are travelling waves (but that in region II is only decaying part) as

$$k = \sqrt{\frac{2mE}{\hbar^2}} \quad (1.68)$$

and $\alpha = \sqrt{\frac{2m(V_0 - E)}{\hbar^2}}$ are positive. (1.69)

We now need to use boundary and normalization conditions to determine the coefficients A_1 , A_2 , B_1 , B_2 , C_1 and C_2 in Eqs. (1.65), (1.66), and (1.67). We may further note that in region I, there is a probability that incident wave gets reflected but in region III there is no reflection. It can be shown after solving the equations that we can get tunnelling probability T defined as

$$T = \frac{|\psi_{III}(x)|^2}{|\psi_I(\text{incident})|^2} = \frac{C_1^2}{A_1^2} = \frac{1}{1 + D \sinh^3(\alpha a)} \quad (1.70)$$

with $D = \frac{V_0^2}{4E(V_0 - E)}$. (1.71)

If $\alpha a \gg 1$ i.e. potential barrier V_0 is extremely higher and the barriers width is very large, then we can use an approximation

$$\sinh \alpha a \approx \frac{1}{2} e^{\alpha a} \quad (1.72)$$

In this case it can be shown that

$$T = T_0 e^{-2\alpha a} \quad (1.73)$$

where $T_0 = \frac{16E(V_0 - E)}{V_0^2}$. (1.74)

Thus we can notice that there is always a probability of finding the particle on the other side of a potential barrier even though its kinetic energy is less than the potential barrier it sees. This result is very useful in understanding many phenomena observed for subatomic particles and cannot be explained by classical mechanics.

One can also show that reflection probability of wave in region I is

$$R = \frac{A_2^2}{A_1^2} = 1 - T \quad (1.75)$$

As mentioned earlier, the phenomenon of tunnelling through a potential barrier is possible only in quantum mechanics with no classical analogue. This means that a particle with kinetic energy smaller than that required to overcome the potential energy of the barrier does have some probability of being on the other side of the barrier. This is quite important to explain quantum well structures, solid state lasers, light emitting diodes, particles inside the nucleus etc.

Further Reading

M. Chandra, *Atomic structure and chemical bond* (Tata McGraw-Hill, New Delhi, 1991)

R. Eisberg, R. Resnick, *Quantum physics of atoms, molecules, solids, nuclei and particles*, 2nd edn. (Wiley, New York, 1985)

C.H. Holbrow, J.N. Lloyd, J.C. Amato, *Modern introductory physics* (Springer, New York, 1998)

S.O. Kasap, *Electronic materials and devices*, 2nd edn. (Tata McGraw-Hill, New Delhi, 1991)

L.I. Schiff, *Quantum mechanics*, 3rd edn. (McGraw Hill International Editions, New York, 1968)

Chapter 2

Structure and Bonding

2.1 Introduction

Matter is composed of atoms and molecules. Gases, liquids and solids are different states of matter (there can be some other states also under some extreme conditions of matter). However they can get converted into one another depending upon their stability at different pressure, temperature or variation of both. For example, as illustrated in Fig. 2.1, water molecule (H_2O) can be in gaseous state (vapour), liquid state (water) or solid state (ice) depending on the temperature and pressure. Water vapour can be converted into liquid water or ice by cooling. Similarly ice can be turned into water and then vapour by heating. At the triple point, solid, liquid and gas co-exist. Critical point as shown in Fig. 2.1 is a point at which some critical values of temperature and pressure are reached. Beyond this point there is no distinction between gas and liquid phase.

Density of gas is usually very low $\sim 10^{19}$ atoms/cc, whereas that of solids and liquids is high viz $\sim 10^{23}$ atoms/cc (we are not considering here porous forms of solids). Correspondingly the distances between atoms (or molecules) in gases are 3 to 4 nm. In solids and liquids the distances between the atoms or molecules are typically 0.2–0.4 nm, an order of magnitude smaller than in gases.

One can see that the density and distances between the atoms or molecules for both liquids and solids are comparable. However a liquid does not have a particular shape and takes the shape of its container. Atoms or molecules in a liquid, as in a gas, are under a state of continuous, random motion known as Brownian motion. In solids the situation is different. There the atoms and molecules also are in a state of motion, but the motion is about their fixed positions. Atoms and molecules vibrate about their fixed positions with some characteristic frequencies.

Let us now consider a solid like Al_2O_3 . It can appear as a transparent, opaque or translucent even without change of its chemical formula or presence of impurities. Why is it so? What is the origin of such a change one may wonder!

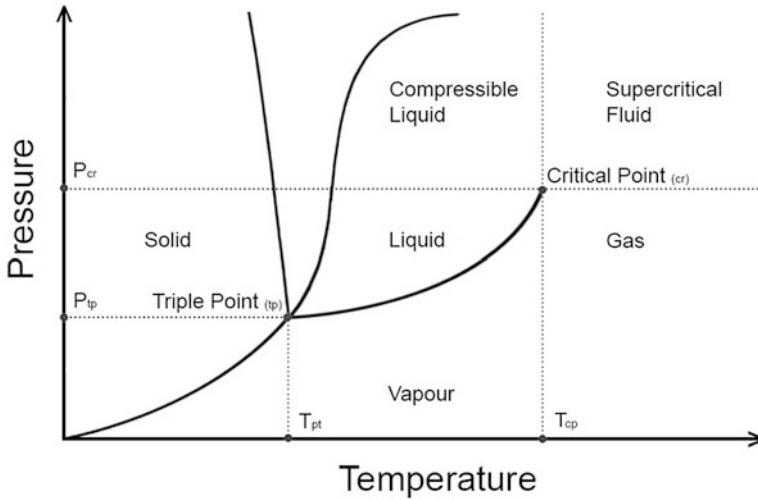


Fig. 2.1 Phase diagram showing different states of matters, which depend upon pressure and temperature

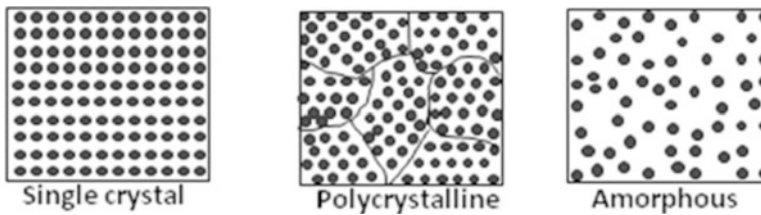


Fig. 2.2 Different types of solids

Figure 2.2 illustrates schematically three different forms of solids known as single crystal, polycrystalline and amorphous. In what is known as a ‘single crystal’ there is almost infinitely long arrangement of atoms or molecules with certain symmetry characteristics of the material. When we say infinite distance, it means a length many a times larger than the distance between two atoms.

In a polycrystalline solid, there are some ‘grain boundaries’. Size of the grain can depend upon the processing and typically can be few μm^3 . Each grain itself is a ‘single crystal’ but the orientations of these different crystals are different or random. Each grain also has a kind of ‘grain wall’ in which atoms may be more or less randomly distributed. The thickness of such walls is often very crucial in determining the various properties of materials such as mechanical, optical or electrical. At the grain boundary the long range periodicity of the solid is broken.

If each grain in the material becomes too small, comparable to the distance between the atoms or molecules, we get what is known as ‘amorphous’ solid. In amorphous solids, the grain boundaries disappear. Although the distances and even

the arrangement between nearest or even next nearest atoms may look similar for most of the atoms, they lack long-range order as in a poly or single crystal. The situation is similar to a liquid where a snapshot of atoms if taken would be similar to what we get in amorphous solid.

Thus transparent, precious form of Al_2O_3 known as ‘sapphire’, a gemstone, is a single crystal; its translucent form is a polycrystalline and the opaque form is an amorphous solid. Further, with same type of atoms or molecules even different single crystals can be formed with altogether different properties. Consider for example diamond and graphite, two well known crystalline forms of carbon. Diamond is transparent, colourless and very hard. Graphite with the same chemical constituent viz. carbon is soft, black and opaque. Diamond is electrically insulating but graphite is conducting. Why two materials made of the same constituents should differ so much? For this, one should try to understand how the atoms in two materials are arranged, how the arrangement of atoms affects the bonding of atoms and ultimately the properties. In fact differences in the properties are usually correlated to the arrangement of atoms or structure and bonding. By studying a large number of solids, scientists have been able to understand how atoms are arranged in different crystal structures. This has led them to make a classification and develop some useful techniques by which one can identify atomic arrangements in some unknown sample. We shall start with some definitions in order to understand all this.

2.2 Arrangement of Atoms

Lattice: It is an arrangement of points repeated in one, two or three directions making it a one dimensional, two dimensional or three dimensional lattice.

However, for the sake of simplicity we may often take examples of two-dimensional lattices (Fig. 2.3).

Crystal: When an atom or a group of atoms are attached to each lattice point, it forms a crystal (Fig. 2.4).

Unit Cell: Crystal structures can be understood in terms of unit cell, which when translated in any direction can fill the complete space without leaving any space in between or without overlapping. *Unit cell is a conventional cell* and can be of any volume, size, orientation or shape (see Fig. 2.5). However it is often taken in such a way that it has atoms at some corners or centre.

Primitive Cell: Unit cell with smallest volume is known as ‘*primitive cell*’ (see Fig. 2.6). It is also necessary to consider the *unit cell vectors* which define the

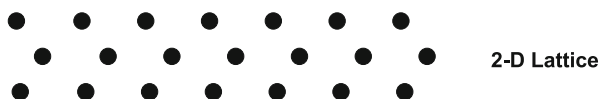


Fig. 2.3 A periodic arrangement of points in two dimensions makes a lattice

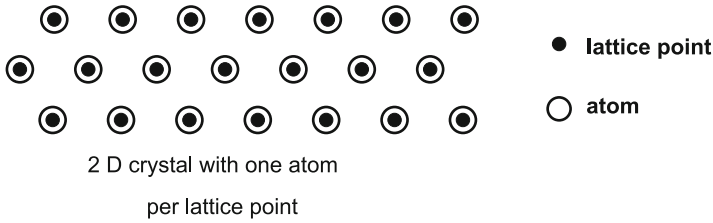


Fig. 2.4 A lattice + atom makes a solid or crystal

Fig. 2.5 Different unit cells in the same arrangement of atoms

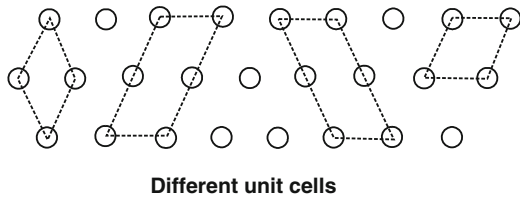
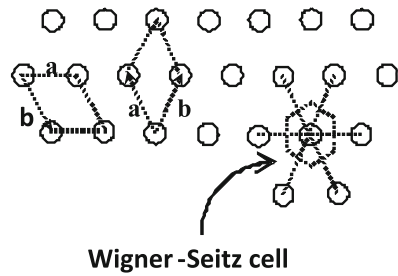


Fig. 2.6 Some primitive cells



boundaries or edges of unit cell. It can be seen that in direction of **a**, if we translate the cell by 'a' unit, the cell is repeated. Same is true for **b**. Thus **a** and **b** are translation vectors.

Translation vectors of primitive or smallest unit cell are known as 'primitive vectors'. It is not always necessary to consider primitive vectors but one can have 'unit vectors' of any convenient length.

$$\text{Volume of the cell } V = \mathbf{a} \cdot \mathbf{b} \times \mathbf{c} \quad (2.1)$$

As is clear from above figure, the choice of unit cell, unit vector, primitive cell or primitive vector is not unique.

In order to classify the crystals further, it is necessary to understand certain symmetry operations around a point (also known as point operation, which helps to classify the crystals).

Symmetry: Symmetry transformation about a point is the one which leaves the system invariant or without any change even after some operation.

Fig. 2.7 A circle can be rotated around an axis passing through its centre through any angle and would appear as if unmoved

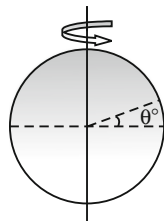


Fig. 2.8 Point symmetry of a square

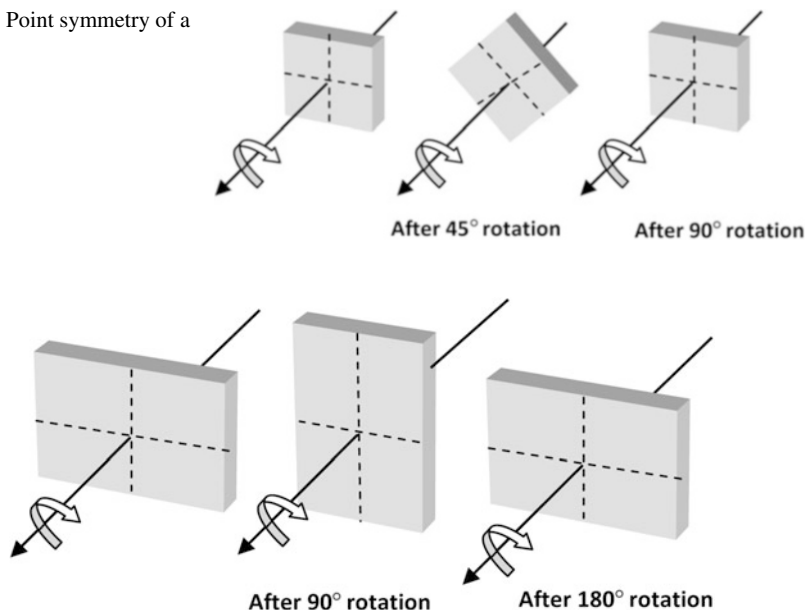


Fig. 2.9 Point symmetry of a rectangle

Rotational Symmetry: Consider for example a circle (Fig. 2.7). Circle is a perfectly symmetric object which can be rotated about an axis through its centre by any angle θ without it appearing as if it was rotated. Thus by rotating through any angle we make an invariant transformation for a circle or a sphere.

But this is not true for any lower symmetry object. Consider now a square (Fig. 2.8). You can rotate it by 90° about an axis passing through its midpoint which appears as if it has not been rotated. This can be done only four times successively in any direction (clockwise or anticlockwise). The square is said to have a fourfold symmetry. The square would also have a twofold symmetry.

Consider now a rectangle. If you rotate it through 90° , you see the difference as shown in Fig. 2.9, but if you rotate it through 180° it looks similar to the original state. This you can do twice. Therefore, rectangle is an object with twofold symmetry.

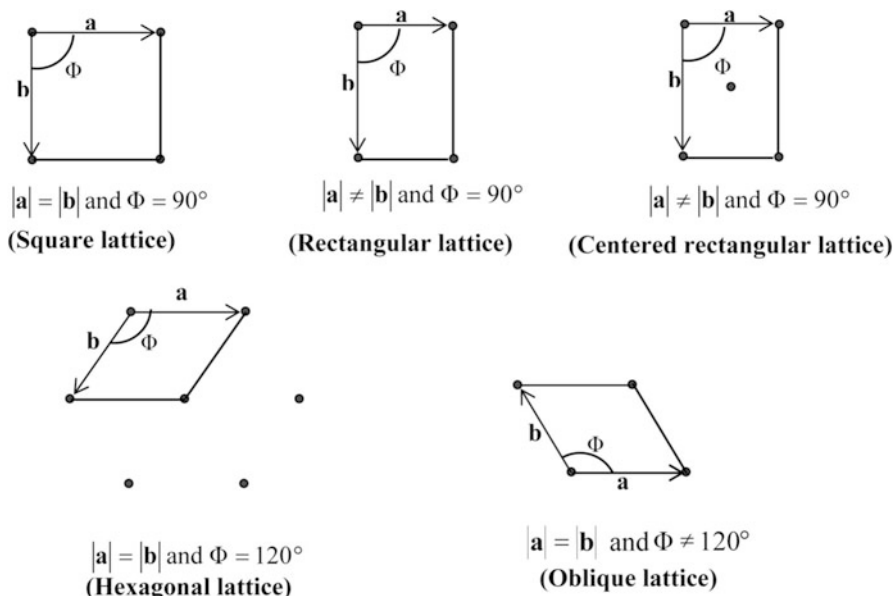


Fig. 2.10 Five two-dimensional Bravais lattices

There are various angles through which one can rotate an object. But it has to be consistent with translational symmetry to make it a crystal. In other words just rotational symmetry is not sufficient to describe a real crystal. It has to be consistent with translational symmetry too. It can be shown that there are only $2\pi/1$, $2\pi/2$, $2\pi/3$, $2\pi/4$ and $2\pi/6$ possible rotations consistent with translation symmetry which can leave the crystal structure invariant.

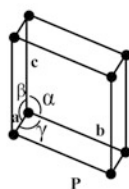
2.3 Two Dimensional Crystal Structures

Although by choosing different vectors, it is possible to construct variety of unit cells, it turns out that there are only *five types* of units or *Bravais lattices* as shown in Fig. 2.10 required to describe all observed two dimensional crystal structures.

2.4 Three Dimensional Crystal Structures

Three dimensional solids are divided into *seven crystal systems* and fourteen Bravais lattices as illustrated in Fig. 2.11. Any three dimensional crystalline solid (except quasicrystals to be discussed later in this chapter) would be one or the other type of the fourteen Bravais lattices.

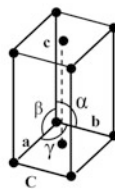
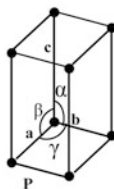
1. Triclinic



$$a \neq b \neq c$$

$$\alpha \neq \beta \neq \gamma$$

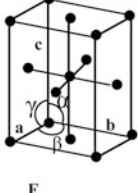
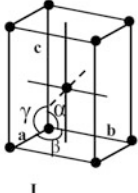
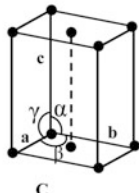
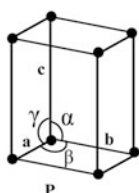
2. Monoclinic



$$a \neq b \neq c$$

$$\alpha = \beta = 90^\circ \neq \gamma$$

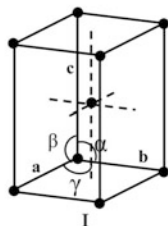
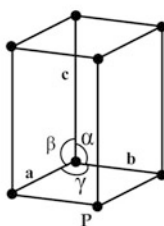
3. Orthorhombic



$$a \neq b \neq c$$

$$\alpha = \beta = \gamma = 90^\circ$$

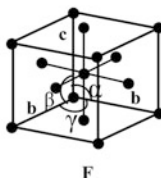
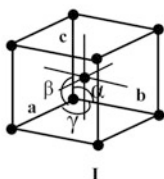
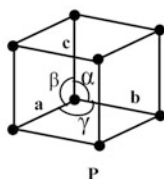
4. Tetragonal



$$a = b \neq c$$

$$\alpha = \beta = \gamma = 90^\circ$$

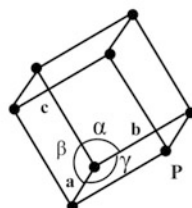
5. Cubic



$$a = b = c$$

$$\alpha = \beta = \gamma = 90^\circ$$

6. Triagonal



$$a = b = c$$

$$\alpha = \beta = \gamma < 120^\circ, \neq 90^\circ$$

Fig. 2.11 Fourteen Bravais lattices in three dimensions (*P* Primitive cell, *I* Body centred cell, *F* Face centred cell, *C* Base centred cell)

7. Hexagonal

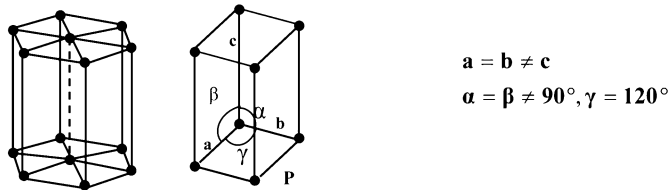


Fig. 2.11 (continued)

2.5 Some Examples of Three Dimensional Crystals

2.5.1 Body Centred Cube (bcc) (Fig. 2.12)

2.6 Planes in the Crystals

Planes in the crystals are denoted by what is known as ‘*Miller indices*’ and are determined as follows.

Consider the intercepts on **a**, **b**, **c** (the primitive axes or non-primitive axes of unit cell). For example as in Fig. 2.13, let the axial lengths be 0.8, 0.4 and 0.4 nm along **a**, **b** and **c** respectively. Let the intercepts be 0.2, 0.4 and 0.4 nm respectively. Fractional intercepts would be $\frac{1}{4}$, 1 and 1. Miller indices of the plane would be 4, 1 and 1 or conventionally written as (411). Miller indices are written in the round brackets.

Few examples are illustrated in Fig. 2.14.

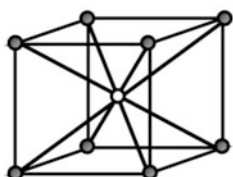
2.7 Crystallographic Directions

Directions are denoted by smallest length of unit vectors i.e. **a** as [100], **b** as [010] and **c** as [001]. Direction numbers are always placed in the square brackets. See Fig. 2.15.

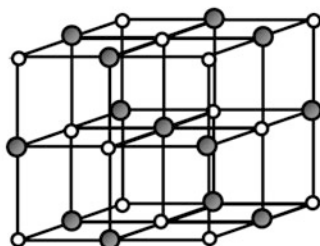
2.8 Reciprocal Lattice

Using X-rays, electrons or neutrons having wavelength ~ 0.1 to 0.2 nm, length comparable to distance between atoms or planes in a solid, it is possible to obtain a diffraction pattern (more details about diffraction pattern will be given in Chap. 7 on Analysis Techniques). Analysis of diffraction patterns leads to the determination

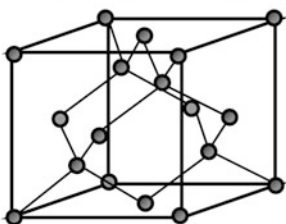
CsCl : Body Centred Cubic (bcc), Cs^+ ion at (000), two atoms/cell, Cl^- ion at $(1/2, 1/2, 1/2)$ co-ordinates. It is a conventional unit cell.



NaCl: Face Centred Cubic (fcc), alternate arrangement of Na^+ and Cl^- ions in x , y and z directions. It is a conventional unit cell.



Diamond: Cubic, interpenetrating cubes, eight atoms/cell



Graphite: Hexagonal hcp ABAB... $c/a = 1.633$

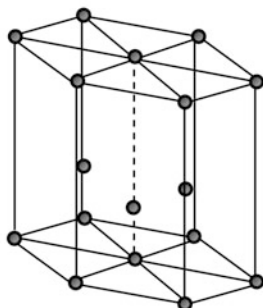


Fig. 2.12 Schematic diagram illustrating CsCl, NaCl, diamond and graphite crystal structures

Fig. 2.13 Plane having Miller indices (411)

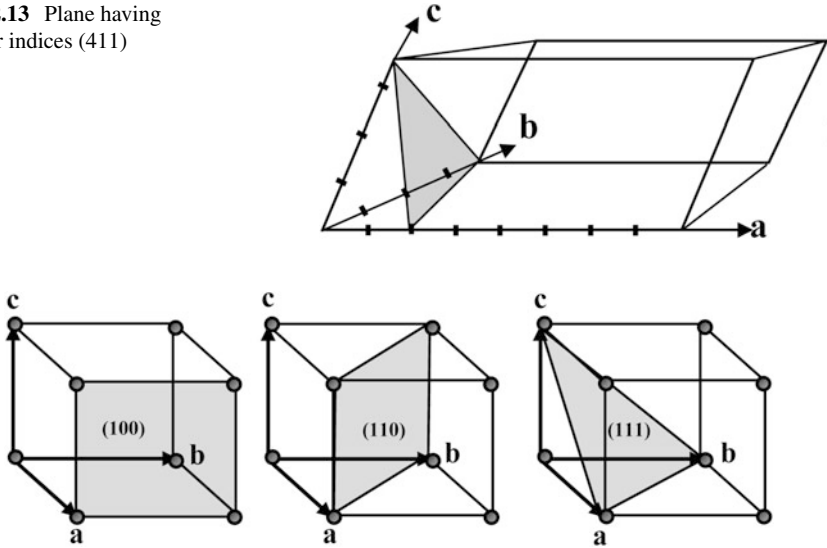
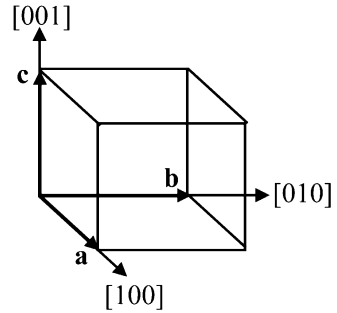


Fig. 2.14 Examples of (100), (110), (111) planes in a cubic crystal

Fig. 2.15 Crystallographic directions



of crystal structure (Bravais lattice). However diffraction pattern is obtained in the *reciprocal space*. Points in the reciprocal space are related to real space. Each point in the reciprocal space represents a set of planes in real lattice.

It is therefore often convenient to use the concept of reciprocal lattice to deal with the crystals. Reciprocal Lattice Vectors **A**, **B** and **C** are defined in terms of the vectors **a**, **b** and **c** in real lattice as

$$\begin{aligned}
 \mathbf{A} &= \frac{2\pi (\mathbf{b} \times \mathbf{c})}{(\mathbf{a} \cdot \mathbf{b} \times \mathbf{c})} \\
 \mathbf{B} &= \frac{2\pi (\mathbf{c} \times \mathbf{a})}{(\mathbf{a} \cdot \mathbf{b} \times \mathbf{c})} \\
 \mathbf{C} &= \frac{2\pi (\mathbf{a} \times \mathbf{b})}{(\mathbf{a} \cdot \mathbf{b} \times \mathbf{c})}
 \end{aligned}
 \tag{2.2}$$

In the resulting network of reciprocal points, position vector $\mathbf{G}$ can be written as

$$\mathbf{G} = h\mathbf{A} + k\mathbf{B} + l\mathbf{C} \quad (2.3)$$

where h, k, l are Miller indices of the planes as discussed in Sect. 2.6.

2.9 Quasi Crystals

As far as point symmetry is concerned, $360/5 = 72^\circ$ is a possible rotation. However if one tries to translate a pentagon in two dimensional space, it cannot fill the space completely as can be seen from the Fig. 2.16. Therefore fivefold rotation does not satisfy the criterion viz. rotation plus translation symmetry to make a crystal. Therefore for quite long time it was thought that crystals having fivefold symmetry cannot exist. Same is the situation with sevenfold symmetry (there are overlapping regions) and you will find that for solids, in the old text books it is considered to be a forbidden geometry for atomic arrangement in crystals. In fact as discussed earlier, only 1, 2, 3, 4 and 6 fold axes of rotation are crystallographically allowed.

However there are now evidences after the experimental work by Dan Shechtman in 1982 that in some crystalline matter fivefold symmetry does exist. Indeed in some alloys like Al-Mn, Al-Pd-Mn and Al-Cu-Co, fivefold, eightfold or tenfold symmetries have been observed. Dan Shechtman got the Nobel Prize for this discovery in 2011. This kind of aperiodic symmetry can be explained by a different kind of space filling method than using simple geometries like rectangles or squares. The credit for it goes to astrophysicist Penrose and is known as '*penrose tiling*' (Interestingly penrose tiling was made by Roger Penrose way back in 1970s just to show an aperiodic tiling. Before Penrose's suggestions such tiling could be seen in Islamic culture). What Penrose suggested is that instead of using a single unit cell for the translation, use two-unit cells like a thin and a fat rhombus or a kite and a dart as shown in Fig. 2.17.

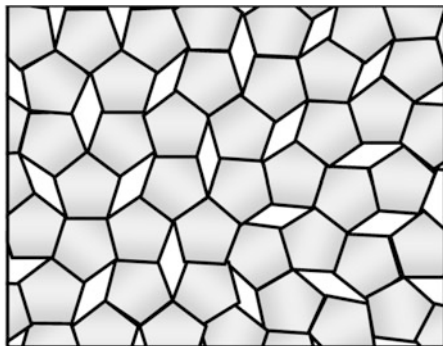


Fig. 2.16 Five-fold pattern

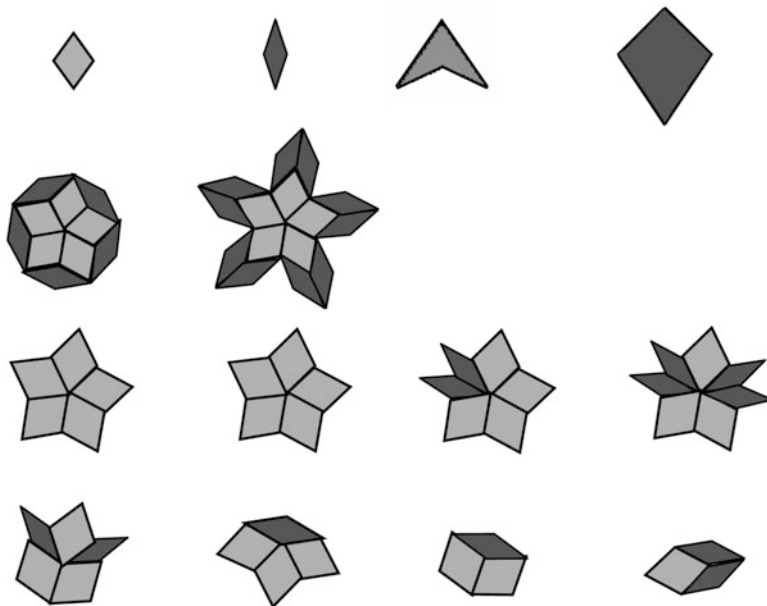


Fig. 2.17 Penrose tiles (*first line*) and various motifs generated using them

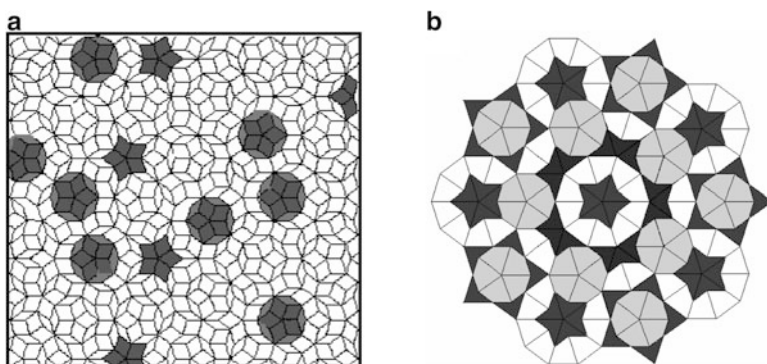


Fig. 2.18 Penrose tiling: (a) using fat and thin rhombus and (b) using kite and dart

Now we can see (Fig. 2.18) how they fit into each other filling a large region without leaving spaces in between or overlapping with each other. Such aperiodic materials which exhibit X-ray diffraction patterns are known as *quasi crystals*. Properties of such alloys are quite interesting. For example they are poor heat and electricity conductors. They also do not get rusted easily. Interestingly Shechtman was not aware of Penrose tiling for a long time.

Box 2.1: Revised Definition of a Crystal

Crystals are defined as solids with periodic arrangement of atoms or molecules. Until 1982 all the known crystals were found to follow the translational and rotational symmetry operations in which only 1, 2, 3, 4 and 6 fold rotations are allowed. However after the observation of quasi crystals by Shechtman it was apparent that even 5, 7 and 10 fold rotations are also possible in solids. Therefore, International Organisation of Crystallographers revised the definition of crystals in 1991. A crystal can be defined now as any solid having a distinct diffraction pattern. We shall learn about the diffraction from crystals in a later chapter.

2.10 Liquid Crystals

Similar to quasicrystals, liquid crystal is another class of unconventional crystals. Unlike quasicrystals the work on liquid crystals began as early as 1888 when Friedrich Reinitzer noticed that the cholesteryl benzoate, when heated, became a cloudy liquid at 145 °C and clear liquid at 179 °C. After hearing about this curious behaviour of cholesteryl benzoate molecules, Otto Lehman looked into more details and not only confirmed Reinitzer's results but also gave the name *liquid crystals* in 1889 to it and similarly behaving materials. He found that some molecules before melting pass through a (liquid) phase in which they flow like liquids. However they have interesting optical properties and have some attributes of crystals. As the name suggests it has liquid-like as well as crystal-like properties. Although credit on the discovery of liquid crystals is given to Reinitzer, he himself has quoted similar observations by Julius Planer from 1861.

Many molecules have rod-like shapes and are the candidates of different types of liquid crystals like nematic, cholesteric, smectic and discotic (see Fig. 2.19). In

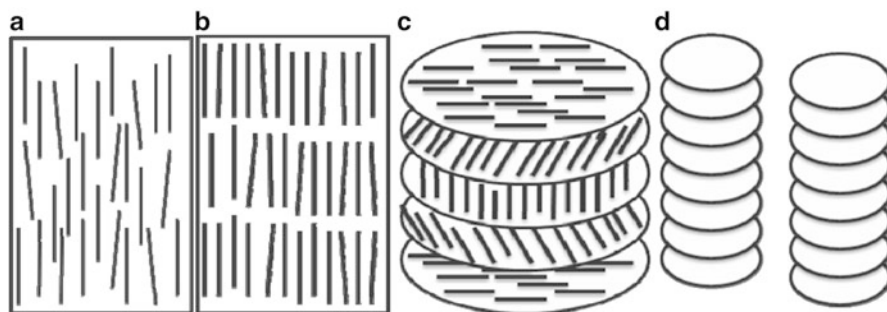


Fig. 2.19 (a) Nematic (b) Smectic (c) Cholesteric (d) Discotic liquid crystals

nematic crystals the molecules are almost parallel to each other in the vertical direction but their positions are random. Thus there is orientation order like in a crystal but positional disorder like in a liquid. The smectic crystals are characterized by slabs of oriented molecules which are disordered with respect to each other. In the cholesteric (also called chiral nematic) the molecules are oriented in the same direction as in a nematic crystal in a single plane but there is rotation of molecules from plane to plane like twisting. There are some disc-like molecules which arrange themselves randomly in a plane but form vertical stacks to form what is known as discotic liquid crystals.

Although for decades the liquid crystal research was driven only by scientific curiosity, their optical properties showed that they have a huge potential in display panels. They are widely used in the display screens of televisions, computers, watches and many other applications.

2.11 Bonding in Solids

When two atoms approach each other, they start interacting. The type and the amount of interaction depend upon the distance between the two atoms and the type (element or electronic configuration) of the atom. However, in general, one can understand this interaction as follows with the help of Fig. 2.20. Consider two atoms. At infinity or a very large distance compared to their sizes, the electrons have no interatomic interaction of any type and both the atoms have all the properties independent of each other. However, as soon as they start coming closer, their electron clouds start interacting with each other. This results into an attractive force and a repulsive force between the two atoms due to mutual electrostatic interaction. The electrostatic interaction between the atoms arises due to the positive charges of nuclei and negative charges of electron clouds. All other forces like gravitation have negligible effect compared to electrostatic interaction and can be neglected.

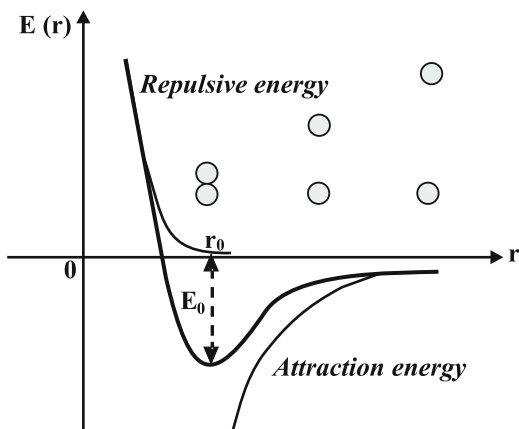


Fig. 2.20 Interatomic potential vs. interatomic distance. Here E_0 is the bond energy and r_0 is the bond length

A distance dependent energy curve has two parts viz. an attractive (negative energy) part and the repulsive (positive energy) part. At certain distance between the two atoms, the attractive interaction dominates as compared to the repulsive interaction. At this distance, say r_0 , the two atoms can form what is known as 'bond' with 'bond energy' E_0 . The two atoms enjoy the state of lower energy by being together than remaining as separate ones. This can give rise to the formation of a molecule. As we shall see below, there can be different types of bonds with different bonding strengths or value of E_0 . The equilibrium distance at which the energy is E_0 is the 'bond length' of the molecule. As the atoms come closer or have a distance smaller than r_0 between them, their electron charge clouds start overlapping. Due to similar charges they start repelling each other. This can be understood in terms of Pauli exclusion principle, which states that no two electrons can have same quantum numbers. Even in solids the atoms are held together due to mutual interaction between the atoms and we have bond energy and bond length as in a molecule. However in a solid there are not just two atoms as we discussed above but a very large number of atoms (even many molecules have more than two atoms and different bonds within a single molecule). It is still possible to consider the interaction between the atoms using a simple diagram as discussed above. The energy difference between energy of free atoms and that of crystal is known as *cohesive energy*.

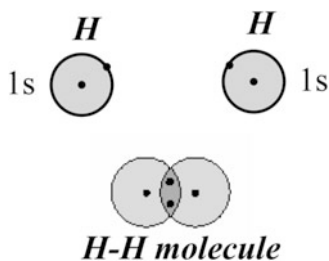
Let us now discuss various types of bonds that occur in some molecules and solids. They are classified into different types of bonds as: (1) Covalent bond, (2) Ionic bond, (3) Metallic bond, (4) Mixed bond and (5) Secondary bond.

Covalent, ionic and metallic bonds are known as primary bonds and others are known as secondary bonds. There are also bonds which are not pure ionic, covalent or metallic but have some partial character of different bonds and are known as mixed bonds. We shall briefly discuss each type of bond with some examples.

2.11.1 Covalent Bond

When two atoms form a molecule by sharing some of their valence electrons, a covalent bond is established. Sometimes atoms may use even all the valence electrons for sharing. Consider a simple example of H_2 molecule. Hydrogen atomic number is 1 and it has its one electron in 1s shell. We know that in atoms stable orbits can be formed with certain electronic configurations, like two electrons in 1s shell. Thus hydrogen atom is more stable if there are two electrons in 1s shell. When two hydrogen atoms come very close to each other so that their 1s shells overlap (See Fig. 2.21), then both have as if two electrons. Therefore they form a molecule attaining a state of lower energy together than both were independent. There are many other molecules like CH_4 , O_2 or N_2 which are covalent in nature.

In solids we have covalent bonds in many crystals like diamond, silicon, germanium etc. We shall discuss the covalent bond in solids with the example of diamond crystal. Diamond is made up of carbon atoms. Carbon has atomic number

Fig. 2.21 Hydrogen molecule**Table 2.1** Examples of some covalent bonds, their equilibrium bond lengths and energies

Bond	Bond length (nm)	Bond energy (eV)
H–H	0.07414	4.5
O–O	0.12074	1.4
P–P	0.18931	2.2
Cl–Cl	0.19879	2.5
C–C	0.12425	3.6
Si–Si	0.2246	1.8
Ge–Ge	0.2403	1.6

6 and its electron configuration is $1s^2 2s^2 2p^2$. Thus, it lacks four atoms in p shell to attain the stable or lowest possible energy configuration. In crystal each carbon atom is coordinated (or surrounded by nearest neighbours as shown in Fig. 2.12) by four carbon atoms. Instead of just using two valence electrons from 2p shell, all four electrons from principal quantum number 2 or L shell are used by each carbon atom. This enables each carbon atom to have four bonds (known as sp^3 bonds) which are then shared with four neighbouring carbon atoms. In this way each carbon atom shares one electron with every other neighbouring carbon atom forming covalent bonds. This also gives rise to a tetrahedral arrangement of carbon atoms so that angle between C–C–C atoms is equal.

In Table 2.1 we give some more examples of covalent bonds and their bond energies along with bond lengths.

2.11.2 Ionic Bond

This type of chemical bond is formed when atoms are in the vicinity and electron(s) from one atom is transferred to another. The one able to transfer the electron or electrons is termed as electropositive and the one which has a tendency to accept electron or electrons is known as electronegative. The ions thus formed are called cations and anions respectively. NaCl, KCl and CaCl_2 are good and simple examples of ionic bond formation. Let us consider the example of NaCl. Sodium (Na) has atomic number 11 and electronic configuration $1s^2 2s^2 2p^6 3s^1$. We can easily see that if there was not an electron in 3s orbit of sodium, it would have been stable like

inert gas element neon. On the other hand chlorine with atomic number 17 (electron configuration $1s^2 2s^2 2p^6 3s^2 3p^5$) lacks one electron to attain a stable configuration like that in argon (Ar) with 18 electrons. Thus we have a situation here that one atom has extra electron and another is short of an electron to attain the stable configuration. When two such atoms are at a sufficiently short distance the electron from Na gets transferred to Cl forming an ionic bond and both gaining a state of lower energy. The bond thus formed is known as *ionic bond*. However there is energy required to remove an electron from metal atom like sodium atom in this example. This energy is known as *energy of ionization*.

In solid of NaCl, as illustrated in Fig. 2.12, the crystal structure is face centre cubic (FCC) with each ion of one type surrounded by six ions of opposite charge. Also each ion has oppositely charged neighbours. Note that there is a difference in the sizes of Na and Cl ions.

However there is ionization energy required to remove an electron from sodium atom. Ionic crystals are easily soluble in polar liquids like water. The ionic crystals have high melting point. In general they are electrically and thermally non-conducting.

2.11.3 Metallic Bond

Atoms like Na, K, Cu, Au, Ag, Fe, Co and Ca in which there are one or more electrons which can be easily removed (valence electrons) from them form metals and the bond that holds the crystal together is known as *metallic bond*. The loose electrons are able to move quite easily from one atom to another and cannot be localized to a particular atom or the other. Such electrons form what is known as an *electron gas* (of free electrons) and are responsible for the *metallic bond* which is an electrostatic interaction between the positive ions of atoms that have lost the electrons and are unable to move themselves (however, ions vibrate about their mean positions) in the crystal. This is schematically shown in Fig. 2.22. In metals the bond is not directional and usually the metal ions try to come as close as possible forming a compact structure with high coordination number of ions. It is therefore possible to move ions in any direction quite easily by applying a force. The metals are thus ductile. The free electrons in metals respond to electric or magnetic field quite easily making them electrically or thermally good conducting materials.

2.11.4 Mixed Bonds

It is often found that bonds have a mixed character. We saw in case of diamond or silicon crystals, we have pure covalent bonds. However if we consider a crystal like that of NaCl, KCl, GaAs, InP or ZnS we find that there is no pure covalent or pure ionic bond amongst the atoms. A class of materials (Al_2O_3 , MgO etc.)

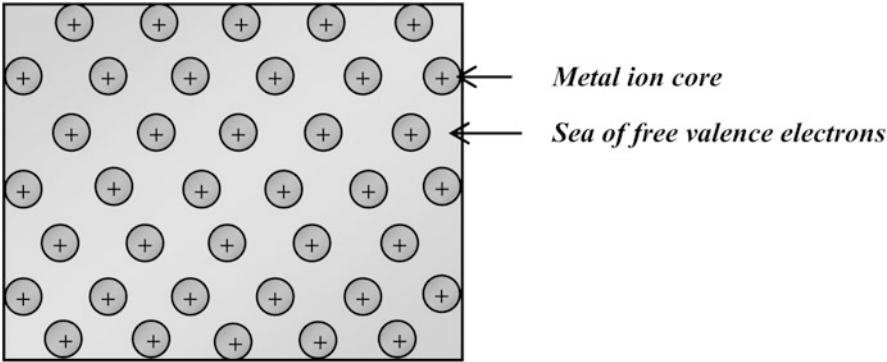


Fig. 2.22 A metal considered as a periodic arrangement of positive ions (*circles*) embedded in an electron gas (*grey background*) of negatively charged electrons contributed by each atom

Table 2.2 Examples of covalent and mixed bonds

Materials	Ionic %	Covalent %
Si	0	100
Ge	0	100
ZnO	62	38
CdS	69	31
GaAs	31	69

known as ‘ceramics’ also have partial metallic and non-metallic characters (ionic and covalent mix bonding). Their properties also change correspondingly. They may not be completely insulators nor conductors of heat and electricity. Such crystals are known to have mixed type of bonding. Often it is useful to find out the percentage of each type (See Table 2.2).

2.11.5 Secondary Bonds

These are very weak bonds compared to ionic, covalent and metallic bonds. Typical examples are found in rare gas atom solids (obtained at low temperatures), solidified CO₂, solidified C₂H₄, water and some molecular crystals. These weak bonds are due to formation of instantaneous dipoles in rare gas atoms and molecules which leads to what is known as *Van der Waals bond*. As illustrated in Fig. 2.23, although rare gas atoms have stable electronic configurations and cannot interact with other atoms when they are very close to each other, at any instance they have small dipoles with positive and negative ends due to movement of charges inside. At high temperature the kinetic energy of atoms is very large and atoms move with high velocities without any interactions between the diploes. However at low temperatures and small distances the dipoles can interact. The interaction between the dipoles can

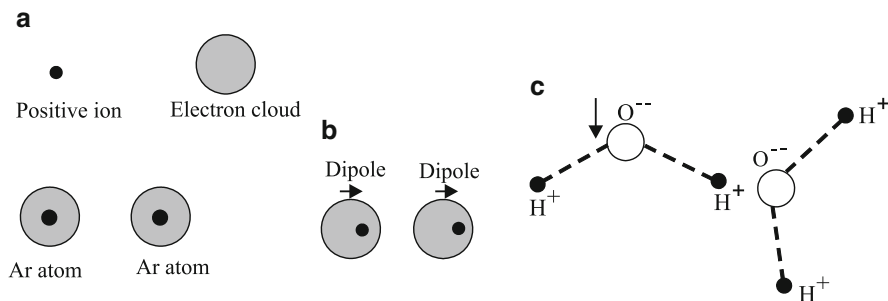


Fig. 2.23 Centre of positive charge and negative charges, over a long period (average), coincides as in (a) but at any instance a dipole can exist due to displacement of electric charge cloud as in (b). Dipole-dipole interaction (c)

be attractive as well as repulsive. But there can be a weak net attractive interaction. This gives a small value of E_0 or bond energy.

It is also likely that some molecules (see Fig. 2.23c) have already some permanent dipoles on them which can then interact with each other. Molecules with permanent dipoles on them are known as *polar molecules*. Consider for example a water molecule H_2O . In water molecule, oxygen attracts two electrons from two hydrogen atoms. Therefore within the molecule we have ionic bonds. Oxygen is negatively charged and two hydrogen atoms are positively charged. This gives rise to a permanent dipole on water molecule (Each oxygen ion can attract positively charged hydrogen atom and each hydrogen atom can in principle get attracted to other H_2O molecule). Such dipoles then can interact with each other forming weak van der Waals bonds but are referred to as *hydrogen bonds*. These are slightly stronger than the fluctuating dipoles discussed earlier but still lie in the category of secondary bonds. In some cases an atom of hydrogen is responsible to bond two electronegative atoms simultaneously. For example two fluorine (chlorine etc.) atoms are simultaneously bonded with single hydrogen atom. Hydrogen atom loses its electron to one of the fluorine atom (ionic bond) and the proton is attracted and is positively charged which bond with the second fluorine atom. Thus a difluoride HF_2^- is formed by *hydrogen bond*.

Van der Waals bonding is also observed in long chains of carbon or C–H. The chains themselves may have various atoms bonded by covalent and ionic bonds but the inter chain reactions may be weak due to Van der Waals interaction, as in polymers.

2.12 Electronic Structure of Solids

We discussed in the previous sections about geometrical structures and bonding in solids. Materials can be identified with appropriate structure and bonding. However to further understand different mechanical, thermal, electrical, magnetic

and optical properties it is necessary to further know the *electronic structure* of solids. The electronic structure of nanomaterials would certainly be different from atoms, molecules or solids. Nanomaterials in fact are intermediate form of solids. They are too small to be considered as bulk, three dimensional solids and too big to be considered as molecules. If we start with an atom, we know that bringing two atoms together to form a molecule changes the electronic structure. If three atoms come together there is further change. In fact addition of each atom would change the electronic structure till the number of atoms becomes too large. Depending upon the type of atoms, at large number the clusters of such atoms would reach bulk electronic structure. Bulk electronic structures or those of atoms and molecules can be learnt in standard text books on these topics. Here we shall briefly discuss how such electronic structures are useful to classify the materials as metals, semiconductors and insulators.

As a simple picture one can consider a one dimensional solid in which atoms are periodically arranged on a line (one dimensional case which can be readily extrapolated to take into account a three dimensional solid). Each atom can then be considered to contribute one electron (becoming a positive ion) which is free to move in the solid.

We can consider first free electron motion and see how such concept leads to band structure of solids.

2.12.1 Free Electron Motion

We saw in Chap. 1 that Schrödinger equation for an electron in one dimension is given as

$$-\frac{\hbar^2}{2m} \left(\frac{\partial^2}{\partial x^2} \right) \psi(x) + V(x)\psi(x) = E\psi(x) \quad (2.4)$$

and the wave function $\psi(x)$ for a free particle is

$$\psi(x) = A \sin \left(\frac{n_x \pi x}{a} \right) \quad (2.5)$$

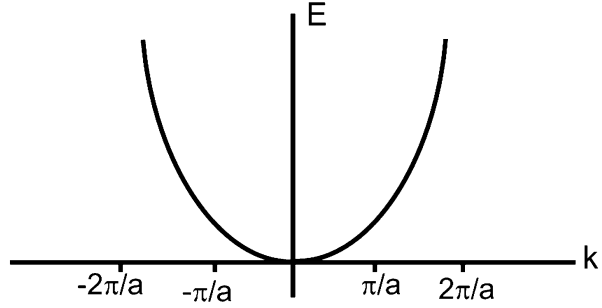
with $n_x = 0, 1, 2 \dots$

2.12.2 Bloch's Theorem

From the free electron motion we go to the electron states in a one dimensional periodic potential, which can be written as

$$V(x) = V(x + a) \quad (2.6)$$

Fig. 2.24 Electron energy levels at different wave vector k



It was realized by Bloch that the wave function satisfies the relation

$$\psi(x + a) = \psi(x)e^{ika} \quad (2.7)$$

2.12.3 Origin of Band Structure

We now give an intuitive argument for the presence of bands. We know that

$$E_k = \frac{\hbar^2 k^2}{2m} \quad (2.8)$$

for travelling wave in one direction.

The electron spectrum is as shown in Fig. 2.24.

$\psi(+)=e^{\frac{i\pi x}{a}}$ and $\psi(-)=e^{\frac{-i\pi x}{a}}$ are oppositely moving waves.

For a pure travelling wave

$$|\psi| = e^{ikx}e^{-ikx} = 1 \quad (2.9)$$

But if electron waves are reflected at $k = \pm\pi n/a$, the standing waves are generated as

$$\psi(+)=e^{\frac{i\pi x}{a}}+e^{\frac{-i\pi x}{a}}=2\cos\left(\frac{\pi x}{a}\right) \quad (2.10)$$

$$\text{and } \psi(-)=e^{\frac{i\pi x}{a}}-e^{\frac{-i\pi x}{a}}=2\sin\left(\frac{\pi x}{a}\right) \quad (2.11)$$

$\psi(+)$ and $\psi(-)$ pile up electrons in different regions as shown in Fig. 2.25.

$$\text{Probability density } \rho(+)=|\psi(+)|^2=4\cos^2\frac{\pi x}{a} \quad (2.12)$$

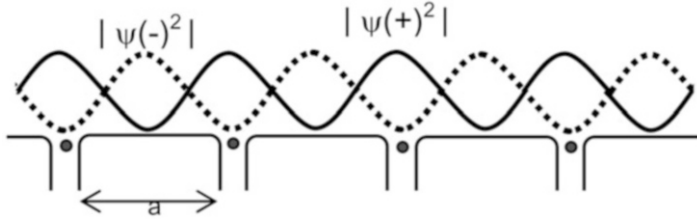
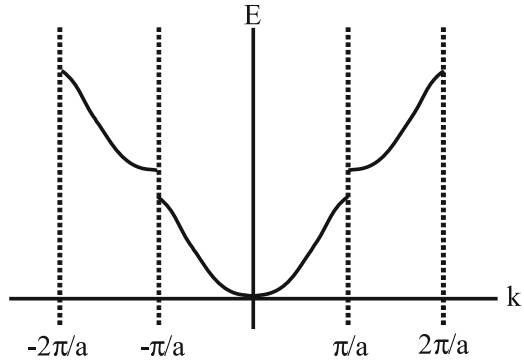


Fig. 2.25 Electron probability distribution in periodic potential

Fig. 2.26 Energy levels separated by band gap at zone boundary



$$\text{and } \rho(-) = |\psi(-)|^2 = 4\sin^2 \frac{\pi x}{a} \quad (2.13)$$

This gives rise to energy gaps in the free electron curve of Fig. 2.24 as illustrated in Fig. 2.26.

This is a very important result as details about amount of energy gap, overlap of energy bands, curvature of bands determine the properties of solids. The bands in solids have closely spaced energy levels. They are filled upto Fermi level according to Pauli exclusion principle viz. no two electrons can exist in an energy state with all the quantum numbers same. The Fermi level is the highest occupied level at absolute zero temperature.

A qualitative picture representing metal, insulator and semiconductors is given in Fig. 2.27. A metal would have its outermost electron filled band (valence band) overlapping with empty band (conduction band). If there is a small gap between valence and conduction band (upto $\sim 2-3$ eV) then it is a semiconductor. If the gap is larger, the material is an insulator. Interestingly, by introducing small amount of dopant (atoms different compared to those in the host material), some localized energy states can be introduced in the energy gap. Depending upon whether such levels are close to valence band or close to conduction band, semiconductor becomes *p* or *n* type semiconductor. The electrical conductivity in many semiconductor devices is controlled by such dopants.

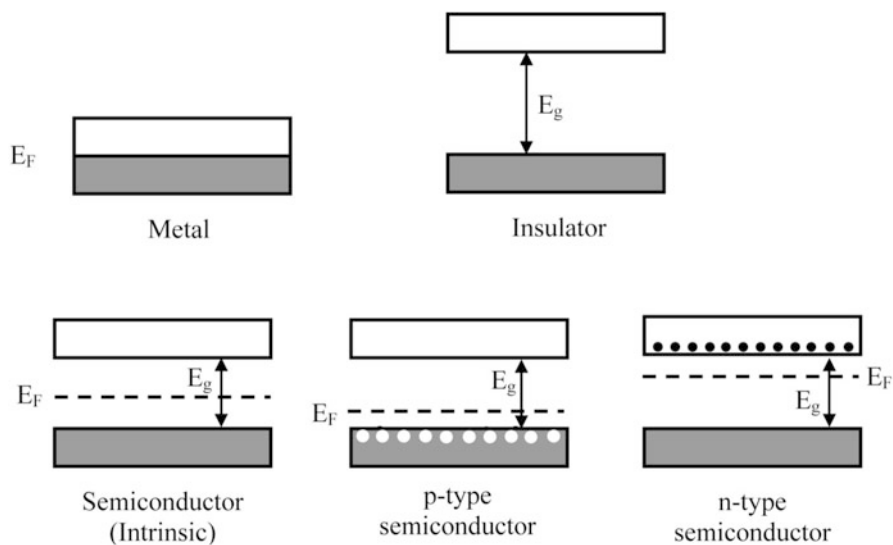


Fig. 2.27 Different types of materials viz. metals, semiconductors and insulators can be understood on the basis of the overlap or separation between outer filled valence and empty conduction bands

A band picture of solids thus enables one to understand differences between metals, semiconductors and insulators. Moreover detailed analysis also helps to understand other physico-chemical properties of solids.

Further Reading

- M. Chandra, *Atomic Structure and Chemical Bond*, 3rd edn. (Tata McGraw Hill Publishing Co Ltd, New Delhi, 1991)
 C. Kittel, *Introduction to Solid State Physics*, 5th edn. (Wiley, New Delhi, 1995)

Chapter 3

Synthesis of Nanomaterials—I

(Physical Methods)

3.1 Introduction

There are a large number of techniques available to synthesize different types of nanomaterials in the form of colloids, clusters, powders, tubes, rods, wires, thin films etc. Some of the already existing conventional techniques to synthesize different types of materials are optimized to get novel nanomaterials and some new techniques are developed. Nanotechnology is an interdisciplinary subject. There are therefore various physical, chemical, biological and hybrid techniques available to synthesize nanomaterials. It can be seen from Box 3.1 that, for each type, there is a large number of possibilities. The list is not complete but gives some commonly used techniques; some of them will be described in this and in the next two chapters. The technique to be used depends upon the material of interest, type of nanomaterial viz. zero dimensional (0-D), one dimensional (1-D) or two dimensional (2-D), their sizes and quantity.

In this chapter we shall discuss some physical as well as related methods to obtain nanomaterials.

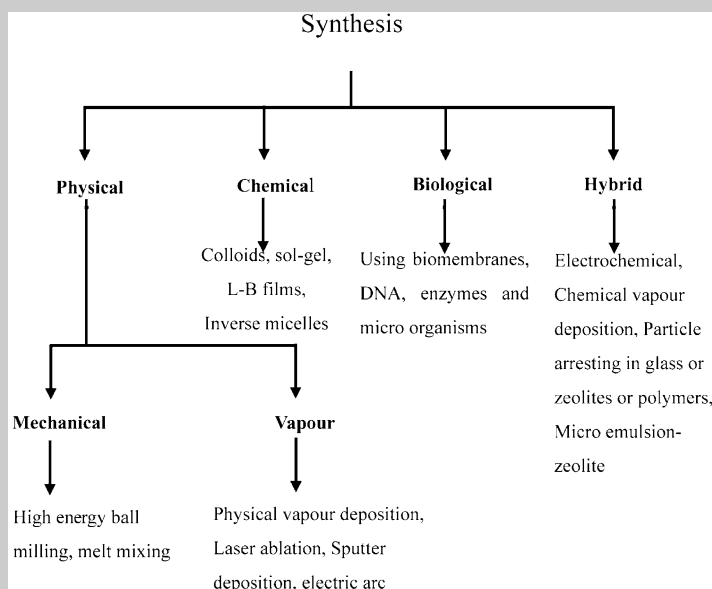
3.2 Mechanical Methods

3.2.1 *High Energy Ball Milling*

It is one of the simplest ways of making nanoparticles of some metals and alloys in the form of powder. There are many types of mills such as planetary, vibratory, rod, tumbler etc. Usually one or more containers are used at a time to make large quantities of fine particles. Size of container, of course, depends upon the quantity of interest. Hardened steel or tungsten carbide balls are put in containers (see Fig. 3.1) along with powder or flakes ($<50\text{ }\mu\text{m}$) of a material of interest. Initial material can

be of arbitrary size and shape. Container is closed with tight lids. Usually 2:1 mass ratio of balls to material is advisable. If the container is more than half filled, the efficiency of milling is reduced. Heavy milling balls increase the impact energy on collision. Larger balls used for milling produce smaller grain size but larger defects in the particles. The process, however, may add some impurities from balls. The container may be filled with air or inert gas. However this can be an additional source of impurity, if proper precaution to use high purity gases is not taken. A temperature rise in the range of 100–1,100 °C is expected to take place during the collisions. Lower temperatures favour amorphous particle formation. The gases like O₂, N₂ etc. can be the source of impurities as constantly new, active surfaces are generated. Cryo-cooling is used sometimes to dissipate the heat generated. During the milling, liquids also can be used. The containers are rotated at high speed (a few hundreds of rpm) around their own axis. Additionally they may rotate around some central axis and are therefore called as ‘planetary ball mill’.

Box 3.1



When the containers are rotating around the central axis as well as their own axis, the material is forced to the walls and is pressed against the walls as illustrated in Fig. 3.2. By controlling the speed of rotation of the central axis and container as well as duration of milling it is possible to ground the material to fine powder (few nm to few tens of nm) whose size can be quite uniform.

Fig. 3.1 A schematic diagram (*sectional*) of a ball mill vessel

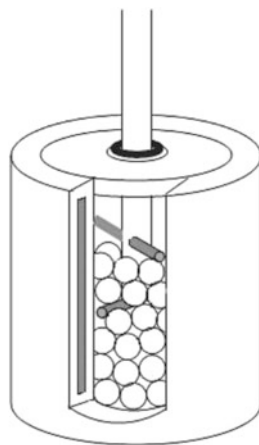
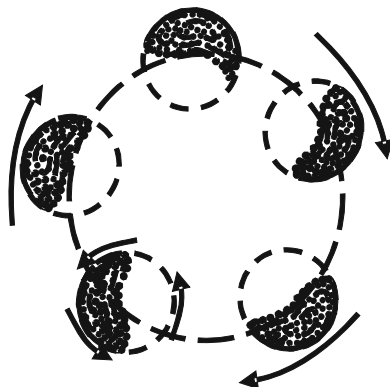


Fig. 3.2 Ball mill in planetary motion (*schematic*). Sketch shows that the material is thrown against the wall during the course of rotation of a single container. *Dark regions* are illustrating the powder material, while the rest is empty

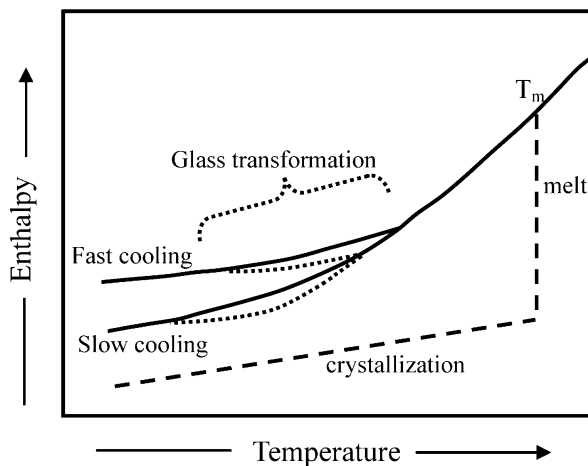


Some of the materials like Co, Cr, W, Ni-Ti, Al-Fe and Ag-Fe are made nanocrystalline using ball mill. Few milligrams to several kilograms of nanoparticles can be synthesized in a short time of a few minutes to a few hours.

3.2.2 Melt Mixing

It is possible to form or arrest the nanoparticles in glass. Structurally, glass is an amorphous solid, lacking large range periodic arrangement as well as symmetry arrangement of atoms/molecules. When a liquid is cooled below certain temperature (T_m), as illustrated in Fig. 3.3, it forms either a crystalline or amorphous solid (glass). Besides temperature, rate of cooling and tendency to nucleate decide whether the melt can be cooled as a glass or crystalline solid with long range order. Nuclei are formed spontaneously with homogeneous (in the melt) or inhomogeneous (on surface of other materials) nucleation which can grow to form ordered, crystalline solid. Metals usually form crystalline solids but if cooled at very high rate

Fig. 3.3 Cooling pattern of glass forming melt



($\sim 10^5$ – 10^6 K/s), they can form amorphous solids. Such solids are known as metallic glasses. Even in such cases the atoms try to reorganize themselves into crystalline solids. Addition of elements like B, P, Si etc. helps to keep the metallic glasses in amorphous state. It is indeed possible to form nanocrystals within metallic glasses by controlled heating (Box 3.2).

It is also possible to form some nanoparticles by mixing the molten streams of metals at high velocity with turbulence. On mixing thoroughly, nanoparticles are found to have been formed. For example a molten stream of Cu-B and molten stream of Ti form nanoparticles of TiB_2 .

Box 3.2: Glass Matrix for Nanoparticle Formation

Although formation of nuclei is necessary for the onset of crystallization, presence of large number of nuclei of different elements creates an inhomogeneous melt favouring glass formation. There are some materials that particularly favour glass formation. Silicates (or germanates) have a central silicon atom in a pyramidal structure with oxygen atoms at the corners as shown in Fig. 3.4. Such silicate units, not sharing the edges, but connected only through the corners are randomly distributed in 3-D to form glassy structure. A window glass has in addition to SiO_2 ($\sim 72\%$), oxides of sodium ($\sim 14.5\%$), calcium ($\sim 8.5\%$), magnesium ($\sim 3.5\%$) and aluminium ($\sim 1.5\%$) as its constituents. Laboratory glassware has silica (80%), boron oxide (10%), sodium oxide (5%), alumina (3%), magnesium oxide (1%) and calcium oxide (1%). Other glasses like fibre glass and quartz glass have different compositions. It can be however noticed that glass formation is largely assisted by the presence of number of different elements

(continued)

Box 3.2 (continued)

(heterogeneous nature). It has been found that beautiful colours in glasses like red, yellow, blue, orange, green and their shades are a result of addition of some elements like gold, cobalt, nickel, iron etc. These colours are attributed to the formation of nanoparticles of these elements. Indeed it is also possible to form nanoparticles of some semiconductors like CdS, CdSe, ZnS and TiO_2 in glass. Usually the percentage of such additives is as small as $<5\%$. The melt of glass forming materials and desired nanoparticle material is well homogenized before the cooling begins. The nanoparticles are usually well separated and immobilized by glassy matrix.

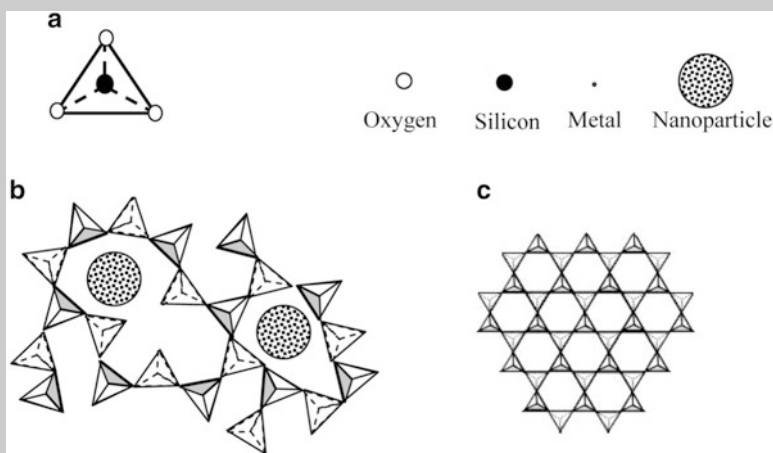


Fig. 3.4 Structure of silicates

- (a) Tetrahedral (pyramidal) building block of typical silicates or germanets
- (b) Disordered arrangement of tetrahedral units to form a 'glassy' substance. This entraps nanoparticles inside appropriate voids
- (c) Tetrahedron can also form crystalline regular arrangement

3.3 Methods Based on Evaporation

There is a variety of methods to form nanostructures by evaporating the materials on some substrates. The nanostructured materials can be in the form of thin films, multilayer films or nanoparticulate thin films (thin films composed of nanoparticles). There are several methods in which material of interest is brought in the gaseous phase (atoms or molecules) which can form clusters and then deposit on appropriate substrates. It is also possible to obtain very thin (even atomic layers, known as monolayers) layers or multilayers (multilayers are layers of two or more materials stacked over each other) forming nanomaterials of wide interest.

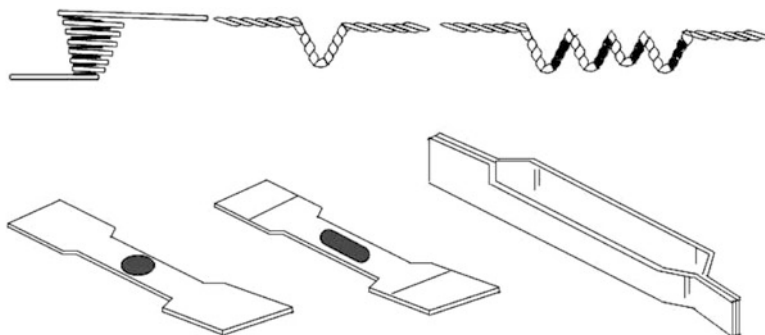


Fig. 3.5 Different shapes of filaments, canoe and baskets used for holding the materials for evaporation

Evaporation can be achieved by various methods like resistive heating, electron beam heating, laser heating, sputtering. It should be remembered that all the synthesis processes need to be carried out in a properly designed vacuum system (see Appendix IV for vacuum techniques), so as to avoid uncontrolled oxidation of source materials and final product as well as that of components of the synthesis system. Mean free path of the particles also increases in vacuum system, which is often desired. Even if some reactive gases are used in certain cases of depositions, it is useful first to evacuate the system to very low pressure so that the materials to be synthesized do not get contaminated by undesired atoms and then pressurize the system to desired value by introducing the high purity gases in the synthesis chamber.

Materials to be evaporated are usually heated from some suitable filament, crucible, boat (collectively called as 'evaporation source' or 'crucible') as shown in Fig. 3.5.

Usually the sources are electrically heated so that enough vapours of the material to be deposited are generated. If the material to be deposited wets the filament material without forming any alloy or compound, the filament is considered to be suitable. Otherwise one needs to melt the material in a basket, canoe-like container. However this type of heating has the disadvantage that the crucible itself and surrounding parts also get heated and become the source of unwanted contamination or impurities. Therefore evaporation by electron beam heating method is desired. Electron beam focuses on the material to be deposited, kept in the crucible as it is generated from a filament that is not in the proximity of the evaporating material. It melts only some central portion of the material in crucible avoiding any contamination from crucible. Thus high purity vapours of materials can be obtained.

Now, let us discuss how the deposition occurs by evaporation. At any given temperature there is some vapour pressure of the material. In evaporation, the number of atoms leaving the surface of solid or liquid material should exceed the

atoms returning to the surface. Evaporation rate from a liquid is given by Hertz-Knudsen equation as follows

$$\frac{dN}{dt} = A\alpha(2\pi mkT)^{-1/2} (p^* - p) \quad (3.1)$$

where N is the number of molecules or atoms leaving a liquid or solid surface, A – area from which atoms evaporate, α – coefficient of evaporation, m – mass of evaporating atoms, k – Boltzmann constant, T – temperature, p^* – equilibrium pressure at the source of evaporation and p is hydrostatic pressure on the surface.

If none of the evaporated molecules or atoms return to the surface of the evaporation source, $p = 0$. The coefficient of evaporation arises as some of the atoms/molecules returning to the surface get reflected back into vapour phase and also change the pressure due to evaporant. Equation (3.1) also suggests that at a given temperature, there would be some specific rate of deposition. It may be noted that evaporation rate equations would in general be more complicated than given by a simple Hertz-Knudsen equation when evaporation from solids, compounds, alloys etc. are taken into account.

It is necessary that the material to be evaporated creates a pressure of $\sim 10^{-1}$ Pa or more, so as to achieve adequate vapour pressure for synthesis. Some materials like Ti, Mo, Fe and Si have large vapour pressure at a temperature, much below their melting points and can be easily evaporated or sublimated from their solids. On the other hand metals like Au and Ag have very low vapour pressures even close to their melting points. Therefore, they need to be melted (for evaporation to occur) to achieve adequate vapour pressure required for deposition.

The constituents of the alloys evaporate at different rates depending upon their pure metal forms. Therefore, the deposited films may have different stoichiometry (% amounts of different constituents in a material) as compared to the original alloy composition.

If a compound is used for evaporation, it is very likely that it would dissociate beforehand and stoichiometry is not maintained in the final product.

Use of reactive gases like oxygen, ammonia and hydrogen during the material synthesis by evaporation is a common practice. Sometimes post-deposition treatments are also used.

In fact in a situation where the material to be deposited is not a single component material, it is preferable to use sputter deposition.

3.3.1 Physical Vapour Deposition with Consolidation

This technique basically involves use of materials of interest as sources of evaporation, an inert gas or reactive gas for collisions with material vapour, a cold finger on which clusters or nanoparticles can condense, a scraper to scrape the nanoparticles

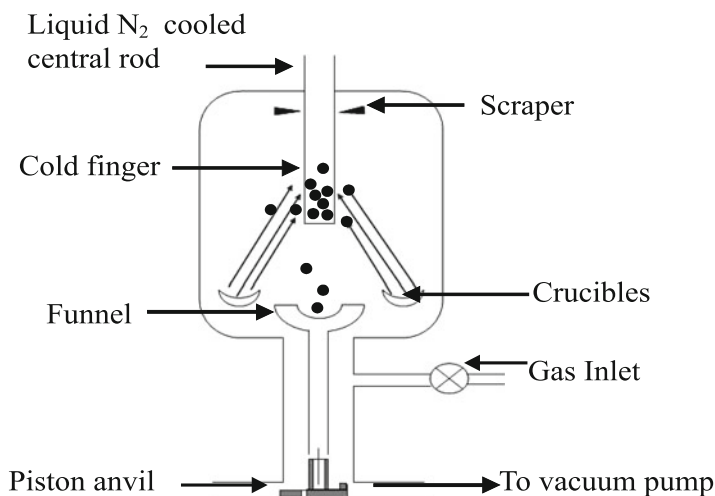


Fig. 3.6 Schematic diagram of synthesis of nanoparticles by physical vapour deposition. Particles condense on the central cooled rod, are scraped and consolidated using a piston-anvil system in vacuum to form a pellet

and piston-anvil (an arrangement in which nanoparticle powder can be compacted). All the processes are carried out in a vacuum chamber so that the desired purity of the end product can be obtained. Figure 3.6 schematically illustrates a set up for carrying out physical vapour deposition and compressing the powder in a pellet form.

Usually metals or high vapour pressure metal oxides are evaporated or sublimated from filaments or boats of refractory metals like W, Ta and Mo in which the materials to be evaporated are held. The density of the evaporated material close to the source is quite high and particle size is small (<5 nm). Such particles would prefer to acquire a stable lower surface energy state. Due to small particle or cluster-cluster interaction bigger particles get formed. Therefore, they should be removed away as fast as possible from the source. This is done by forcing an inert gas near the source, which removes the particles from the vicinity of the source. In general the rate of evaporation and the pressure of gases inside the chamber determine the particle size and their distribution. Distance of the source from the cold finger is also important. Evaporated atoms and clusters tend to collide with gas molecules and make bigger particles, which condense on cold finger. While moving away from the source to cold finger the clusters grow. If clusters have been formed on inert gas molecules or atoms, on reaching the cold finger, gas atoms or molecules may leave the particles there and then escape to the gas phase. If reactive gases like O_2 , H_2 and NH_3 are used in the system, evaporated material can interact with these gases forming oxide, nitride or hydride particles. Alternatively one can first make metal nanoparticles and later make appropriate post-treatment to achieve desired metal compound etc. Size, shape and even the phase of the evaporated material

can depend upon the gas pressure in deposition chamber. For example using gas pressure of H_2 more than 500 kPa, TiH_2 particles of ~ 12 nm size were produced. By annealing them in O_2 atmosphere, they could be converted into titania (TiO_2) having rutile phase. However if titanium nanoparticles were produced in H_2 gas pressure less than 500 kPa, they could not be converted into any crystalline oxide phase of titanium but always remained amorphous.

Clusters or nanoparticles condensed on the cold finger (water or liquid nitrogen cooled) can be scraped off inside the vacuum system. The process of evaporation and condensation can be repeated several times until enough quantity of the material falls through a funnel in which a piston-anvil arrangement has been provided. One can even have separate low and high pressure presses. A pressure of few mega pascal (MPa) to giga pascal (GPa) is usually applied depending upon the material. Low porosity pellets are easily obtained. Density of the material thus can be ~ 70 to 90 % of the bulk material.

3.3.2 Ionized Cluster Beam Deposition

This method was first developed by Takagi and Yamada around 1985 and is also useful to obtain adherent and high quality single crystalline thin films. The set up consists of a source of evaporation, a nozzle through which material can expand into the chamber, an electron beam to ionize the clusters, an arrangement to accelerate the clusters and a substrate on which nanoparticle film can be deposited, all housed in a suitable vacuum chamber. A schematic arrangement is illustrated in Fig. 3.7.

Small clusters from molten material are expanded through the fine nozzle. The vapour pressure ~ 1 Pa to 1 kPa needs to be created in the source and the nozzle needs to have a diameter larger than the mean free path of atoms or molecules in

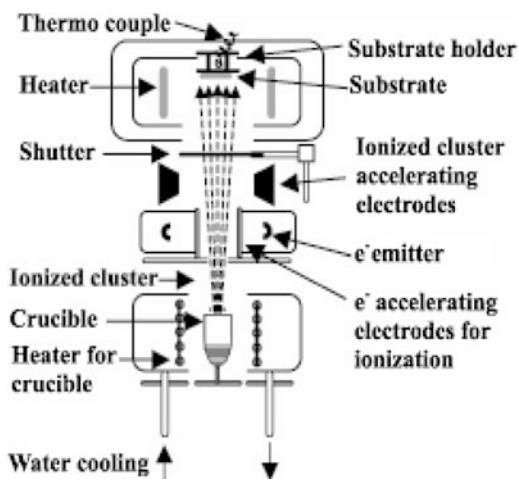


Fig. 3.7 Ionized cluster beam apparatus

vapour form in the source to form the clusters. On collision with electron beam, clusters get ionized. Due to applied accelerating voltage, the clusters are directed towards the substrate. By controlling the accelerating voltage, it is possible to control the energy with which the clusters hit the substrate. Stable clusters of some materials would require considerable energy to break their bonds and would rather prefer to remain as small clusters of particles. Thus it is possible to obtain the films of nanocrystalline material using ionized cluster beam. However it is not unlikely that some neutral atoms also get incorporated in the film. Besides, the clusters are not mass selected. Therefore they can have wide distribution of particle sizes. In fact anything that can cross the accelerating voltage can get incorporated in the film.

3.3.3 Laser Vapourization (Ablation)

In this method, vapourization of the material is effected using pulses of laser beam of high power. The set up along with the interaction with the evaporation source is depicted in Fig. 3.8.

The set up is an Ultra High Vacuum (UHV) or high vacuum system equipped with inert or reactive gas introduction facility, laser beam, solid target and cooled substrate. Clusters of any material of which solid target can be made are possible to synthesize. Usually laser operating in the UV range such as excimer (excited monomers) laser (see Table 3.1) is necessary because other wavelengths like IR or visible are often reflected by surfaces of some metals.

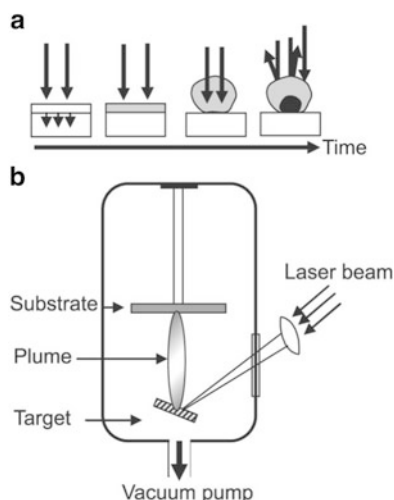
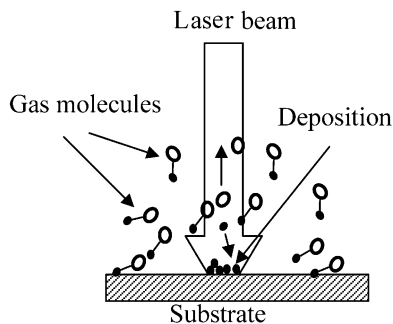


Fig. 3.8 (a) Sequence of material evaporation by laser beam interaction with a target material. (b) Laser deposition schematic apparatus

Table 3.1 Wavelengths of some commonly used excimer lasers

Gas	F ₂	ArF	KrCl	KrF	XeCl	XeF
λ , nm	157	193	222	249	308	350

Fig. 3.9 Schematic diagram depicting laser pyrolysis



A powerful beam of laser evaporates the atoms from a solid source and atoms collide with inert gas atoms (or reactive gases) and cool on them forming clusters. They condense on the cooled substrate. The method is often known as laser ablation. Gas pressure is very critical in determining the particle size and distribution. Simultaneous evaporation of another material and mixing the two evaporated materials in inert gas leads to the formation of alloys or compounds. This method can produce some novel phases of the materials which are not normally formed. For example Single Wall Carbon Nanotubes (SWNT) are mostly synthesized by this method.

3.3.4 Laser Pyrolysis

Another method of thin films synthesis using lasers is known as ‘laser pyrolysis’ or laser-assisted deposition. Here a mixture of reactant gases is decomposed using a powerful laser beam in presence of some inert gas like helium or argon. Atoms or molecules of decomposed reactant gases collide with inert gas atoms and interact with each other, grow and are then deposited on cooled substrate. Schematic diagram is given in Fig. 3.9.

Many nanoparticles of materials like Al_2O_3 , WC and Si_3N_4 are synthesized by this method. Here too, gas pressure plays an important role in deciding the particle sizes and their distribution.

3.4 Sputter Deposition

Sputter deposition is a widely used thin film deposition technique, specially to obtain stoichiometric thin films (i.e. without changing the composition of the original material) from target material. Target material may be some alloy, ceramic or compound. Sputtering is also effective in producing non-porous compact films. It is a very good technique to deposit multilayer films for mirrors or magnetic films for spintronics applications.

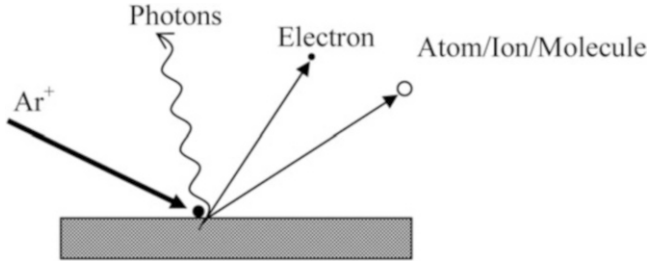


Fig. 3.10 Interaction of an ion with target

In sputter deposition, some inert gas ions like Ar^+ are incident on a target at a high energy. Depending on the energy of ions, ratio of ion mass to that of target atoms mass, the ion-target interaction can be a complex phenomenon. The ions become neutral at the surface but due to their energy, incident ions may get implanted, get bounced back, create collision cascades in target atoms, displace some of the atoms in the target creating vacancies, interstitials and other defects, desorb some adsorbates, create photons while losing energy to target atoms or even sputter out some target atoms/molecules, clusters, ions and secondary electrons. Figure 3.10 shows a schematic picture of various possibilities.

For deposition of materials, one is interested in sputtering out the target material. Target material sputter yield is given by

$$Y = \frac{3\alpha}{4\pi^2} \frac{4M_1M_2}{(M_1 + M_2)^2} \frac{E_1}{E_b} \quad (3.2)$$

$$\text{for } E_1 < 1 \text{ keV}$$

$$\text{and } Y = 3.56\alpha \frac{Z_1Z_2}{Z_1^{2/3} + Z_2^{2/3}} \left(\frac{M_1}{M_1 + M_2} \right) \frac{S_n(E_2)}{E_b} \quad (3.3)$$

$$\text{for } E_2 > 1 \text{ keV}$$

where α is efficiency of momentum transfer, M_1 – mass of incident ion, M_2 – mass of target atom, Z_1 – atomic number of incident ion, Z_2 – atomic number of target atom, E_1 – energy of incident ion, E_b – binding energy of target atom and S_n is known as stopping power. It is an energy loss per unit length due to nuclear collisions.

Sputter yield for different elements with same incident ion having same energy varies in general. This would lead one to think that from a target consisting of two different elements or more, the one having higher sputter yield should get incorporated in larger quantity than the others. However high sputter yield elements get depleted fast and other elements also make contribution. Thus the stoichiometry is achieved in the deposited film.

Sputter deposition can be carried out using Direct Current (DC) sputtering, Radio Frequency (RF) sputtering or magnetron sputtering. In all the above cases one can also use discharge or plasma of some inert gas atoms or reactive gases. The deposition is carried out in a high vacuum or ultra high vacuum system equipped with electrodes, one of which is a sputter target and the other is a substrate, gas introduction facility etc. Although the system during deposition is at high gas pressure, low base pressure ensures that the adequate purity is obtained.

3.4.1 DC Sputtering

This is a very straight forward technique of deposition, in which sputter target is held at high negative voltage and substrate may be at positive, ground or floating potential (see Fig. 3.12). Substrates may be simultaneously heated or cooled depending upon the material to be deposited. Once the required base pressure is attained in the vacuum system, usually argon gas is introduced at a pressure <10 Pa. A visible glow is observed and current flows between anode and cathode indicating the deposition onset. When sufficiently high voltage is applied between anode and cathode with a gas in it a glow discharge is set up with different regions as cathode glow, Crooke's dark space, negative glow, Faraday dark space, positive column, anode dark space and anode glow (Box 3.3). These regions are the result of plasma, i.e. a mixture of electrons, ions, neutral atoms and photons released in various collisions. The density of various particles and the length over which they are distributed depends upon the gas pressure. Energetic electron impacts cause gas ionization. Ratio of ions/neutrals can be typically $\sim 10^{-4}$. Thus at a few Pa pressure, sufficiently large number of ions are generated that can be used to sputter the target.

Box 3.3: Creating Plasma

The terms glow discharge and plasma are often used to mean the same thing. Plasma is a mixture of free electrons, ions and photons. Plasma is overall neutral but there can be regions which are predominantly of positive or negative charges. One can get plasma in different gases at different pressures by using a vacuum tube with two metal electrodes and applying high DC or AC voltage. Different regions of plasma generated by applying a high DC voltage can be divided into five regions as

- (i) Cathode or Crook's dark space,
- (ii) Negative space glow,
- (iii) Faraday's dark space,
- (iv) Positive column,
- (v) Anode dark space.

(continued)

Box 3.3 (continued)

The extent of these regions depends upon the pressure of the gas.
The glow discharge in a glass tube with two electrodes at the end filled with air has following characteristics at different pressures (Fig. 3.11).

Pressure (torr)	Appearance of discharge spark
~10	General glow discharge
~1–0.5	Closely spaced striations in positive column
~0.5	10 mm spacing of striations
~0.3	Crook’s dark space 5 mm long
~0.1	Crook’s dark space 10 mm long
~0.05	Crook’s dark space 20 mm long
~0.005	Discharge disappearance

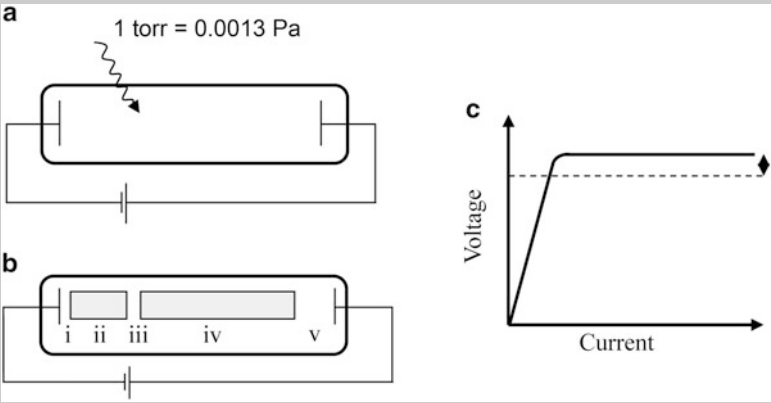


Fig. 3.11 (a) Generation of plasma by applying high voltage between two electrodes in an evacuated glass tube; (b) Different regions of plasma and (c) Typical current-voltage characteristic

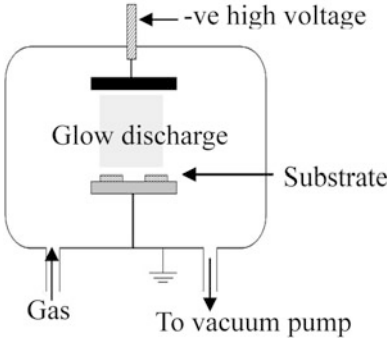
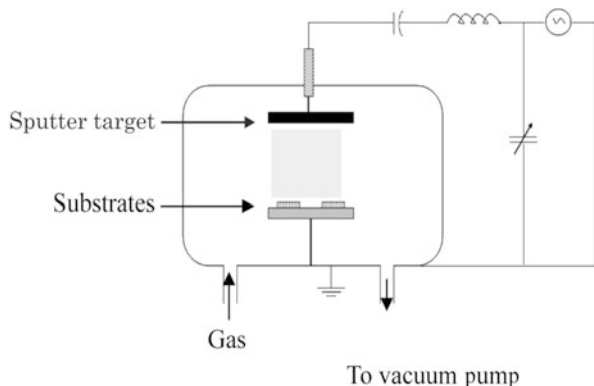


Fig. 3.12 Schematic layout of a typical DC sputtering unit

Fig. 3.13 RF sputtering set up with matching network



3.4.2 RF Sputtering

If the target to be sputtered is insulating, it is difficult to use DC sputtering. This is because it would mean the use of exceptionally high voltage ($>10^6$ V) to sustain discharge between the electrodes. (In DC discharge sputtering 100–3,000 V is usual.) However if some high frequency voltage is applied the cathode and anode alternatively keep on changing the polarity and oscillating electrons cause sufficient ionization.

In principle, 5–30 MHz frequency can be used and the electrodes can be insulating. However, 13.56 MHz frequency is commonly used for deposition, as this frequency is reserved worldwide for this purpose and many others are available for communication. Target itself biases to negative potential becoming cathode when the arrangement as depicted in Fig. 3.13 is used.

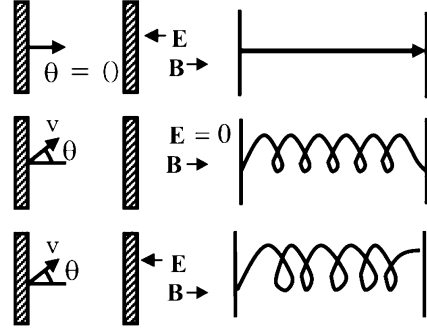
3.4.3 Magnetron Sputtering

RF/DC sputtering rates can further be increased by using magnetic field. When both electric and magnetic fields act simultaneously on a charged particle, the force on it is given by Lorentz force.

$$\mathbf{F} = q (\mathbf{E} + \mathbf{v} \times \mathbf{B}) \quad (3.4)$$

If $\mathbf{E}$ and $\mathbf{B}$ are parallel to each other as illustrated in Fig. 3.14 and an electron leaves the cathode at angle $\theta = 0^\circ$, then $\mathbf{v} \times \mathbf{B} = 0$ and only electric field vector acts on an electron. However an electron making an angle θ would have both electric and magnetic fields acting on it. The magnitude of velocity component in the direction of electric field would be $q\mathbf{v} \cos \theta$. The component due to magnetic field would be

Fig. 3.14 Effect of $\mathbf{E}$ and $\mathbf{B}$ on electron



perpendicular to both $\mathbf{v}$ and $\mathbf{B}$ with a magnitude $qvB \sin \theta$. Electron orbits around the magnetic field axis with radius r . The centrifugal force on it would be $m (v \sin \theta)^2/r$

$$\frac{mv^2 \sin^2 \theta}{r} = qvB \sin \theta \quad (3.5)$$

$$\text{Thus } r = \frac{mv \sin \theta}{qB} \quad (3.6)$$

Electron moves in a helical path and is able to ionize more atoms in the gas. In practice, both parallel and perpendicular magnetic fields to the direction of electric field are used to further increase the ionization of the gas, increasing the efficiency of sputtering. By introducing gases like O_2 , N_2 , NH_3 , CH_4 and H_2S , while metal targets are sputtered, one can obtain metal oxides like Al_2O_3 , nitrides like TiN and carbides like WC . This is known as ‘reactive sputtering’.

3.4.4 ECR Plasma Deposition

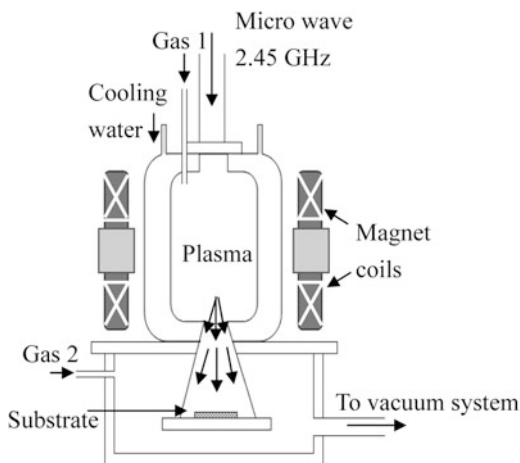
The plasma density can be further enhanced using microwave frequency and coupling the resonance frequency of electrons in magnetic field as

$$\omega = \frac{eB}{m} \quad (3.7)$$

where ω is microwave frequency, often 2.45 GHz, e – charge of electron, m – mass of free electron and B is magnetic field (870 G).

Ionization density using ECR plasma is about 2–3 orders of magnitude larger than that by previous methods. Thin films and nanoparticles of SiO_2 , SiN and GaN have been obtained by using this technique. Schematic of this technique is shown in Fig. 3.15.

Fig. 3.15 Schematic diagram for Electron Cyclotron Resonance (ECR) for deposition



3.5 Chemical Vapour Deposition (CVD)

Chemical vapour deposition, a hybrid method using chemicals in vapour phase is conventionally used to obtain coatings of a variety of inorganic or organic materials. It is widely used in industry because of relatively simple instrumentation, ease of processing, possibility of depositing different types of materials and economical viability. Under certain deposition conditions nanocrystalline films or single crystalline films are possible. There are many variants of CVD like Metallo Organic CVD (MOCVD), Atomic Layer Epitaxy (ALE), Vapour Phase Epitaxy (VPE), Plasma Enhanced CVD (PECVD). They differ in source gas pressure, geometrical layout and temperature used. Basic CVD process, however, can be considered as a transport of reactant vapour or reactant gas towards the substrate (see Fig. 3.16) kept at some high temperature where the reactant cracks into different products which diffuse on the surface, undergo some chemical reaction at appropriate site, nucleate and grow to form the desired material film. The by-products created on the substrate have to be transported back to the gaseous phase removing them from substrate.

Vapours of desired material may be often pumped into reaction chamber using some carrier gas. In some cases the reactions may occur through aerosol formation in gas phase. There are various processes such as reduction of gas, chemical reaction between different source gases, oxidation or some disproportionate reaction by which CVD can proceed. However it is preferable that the reaction occurs at the substrate rather than in the gas phase. Usually $\sim 300\text{--}1,200\text{ }^{\circ}\text{C}$ temperature is used at the substrate. There are two ways viz. 'hot wall' and 'cold wall' by which substrates are heated as shown in Fig. 3.17.

In hot wall set up the deposition can take place even on reactor walls. This is avoided in cold wall design. Besides this, the reactions can take place in gas phase with hot wall design which is suppressed in cold wall set up. Further, coupling of plasma with chemical reaction in cold wall set up is feasible.

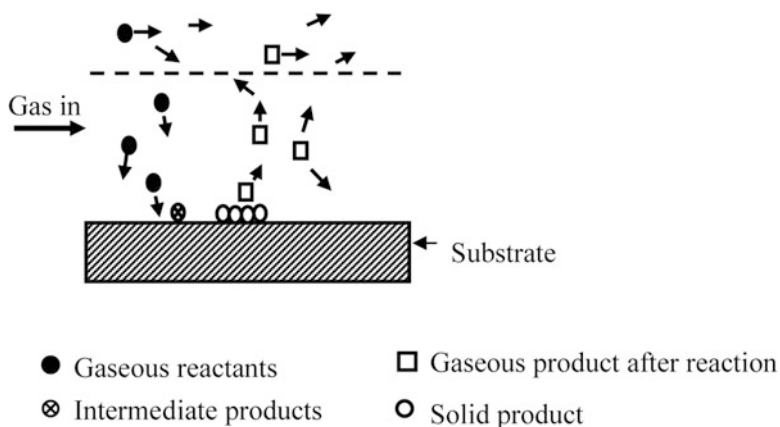


Fig. 3.16 Basic concept of Chemical Vapour Deposition (CVD) process

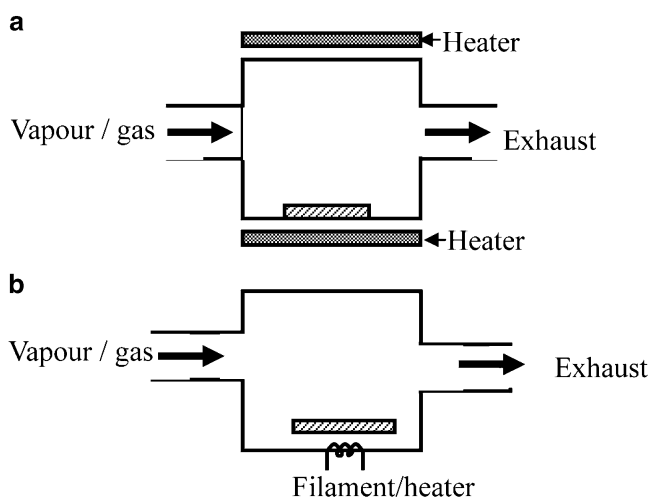


Fig. 3.17 (a) Hot wall and (b) cold wall CVD

Usually gas pressures in the range of $100\text{--}10^5$ Pa are used. Growth rate and film quality depend upon the gas pressure and the substrate temperature. When the growth takes place at low temperature, it is limited by the kinetics of surface reaction. At intermediate temperature it is limited by mass transport i.e. supply of reacting gases to the substrate. Here the reaction is faster and supply of reactants is slower. At high temperature, growth rate reduces due to desorption of precursors from the substrate.

When two types of atoms or molecules say P and Q are involved in the desired film, there are two ways in which growth can take place. In what is known as Langmuir-Hinshelwood mechanism, both P and Q type of atoms/molecules are

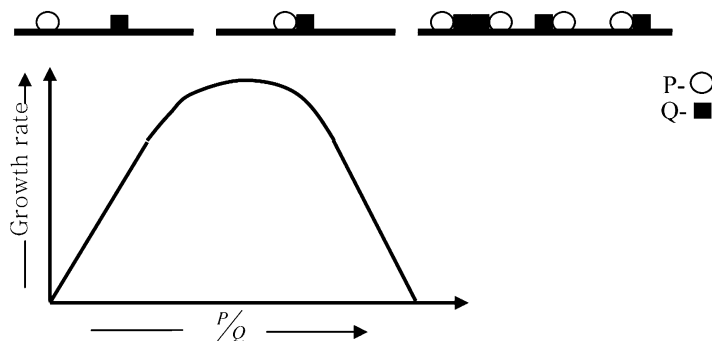


Fig. 3.18 Langmuir-Hinshelwood mechanism of growth

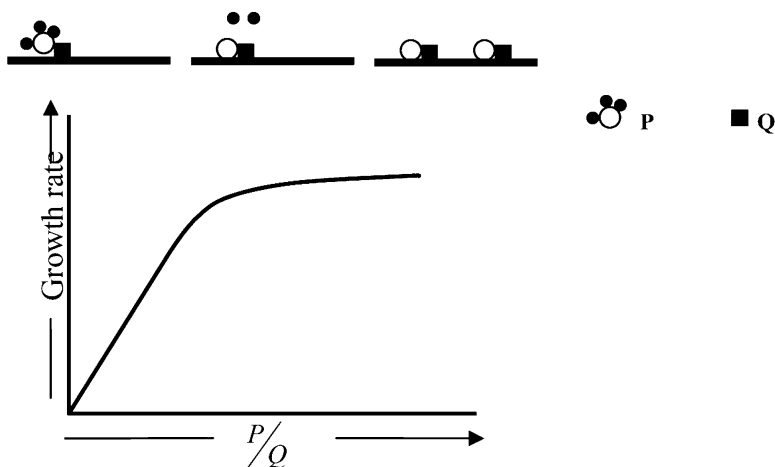


Fig. 3.19 Elay-Riedel mechanism

adsorbed on the substrate surface and interact there to produce the product PQ . When one species is adsorbed in excess of the other, the growth depends on the availability of adsorption sites for both P and Q as shown in a schematic diagram (Fig. 3.18).

However it is also possible to have another way in which reaction can occur i.e. one species say P adsorbs on the substrate and the species Q from gas phase interacts with P . Thus there is no sharing of sites. This type of mechanism is known as Elay-Riedel mechanism as shown in Fig. 3.19.

3.6 Electric Arc Deposition

This is one of the simplest and useful methods which leads to mass scale production of Fullerenes, carbon nanotubes (to be discussed in a later chapter) etc.

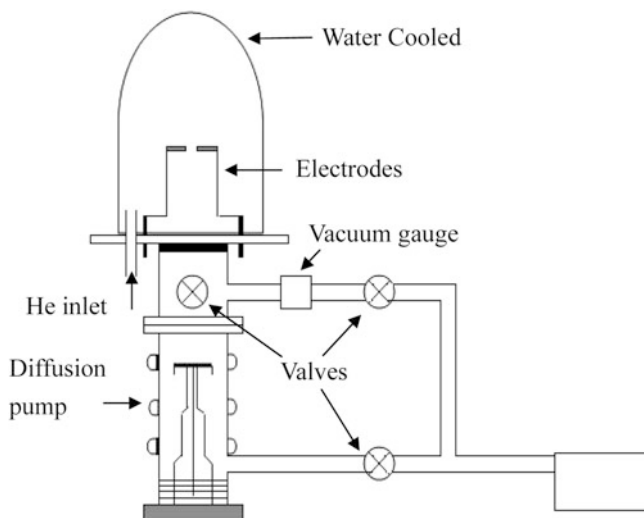


Fig. 3.20 Arc deposition set up

The set up as shown in Fig. 3.20 requires basically water cooled vacuum chamber and electrodes to strike an arc between them. The positive electrode itself acts as the source of material. If some catalysts are to be used, there can be some additional thermal source of evaporation. Depending upon the desired products inert gas or reactive gas introduction is necessary.

Usually the gap between the electrodes is ~ 1 mm and high current ~ 50 – 100 A is passed from a low voltage power supply (~ 12 – 15 V). Inert gas pressure is maintained in the vacuum system. When an arc is set up, anode material evaporates. This is possible as long as the discharge can be maintained. The adjustment of the electrode gap without breaking the vacuum becomes essential, as one of the electrode burns and gap increases.

As will be discussed later, this method was the one in which by striking the arc between the two graphite electrodes, it was possible to get fullerenes in large quantity. Fullerenes are formed at low helium pressure and nanotubes are formed at high pressure. Also, fullerenes are obtained by purification of soot collected from inner walls of vacuum chamber, whereas nanotubes are found to be formed only at high He gas pressure and in the central portion of the cathode. No carbon nanotubes are found on the chamber walls. Some nanoparticles of carbon also are usually found around the region where nanotubes are formed. Temperature reaches as high as $\sim 3,500$ °C as arc discharge takes place. In principle, other nanocrystals or tubes also should be possible to be obtained by this method. However this method is mostly found to be suitable for deposition of fullerenes or carbon nanotubes.

3.7 Ion Beam Techniques (Ion Implantation)

There are many examples in which high energy (few keV to hundreds of keV) or low energy (<200 eV) ions are used to obtain nanoparticles. Ions of interest are usually formed using an ion gun specially designed to produce metal ions which are accelerated to high or low energy towards the substrate heated to few hundreds of $^{\circ}\text{C}$. Depending upon the energy of the incident ions, various other processes like sputtering and electromagnetic radiation may take place as already discussed in connection with sputter deposition in Sect. 3.3. It is possible to obtain single element nanoparticles or compounds and alloys of more than one element. Post-annealing also is used sometimes to improve the crystallinity of the materials. In some experiments it has been possible to even obtain doped nanoparticles (i.e. nanoparticles in which some foreign atoms are intentionally introduced to alter the properties of the host material in the controlled fashion) using ion implantation.

There is also a possibility of making nanoparticles using swift heavy ions (few MeV energy) employing high energy ion accelerators like pelletrons.

3.8 Molecular Beam Epitaxy (MBE)

This technique can be used to deposit elemental or compound quantum dots, quantum wells as well as quantum wires in a very controlled manner. A schematic diagram is given in Fig. 3.21. High degree of purity is achievable using ultra high vacuum (better than 10^{-8} Pa). Special sources of deposition known as Knudsen cell (K-cell) or effusion cell are employed to obtain molecular beams of the constituent elements. The rate of deposition is kept very low and substrate temperature is rather high in order to achieve sufficient mobility of the elements on the substrate and layer by layer growth to obtain nanostructures or high purity thin films. Technique like Reflected High Energy Electron Diffraction (RHEED) is incorporated to monitor the high crystallinity of the growing film.

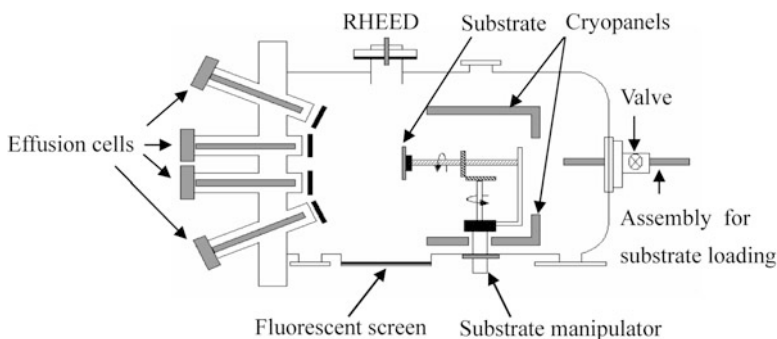


Fig. 3.21 Schematic diagram of molecular beam epitaxy

Further Reading

K.L. Chopra, *Thin Films Phenomenon* (McGraw-Hill Book, New York, 1969)

H. Gleiter, Prog. Mater. Sci. **33**, 223 (1989)

L.I. Maissel, R. Glang, *Handbook of thin film Technology* (McGraw-Hill Book, New York, 1970)

P. Milani, S. Iannotta, *Cluster Beam Synthesis of Nanostructured Materials* (Springer, Berlin/Heidelberg/New York, 1999)

M. Ohring, S. Iannotta, *Material Science of Thin Films, Deposition and Structure*, 2nd edn. (Academic, Boston, 2002)

R. Vyeda, Prog. Mater. Sci. **35**, 1 (1991)

Chapter 4

Synthesis of Nanomaterials—II

(Chemical Methods)

4.1 Introduction

In this chapter we shall discuss some of the wet chemical methods to synthesize nanomaterials. There are numerous advantages of using chemical methods which are summarized in Box 4.1. In some cases nanomaterials are obtained as colloidal particles in solutions, which can be filtered and dried to obtain powder. We can obtain thin films or nanoporous materials by electrodeposition and etching. Advantages of chemical synthesis are manifold. In many cases very well known chemical reaction route can be optimized to obtain nanoparticles. Particles of different shapes and sizes are possible depending upon the chemicals used and reaction conditions.

Box 4.1: Some Advantages of Chemical Synthesis

- Simple techniques
- Inexpensive, less instrumentation compared to many physical methods
- Low temperature (<350 °C) synthesis
- Doping of foreign atoms (ions) possible during synthesis
- Large quantities of the materials can be obtained
- Variety of sizes and shapes are possible
- Materials are obtained in the form of liquid but can be easily converted into dry powder or thin films
- Self assembly or patterning is possible

Low to high temperature routes are possible. In many cases, particles can be doped with different metal ions quite easily. Coupled, coated, chemically capped (by some molecules or passivated) particles can be made. Most important is that very narrow size distributed materials are possible to synthesize and in many cases

large quantities of materials can be obtained. The instrumentation involved in the chemical synthesis can be relatively simple and inexpensive as compared to many physical methods. As in many cases nanoparticles synthesized by chemical method form what is known as ‘colloids’. We shall first try to understand them and then proceed to some specific chemical routes to obtain them as nanoparticles.

4.2 Colloids and Colloids in Solutions

Colloids are known since very long time. A class of materials, in which two or more phases (solid, liquid or gas) of same or different materials co-exist with the dimensions of at least one of the phases less than a micrometre is known as colloids. Colloids may be particles, plates or fibres (see Fig. 4.1). Nanomaterials are a subclass of colloids, in which one of the dimensions of colloids is in nanometre range.

There are several examples around us, having different combinations of phases, in the form of colloids like liquid in gas (fog), liquid in liquid (fat droplets in milk), solid in liquid (tooth paste), solid in solid (tinted glass), gas in liquid (foam). There can be multiple existing colloids like water and oil bubbles in porous mineral rocks. Organic and inorganic materials can be dispersed into each other to form colloids. Several examples exist even of bio-colloids. Blood and bones are good examples of bio-colloids. Blood has corpuscles dispersed in serum and bone has colloids of calcium phosphate embedded in collagen.

Colloids may even form networks. For example aerogels (discussed in more details in Chap. 9) are a network of silica colloidal particles, pores of which are filled with air.

4.2.1 Interactions of Colloids and Medium

Colloids are particles with large surface to volume ratio. Correspondingly there are large number of atoms/molecules on the surface of a colloidal particle, which

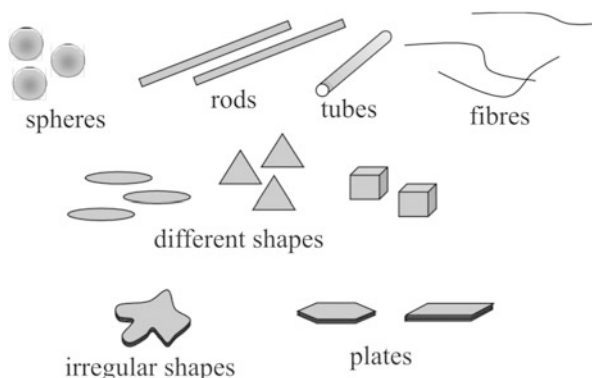


Fig. 4.1 Different shapes of colloids

do not have as many neighbours as those for an atom/molecule inside the interior. Therefore atoms on the surface are in a highly reactive state, which easily interact to form bigger particles or tend to coagulate. It is thus necessary to understand the stability of colloids i.e. how the colloids dispersed in a medium can remain as separated particles. In general there are a number of interactions involved. For the sake of understanding these interactions, we consider the inorganic spherical colloids of equal size, dispersed in a liquid medium. When fine particles are dispersed in a liquid medium, it is known that they undergo *Brownian motion* (Box 4.2). If we are able to tag a particle in the solution, as depicted in Fig. 4.2, it would appear as if it is making a random motion. All other particles also execute random motion, hitting each other and changing direction of motion in solution. Distance travelled between successive collisions is random too. However an average distance travelled by a colloidal particle can be found as

$$\Delta \bar{R}^2 = \left(\frac{kT}{3\pi r\eta} \right) \Delta t \quad (4.1)$$

where $\Delta \bar{R}$ is distance travelled by a particle from its original position in time Δt , k – Boltzmann's constant, T – temperature of liquid, r – particle radius and η is viscosity of the liquid.

Box 4.2: Robert Brown (1773–1858)

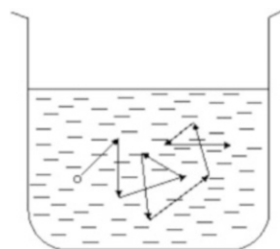
Robert Brown who discovered Brownian motion in 1827 was born on 21st December 1773 in Montrose, Scotland. He studied medicine at the University of Edinburgh. Immediately after completing his education he joined Fifeshire Regiment of Fencibles as Ensign and Surgeon's Mate. In his free time he used to study botany. On 18th July 1801 he joined an eminent botanist Sir Joseph Banks to go to Australia on an expedition. Their goal was to collect some rare plants and study them. They returned back to England in October 1805 with a huge collection of some 4,000 species of plants, some zoological specimens and numerous drawings and notes. Brown spent next five years working on this material. Brown used a microscope throughout his studies. He is considered to be a 'gifted observer'. He identified naked ovule in the gymnosperm which is a rather difficult task even using modern microscopes. He however is famous due to 'Brownian motion' which he observed in pollen grains using his simple optical microscope. Others also had observed such motion of particles under microscope but they simply attributed it to 'life' itself. Careful experiments by Brown showed that they were not the consequence of 'life' nor any currents in fluids or evaporation of liquid.

(continued)

Box 4.2 (continued)

From 1806 to 1822 Robert Brown served as a clerk, librarian and housekeeper of the Linnen Society of London. He was elected as fellow of Royal Society in 1810 and fellow of Linnen Society in 1822. He remained the president of Linnen Society from 1849 to 1853. He died on 10th June 1858 in London.

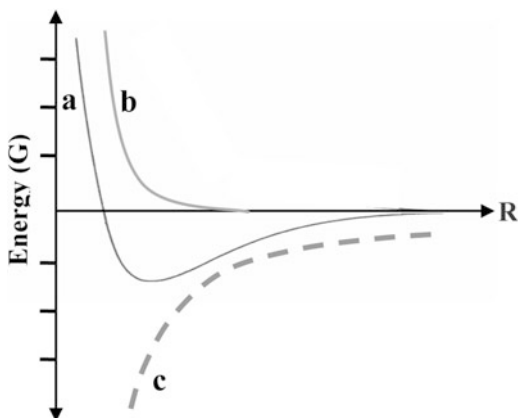
Fig. 4.2 Brownian motion of colloidal particles



Interactions of such constantly and randomly moving particles with each other and with liquid in general would be quite complex. Consider for example the van der Waals interaction (see Chap. 2) between two atoms or molecules. There are two parts in the interaction viz. attractive and repulsive parts given by the Eq. (4.2), irrespective of whether there exist permanent dipoles or not. The interaction is given as

$$dG_1 = \left(\frac{A}{R^{12}} \right) - \left(\frac{B}{R^6} \right) \quad (4.2)$$

Fig. 4.3 Van der Waals interaction: (a) Resulting energy, (b) Repulsive energy and (c) Attractive energy



where dG_I is the interaction energy, A and B are constants and R is the distance between two particles.

Here the first term is repulsive interaction (Born repulsive interaction) effective only at short distance and second term represents long range attractive interaction (van der Waals attraction). Repulsive part arises due to repulsion between electron clouds in each atom and attractive part is due to interaction between fluctuating or permanent dipoles of atoms/molecules. Schematically it is shown in Fig. 4.3. Equation (4.2) is known as Lennard-Jones equation (Box 4.3).

Box 4.3: Free Energy

Free energy of a body is a measure of its ability to do work. A body always tends to attain the state of lower energy by releasing 'free energy' and is given by Gibb's free energy

$$G = H - TS \quad (4.3)$$

where G is the energy absorbed or released by a body, H – enthalpy, T – temperature and S is entropy.

In order to understand the meaning of surface free energy, consider that a cylinder with cross section ' A ' and length ' $2l$ ' is cut into two equal pieces as shown in Fig. 4.4. Work has to be done in order to break the cylinder into two pieces and separate them by a distance R .

Energy increases as the distance between the two pieces increases. The two pieces keep on attracting each other due to intermolecular forces at the broken interface. Therefore, more and more energy has to be supplied in order to separate them. Interfacial energy is related to the surface tension ' γ '.

(continued)

Box 4.3 (continued)

Fig. 4.4 A cylinder with cross section A and length 2ℓ has been cut into two pieces of equal length

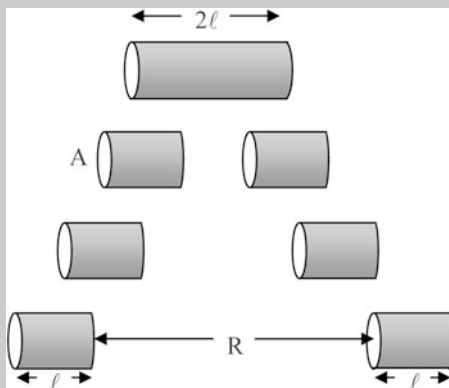
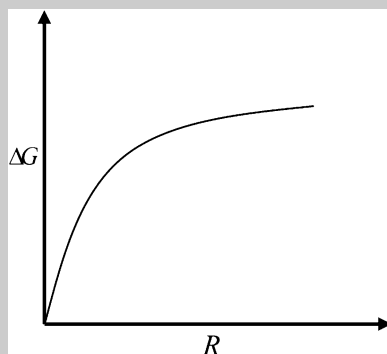


Fig. 4.5 Change in surface free energy with separation R between two pieces of the cylinder shown in Fig. 4.4



Surface tension γ is along the normal to surface and is expressed as force per unit length (N/m or Joule/m²). Surface free energy is given by

$$\Delta W = \Delta G - 2\gamma A \quad (4.4)$$

where ΔW is the work done to break the cylinder, γ —surface tension and A is the surface area. (Factor ‘2’ arises because while creating two surfaces, surface area would be $2A$ for two surfaces as shown in Fig. 4.4).

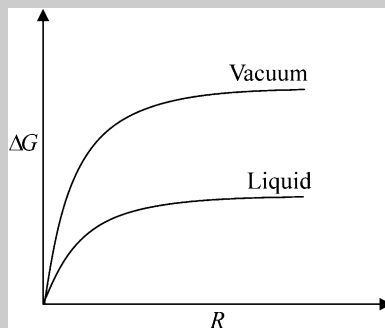
The surface free energy per unit volume would increase as shown in Fig. 4.5 with separation R between two particles.

The force of attraction would increase in vacuum as shown in Fig. 4.6. In a liquid the force of attraction would reduce in general.

(continued)

Box 4.3 (continued)

Fig. 4.6 Change in surface free energy with separation R between two pieces of the cylinder shown in Fig. 4.4, placed in vacuum or liquid



4.2.2 Colloids in Vacuum

Lennard-Jones equation (4.2) is sufficient to describe an interaction between two atoms or molecules. When we consider colloids with large number of atoms, we need to take into account all the atoms and their interactions with each other. This is quite a complex situation as shown in Fig. 4.7. To describe the interaction between colloidal particles Derjaguin, Landau, Verwey and Overbeek proposed a theory known as DLVO theory. In order to reduce the complexity of the problem, they assumed two interacting spherical particles of equal size. Let the radius of each particle be ' r ' and let two particles be separated by a distance ' R '.

It was shown that for two similar spherical particles in vacuum the attractive interaction is given by

$$dG(\text{attraction}) = - \left(\frac{A_H \cdot r}{12R} \right) \left[1 + \frac{3}{4} \cdot \frac{R}{r} \right] \quad (4.5)$$

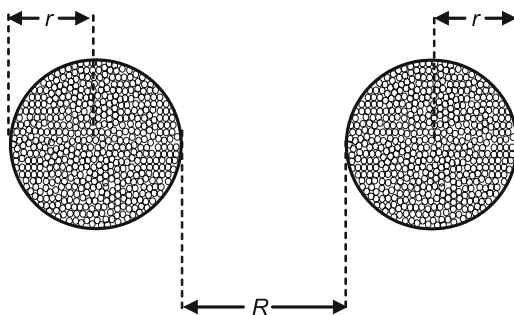


Fig. 4.7 Interaction even between two spherical particles of same material and same size is complex due to presence of large number of atoms in each particle

where A_H is known as Hamaker constant and is given by

$$A_H = A' \pi^2 n^2 \quad (4.6)$$

where A' is a constant related to A in the Lennard-Jones equation (4.2) and n is the number of atoms/molecules per unit volume in a colloid.

4.2.3 Colloids in a Medium

So far we considered the colloids just in vacuum. Consider now a situation, in which inorganic spherical colloids are immersed in a liquid (and do not dissolve). The attractive interactions between the colloids get modified through the change of Hamaker constant as A_H , which can be written now as

$$A_H = \left(\sqrt{A_{1v}} - \sqrt{A_{2v}} \right)^2 \quad (4.7)$$

where A_{1v} is the Hamaker constant for particle of inorganic solid under consideration, in vacuum and A_{2v} is Hamaker constant of colloid of medium in vacuum. It can be seen from above equation that in general the effect of liquid medium is to reduce the Hamaker constant of colloid particle. Hence the attractive force between colloid particles will in general reduce.

4.2.4 Effect of Charges on Colloids

Colloids in liquid may be positively charged, negatively charged or even neutral. But in most of the cases they are charged. There are various sources for colloids by which they acquire charges on their surfaces viz. through composition of colloidal material, properties of dispersing medium including the type and concentration of dissolved ions in the solution. In any case as soon as there are some charges on particles, ions of opposite charges accumulate around them. Oppositely charged ions are known as *counter ions*.

This accumulation of counter ions leads to formation of an electric double layer. Helmholtz considered that the situation is like that in a parallel plate condenser, where there are opposite charges as plates separated by a medium (for example air as shown in Fig. 4.8). Due to their Brownian motion, counter ions are not fixed nor are colloidal particles at rest. They execute their own Brownian motion and form a dynamic double layer around them as schematically shown in Fig. 4.9.

The changes occur when concentrations of charges on colloid and local ionic charges in solution contribute to double layer change. The concentration of electrolyte would strongly affect the electric potential curve and is shown schematically in Fig. 4.10.

Fig. 4.8 Double layer of charges and fall of electrical potential from negatively charged plate with positive ion accumulation

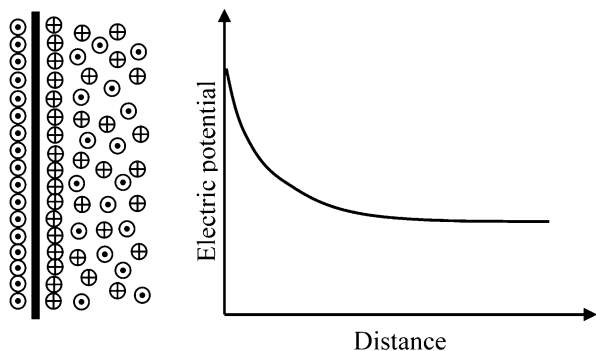


Fig. 4.9 Diffuse double layer because of Brownian motion

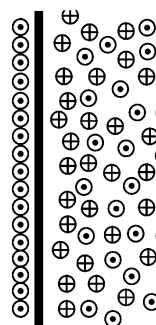
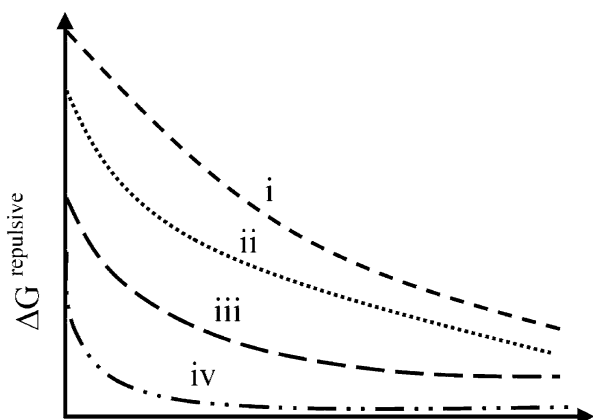


Fig. 4.10 Effect of concentration on change of potential from a particle.
Curves i, ii, iii and iv refer to increasing electrolyte concentrations



Consider now a situation in which two charged colloidal particles come closer with their electric double charge layers. As they approach each other force of repulsion increases as shown in Fig. 4.11. It is easy to see that the difference in concentrations of double layer charges would play an important role.

The repulsive interaction for low concentration would set at longer distance and that for higher concentration at smaller distance.

Fig. 4.11 Colloid-colloid interaction with electric double charge layer

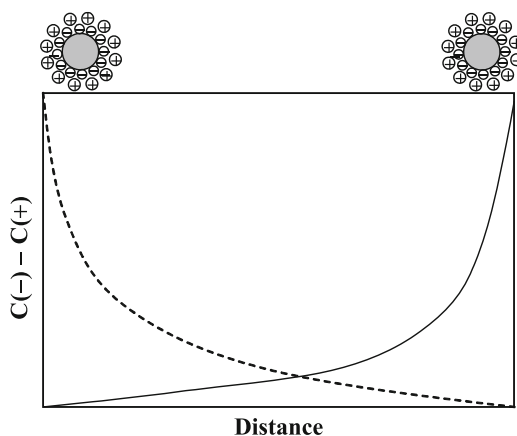


Fig. 4.12 Colloids with coatings. Separation between the particles, each coated with coating of thickness δ , increases by 2δ

4.2.5 Stearic Repulsion

Stability of colloids can be increased by creating Stearic hindrance or repulsion (see Fig. 4.12).

By adsorbing some layers of a different material on colloidal particles e.g. polymer or organic molecules on inorganic colloidal particles, it is possible to reduce the attractive forces between them. With addition of adsorbed layers, the effective sizes of the particles change which helps them to stay at a longer distance from each other, reducing the attractive interaction. However, in case where the coating material is similar in properties to that of the solvent, the effect of coating would be negligible. By anchoring long chain molecules on the particles, it is possible to keep them apart with negligible interaction. This idea is the basis of 'capped nanoparticles' discussed in the next section.

Thus the colloids interact with each other dynamically and are affected by van der Waals forces, colloid-colloid interaction mediated through dispersing medium, electric double layer and Stearic interactions. All the interactions may not be set in for every case. In general, contributions of these various attractive and repulsive

interactions—also dictated by temperature, concentration of colloids and dispersing medium—are additive and can be written as follows.

$$\begin{aligned}\Delta G = & \Delta G_1 \text{ (attractive and repulsive)} + \Delta G_2 \text{ (colloid-colloid attractive)} \\ & + \Delta G_3 \text{ (electrostatic repulsion)} + \Delta G_4 \text{ (Stearic repulsion)} \\ & + \Delta G_5 \text{ (any other)}\end{aligned}\quad (4.8)$$

If the repulsive forces are strong enough, colloids would be stabilized. Otherwise ripening, coagulation, or network formation may take place.

4.2.6 *Synthesis of Colloids*

Colloids are thus phase separated submicrometre particles in the form of spherical particles or particles of various shapes and sizes like rods, tubes, plates. They are the particles suspended in some host matrix. Metal, alloy, semiconductor and insulator particles of different shapes and sizes can be synthesized in aqueous or non-aqueous media. Colloidal particles in liquids are stabilized as discussed above by Coulombic repulsion, which arises due to similar charges they may have acquired on their surfaces. In some cases surface passivating molecules may be used which provide sufficient stearic hindrance inhibiting coalescence or aggregation. Nanomaterials are a special class of colloidal particles which are few hundreds of nanometre or smaller in size.

Synthesis of colloids is a very old method. Making nanoparticles using colloidal route goes back to nineteenth century when M. Faraday synthesized gold nanoparticles by wet chemical route. The particles are so stable that even today the colloidal solution made by him can be seen in the British Museum in London.

Here we shall discuss some commonly used synthesis methods of metal, semiconductor and insulator nanoparticles with some examples.

Chemical reactions in which colloidal particles are obtained are carried out in some glass reactor of suitable size. Glass reactor usually has a provision to introduce some precursors, gases as well as measure temperature and pH during the reaction. It is usually possible to remove the products at suitable time intervals. Reaction is usually carried out under inert atmosphere like argon or nitrogen gas so as to avoid any uncontrolled oxidation of the products. There is also a provision made to stir the reactants during the reaction by using teflon-coated magnetic needle. Figure 4.13 illustrates a simple chemical synthesis set up to obtain nanoparticles by colloidal route.

4.3 Nucleation and Growth of Nanoparticles

Synthesis of nanoparticles of different shapes and sizes may appear as a complex process. Over several decades, scientists have tried to understand the process of atom-by-atom nucleation and growth of small to large particles in melts, aqueous or

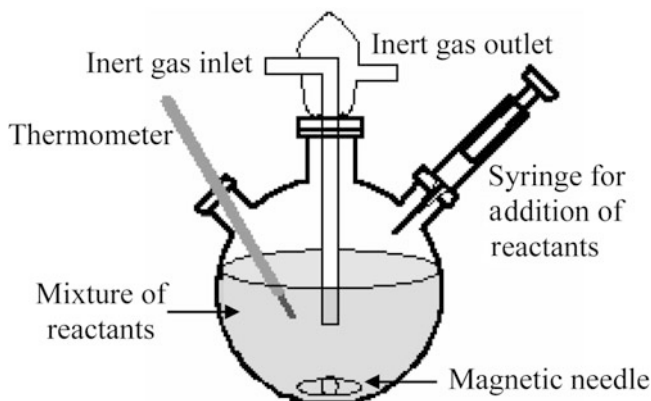


Fig. 4.13 A typical chemical reactor to synthesise nanoparticles

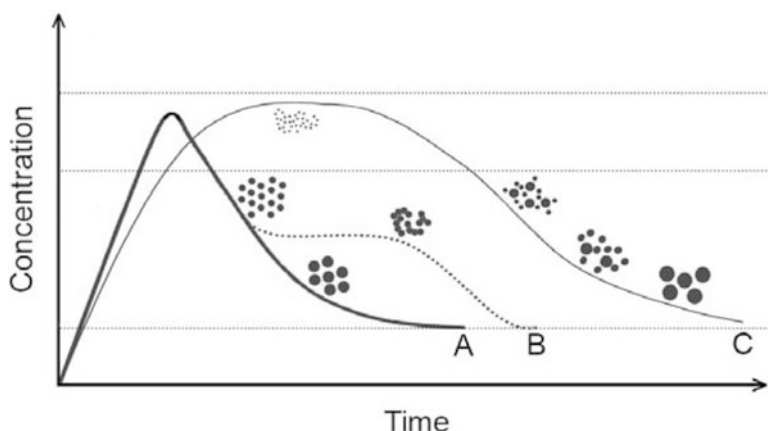


Fig. 4.14 Nucleation and growth of nanoparticles (LaMer diagram). All the dots appearing in the sketch are small/big nanoparticles

non-aqueous media from gas phase, or even in solids. The process of nucleation is a ‘bottom-up’ approach in which atoms and/or molecules come together to form a solid. The process can be spontaneous, and it may be homogenous or heterogeneous nucleation. Homogenous nucleation is said to take place when it involves the nucleation around the constituent atoms or molecules of the resultant particles. Heterogeneous nucleation, on the other hand, can take place on a foreign particle like dust, deliberately adding seed particles, templates, or the walls of the container.

Nucleation can also occur by a cavitation process. In cavitation, if there are some bubbles (which can be formed deliberately, as discussed in a later section) in the solution which then collapse, the high local temperature and pressure thus generated may be sufficient to cause homogenous nucleation. It can be seen in curve A of Fig. 4.14 that fast nucleation takes place as the solute concentration approaches

super saturation. If the nuclei acquire atoms quickly by diffusion through solution thus reducing the solute concentration, particles of uniform size are formed in a relatively shorter time compared to aggregated particles in curve B, or *Ostwald-ripened* particles as in C.

In the Ostwald ripening process, the nucleation proceeds for a long time bringing in smaller and larger nuclei to co-exist when super-saturation region exists over a longer period and then the solute concentration decreases. Larger particles tend to grow at the expense of smaller ones, thus becoming even larger. This is because of the lowering of total surface energy. It should be remembered that the concentration of solutes and temperature of the solution would strongly affect the growth. Additionally, the crystalline structure, defects, favourable sites, etc. would strongly affect the final products.

As schematically illustrated in Fig 4.14, once the nuclei are formed they may take different routes—A, B, or C—depending on the growth conditions. The growth route depicted by curve A is a classical route suggested by LaMer, and is hence called as a LaMer diagram.

Nucleation process is controlled (except in bio-minerals, to be discussed in Chap. 5) thermodynamically. The size of a nucleus is determined by both the free energy change occurring during the formation of the solid (from a liquid), as well as the surface energy of the nucleus. A stable nucleus (of a critical radius r^*) needs to be formed so that it can grow into a larger stable particle. Let us refer to particles of radii smaller than r^* as embryos; the energy for such an embryo formation (ΔG_r) is given by:

$$\Delta G_r = \frac{4}{3}\pi r^3 \Delta G_v + 4\pi r^2 \gamma_{SL} \quad (4.9)$$

where r is the radius of the embryo, ΔG_v is the free energy change per unit volume between the liquid and solid, and γ_{SL} is the interfacial free energy of the liquid and solid.

Below the melting temperature (T_m) of the solid, ΔG_v is negative, whereas the surface free energy, or surface tension γ_{SL} is positive. The two energies compete with each other with increasing value of the embryo radius r . The nature of the resulting curve for ΔG_r is illustrated schematically in Fig. 4.15.

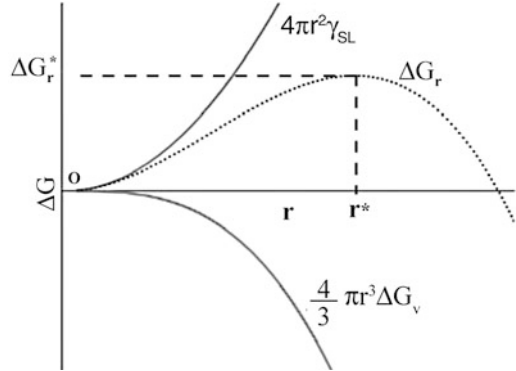
Differentiating Eq. (4.9) with respect to r and equating it to zero at $r = r^*$, we obtain:

$$4\pi r^{*2} \Delta G_v + 4\pi (2r^*) \gamma_{SL} = 0 \quad (4.10)$$

Hence,

$$r^* = \frac{-2\gamma_{SL}}{\Delta G_v} \quad (4.11)$$

After the critical size r^* , the free energy starts decreasing and the growth begins.

Fig. 4.15 Nucleation process

The energy ΔG_v depends on latent heat of fusion and the degree of undercooling. Undercooling is a result of faster cooling rate than required for the equilibrium cooling. Some finite time is required in any system so that atoms/molecules adjust themselves and acquire the position of minimum energy. The undercooling temperature is given as:

$$\Delta T = T - T_m \quad (4.12)$$

where T_m is equilibrium melting temperature, T – bulk temperature, and ΔH_f is the heat of fusion per unit volume.

It can be shown that

$$\Delta G_v = \frac{\Delta H_f \Delta T}{T_m} \quad (4.13)$$

Therefore,

$$r^* = \frac{-2\gamma_{SL} T_m}{\Delta H_f \Delta T} \quad (4.14)$$

If ΔG_v is the total free energy of the nucleus of radius r , number of clusters n_r can be obtained using

$$n_r = N e^{\left(\frac{-\Delta G_v}{kT} \right)} \quad (4.15)$$

where N is the total number of atoms and k is Boltzmann constant.

In case there is strain caused in the nucleus formation, additional term $\frac{4}{3}\pi r^3 \epsilon$ can be added to the Eq. (4.9). Here ϵ is the strain energy and is positive.

When the nucleation occurs on some foreign particle or surface (e.g. of the container wall or substrate), heterogeneous nucleation is said to occur. This lowers the energy necessary for the nucleation of a particle. Consequently the critical size

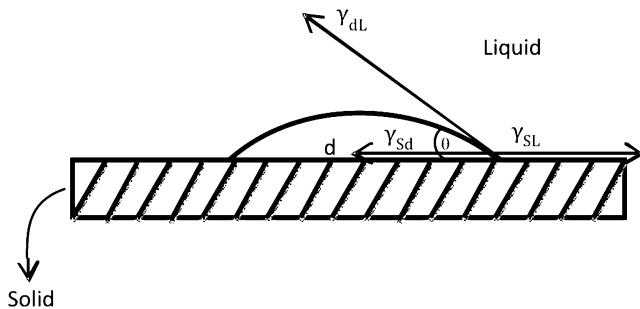


Fig. 4.16 Nucleation on a solid surface

or radius r^* for nucleation is smaller in heterogeneous nucleation than that for homogeneous nucleation. This can be understood as follows.

Consider a solid substrate on which a drop of liquid forms as illustrated in Fig 4.16. γ_{SD} , γ_{SL} and γ_{DL} are the interfacial surface tensions between substrate and liquid drop, surface and liquid and the drop and liquid respectively. It can be seen from Fig. 4.16 that:

$$\gamma_{DL} \cos \theta = \gamma_{SL} - \gamma_{SD} \quad (4.16)$$

Therefore,

$$\Delta G = V\Delta G_v + A_{dL}\gamma_{dL} - \pi r^2\gamma_{dL} \cos \theta \quad (4.17)$$

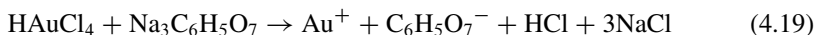
where V is the volume of the liquid drop; it may be noted that there is change in the volume free energy ΔG_v as discussed with reference to homogeneous nucleation. Hence the critical radius r^* which is given as

$$r^* = \frac{-2\gamma_{dL}}{\Delta G_v} \quad (4.18)$$

also changes. In general, once the nuclei with critical radii r^* are generated, stable nuclei and particle growth starts by addition of atoms or molecules from the solute.

4.4 Synthesis of Metal Nanoparticles by Colloidal Route

Colloidal metal nanoparticles are often synthesized by reduction of some metal salt or acid. For example highly stable gold particles can be obtained by reducing chloroauric acid (HAuCl_4) with tri sodium citrate ($\text{Na}_3\text{C}_6\text{H}_5\text{O}_7$). The reaction takes place as follows (Fig. 4.17):



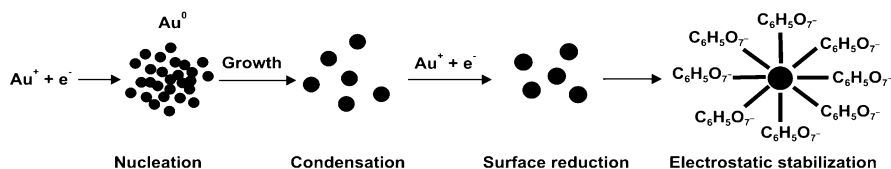


Fig. 4.17 Stabilization by electrochemical double layer formation

The reaction can be carried out in water using the set up shown in Fig. 4.13.

Metal gold nanoparticles exhibit intense red, magenta and other colours, depending upon the particle size. The size dependent optical properties of metal nanoparticles are discussed in Chap. 8. Gold nanoparticles discussed above are stabilized by repulsive Coulombic interactions. It is also possible to stabilize gold nanoparticles using thiol or some other capping molecules.

In a similar manner, silver, palladium, copper and other metal nanoparticles can be synthesized using appropriate precursors, temperature, pH, duration of synthesis etc. Particle size, size distribution and shape strongly depend on the reaction parameters and can be controlled to achieve desired results. It is also possible to synthesize alloy nanoparticles using appropriate precursors.

4.5 Synthesis of Semiconductor Nanoparticles by Colloidal Route

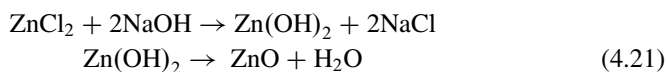
Compound semiconductor nanoparticles can be synthesized by wet chemical route using appropriate salts. Here we shall discuss some methods to obtain semiconductor nanoparticles.

Sulphide semiconductors like CdS and ZnS can be synthesized as nano-particles simply by coprecipitation. For example, to obtain ZnS nanoparticles any zinc salt like zinc sulphate ($ZnSO_4$), zinc chloride ($ZnCl_2$), zinc nitrate ($Zn(NO_3)_2$) or zinc acetate ($Zn[CH_3COO]_2$) can be dissolved in aqueous medium and Na_2S is added to the solution. (One can even dissolve H_2S gas in the Zn salt solution).

Following simple reaction



results to give solid colloidal particles of ZnS. To obtain zinc oxide particles one can use NaOH. Following reaction takes place.



Selenide particles can be obtained using appropriate selenium giving salt.

However all these nanoparticles need to be surface passivated as colloids formed in liquids have a tendency to coagulate or ripen due to attractive forces existing between them. The electrostatic and other repulsive forces may not be sufficient to keep them apart. However, as it was also discussed earlier, steric hindrance can be created by appropriately coating the particles to keep them apart. This is often known as ‘chemical capping’ and has become a widely used method in the synthesis of nanoparticles. Advantage with this chemical route is that, one can get stable particles of variety of materials not only in the solution, but even after drying off the liquid. One can even make thin films of the capped particles by spin coating or dip coating techniques. The coating, however, has to be stable and non-interactive with the particle itself except at the surface. Coatings may be a part of post-treatment or a part of the synthesis reaction to obtain nanoparticles. If it is a part of the synthesis reaction, the concentration of capping molecules can be used in two ways i.e. to control the size as well as to protect the particles from coagulation.

Chemical capping can be carried out at high or low temperature depending on the reactants. In high temperature reactions, cold organometallic reactants are injected in some solvent like trioctylphosphineoxide (TOPO) held at a temperature of 300–400 °C. For example when dimethyl cadmium $[\text{Cd}(\text{CH}_3)_2]$ and Se powder were injected in TOPO, CdSe nanoparticles capped with PO_4 groups were obtained. There are, however, other chemicals also which can be used as precursors to obtain high quality particles. It is possible to remove the aliquots at different intervals, as the reaction proceeds, to obtain the particles of different sizes. The particles with high quality and as narrow size distribution as <5 % have been achieved by this method.

Although, this is a very good route of synthesizing the nanoparticles, most of the organo-metallic compounds are prohibitively expensive. Besides they are also toxic and difficult to handle. Such synthesis should, therefore, be carried out only under expert guidance.

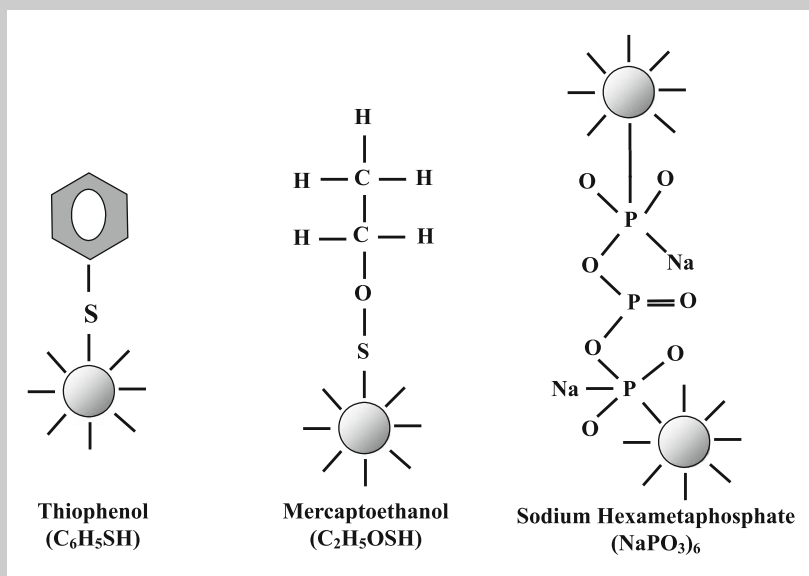
A wide range of metal oxides and other insulators with wide band gap can be synthesized by chemical precipitation method along with suitable surface passivant, if necessary. Some of the oxides and insulators may be stable and may not agglomerate or coalesce easily (Box 4.4).

Box 4.4: Chemical Capping of Nanoparticles

A variety of molecules can be used to cap the nanoparticles. For example capping of metal-sulphide nanoparticles by few organic and inorganic molecules is illustrated in Fig. 4.18.

In another variation of chemical capping method, reactions of inexpensive and non-toxic chemicals like metal chlorides, nitrates, acetates and inorganic salts are performed at moderate temperatures like 80–120 °C. In some cases, it is even possible to synthesize nanoparticles at room temperature. However the initial size distribution can be quite broad. The size distribution

(continued)

Box 4.4 (continued)**Fig. 4.18** Capping of nanoparticles by different molecules

can be narrowed down by a method known as ‘size selective precipitation’. For this one needs to use a proper pair of solvent-nonsolvent liquids. Some of the solvent-nonsolvent pairs are pyridine-hexane, chloroform-methanol or dimethyl sulfoxide-diethyl ether.

The nanoparticles are dispersed in a solvent so as to get an optically clear solution. Nonsolvent solution is then added so that flocculation occurs. Supernatant and flocculate are separated by centrifugation. Precipitate has larger particles and can be separated. Bigger particles, therefore, are first separated from smaller particles and again dispersed in the solvent solution. The process is continued until no change in the size distribution is observed by repeating the procedure.

Advantage of chemical capping method is mainly that nanoparticles are chemically stable over a long time. The thermal stability depends upon the capping molecules used. In most cases, where organic molecules are used, particles are stable upto about 200–250 °C and may find considerable range of applications. Another advantage with both the methods is that while synthesizing the nanoparticles, they can be doped with some metal ions so as to get fluorescent particles at relatively low temperatures. A wide range of semiconductor nanoparticles can be synthesized by this way and are found to be useful in many applications.

4.6 Langmuir-Blodgett (LB) Method

This technique to transfer organic overlayers at air-liquid interface onto solid substrates is known for nearly 70 years. The technique was developed by two scientists Langmuir and Blodgett and bears their names.

In this technique, one uses amphiphilic long chain molecules like that in fatty acids. An amphiphilic molecule (see Fig. 4.19) has a hydrophilic group (water loving) at one end and a hydrophobic group (water hating) at the other end. As an example consider the molecule of arachidic acid, which has a chemical formula $[\text{CH}_3(\text{CH}_2)_{16}\text{COOH}]$. There are many such long organic chains with general chemical formula $[\text{CH}_3(\text{CH}_2)_n\text{COOH}]$, where n is a positive integer. In this case $-\text{CH}_3$ is hydrophobic and $-\text{COOH}$ is hydrophilic in nature.

Usually molecules with $n > 14$ are candidates to form L-B films. This is necessary in order to keep hydrophobic and hydrophilic ends well separated from each other. Figure 4.20 illustrates few examples of different types of molecules which have successfully been used in L-B film deposition.

When such molecules are put in water, the molecules spread themselves on surface of water in such a way that their hydrophilic ends, often called as ‘heads’ are immersed in water, whereas the hydrophobic ends called as ‘tails’ remain in air. They are also surface active agents or surfactants. Surfactants are amphiphilic molecules that is an organic chain molecule in which at one end there is a polar, hydrophilic (water loving) and at the other a nonpolar, hydrophobic (water hating) group of atoms.

Using a movable barrier, it is possible to compress these molecules to come closer together to form a ‘monolayer’ and align the tails. It is, however, necessary that

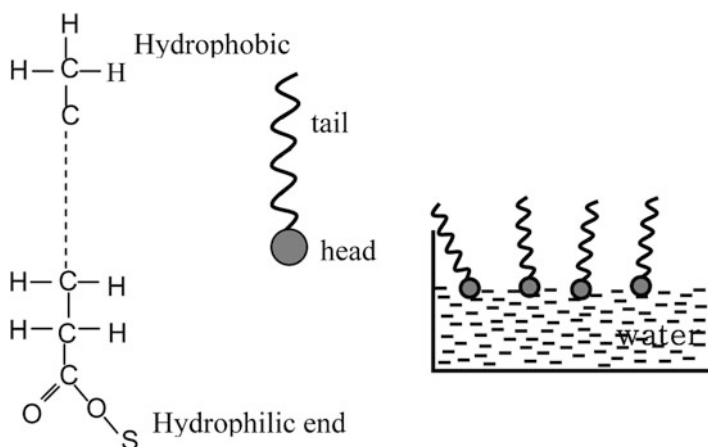


Fig. 4.19 Amphiphilic molecules with hydrophilic and hydrophobic ends to stay with head group immersed in water and tail group in air

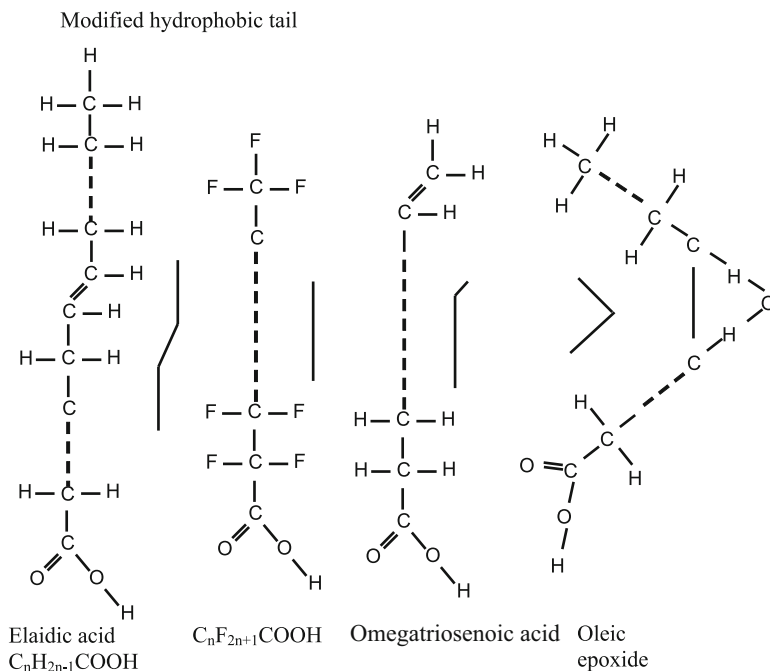


Fig. 4.20 A variety of organic molecules used for L-B thin film deposition

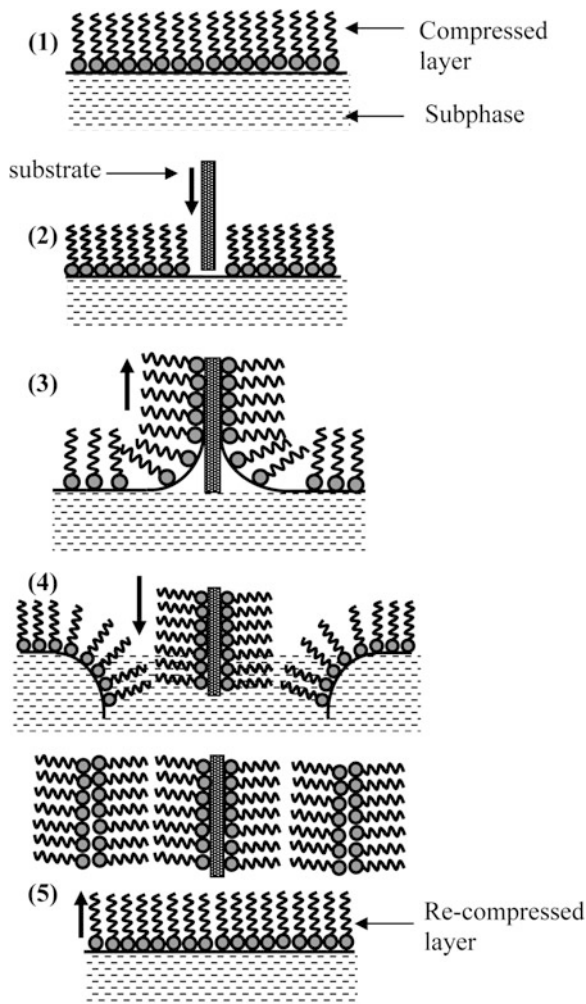
hydrophilic and hydrophobic ends are well separated with $n > 14$. Such monolayers are two dimensionally ordered and can be transferred on some suitable solid substrates like glass, silicon etc. This is done simply by dipping the solid substrate inside the liquid in which ordered organic molecular monolayer is already formed, as shown in Fig. 4.21.

Depending upon the nature of substrate material i.e. whether hydrophobic or hydrophilic, layers are transferred on the solid substrates. A glass slide when dipped in the solution becomes wet with water. Therefore, while it is withdrawn from the liquid the head groups can be easily attached to glass surface. As a result the whole monolayer gets transferred in a manner as if a carpet is pulled. Now the glass substrate has tail groups, which are hydrophobic on outer side.

Therefore as it is dipped in the liquid again, it acquires a second layer with tail-tail coming closer together and while it is pulled back to air, another monolayer of molecules with head-head groups coming together is pulled. The process of dipping-pulling the substrate can be repeated several times to obtain ordered multilayers of molecules. However to keep ordered layers available on water surface, it is necessary to keep constant pressure on the molecules.

In general there are three types of L-B films with different multilayer sequence, as shown in Fig. 4.22, identified. These are known as X, Y and Z type. Y types of films are most commonly found. Although the layers are ordered, there is only the

Fig. 4.21 Deposition of LB films by following steps:
 (1) A monolayer of amphiphilic molecules is formed.
 (2) A substrate is dipped in the liquid.
 (3) The substrate is pulled out, during which ordered molecules get attached to the substrate.
 (4) When the substrate is again dipped, molecules again get deposited as the substrate forming a second layer on the substrate.
 (5) As the substrate is again pulled out, a thin layer gets deposited. By repeating the procedure a large number of ordered layers can be transferred on a substrate



weak Van der Waals interaction between different layers. In this sense even with large number of layers present, the film preserves its two dimensional properties. Lengths of organic molecules as discussed above are typically 2–5 nm. Thus L-B films themselves are good examples of nanostructured materials.

It is also possible to obtain nanoparticles using LB technique. As shown in Fig. 4.23, a metal salt like CdCl_2 or ZnCl_2 is dissolved in water on surface of which a compressed uniform monolayer (single layer of molecules) of surfactant is spread. When H_2S gas is passed in the solution, CdS or ZnS nanoparticles of few tens of nanometers can be formed. Particles are monodispersed (almost one size) in size. If surfactant molecules are not present, uniform nanoparticles are not formed.

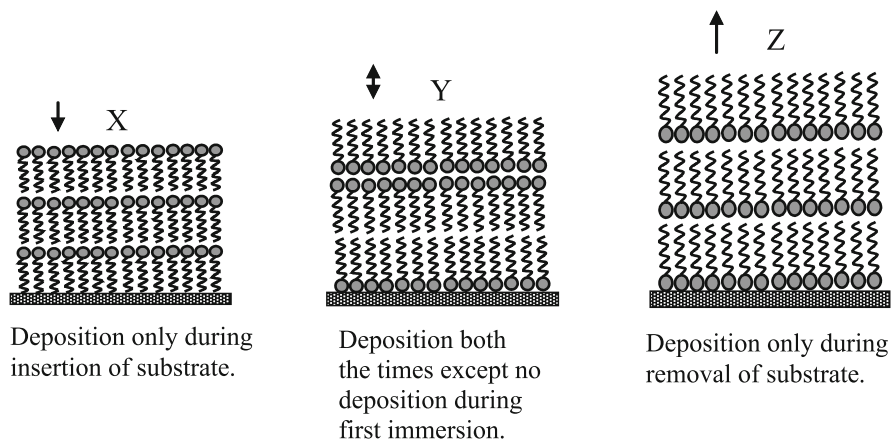


Fig. 4.22 X, Y and Z type L-B films

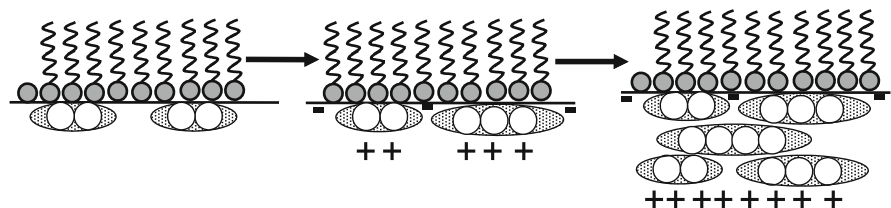


Fig. 4.23 Metal sulphide nanoparticles obtained under the monolayers at water-air interface. The time of hydrogen sulphide treatment increases from top to bottom

4.7 Microemulsions

Synthesis of nanoparticles in the cavities produced in microemulsion is a widely used method. Advantage of this method is the biocompatibility and biodegradability of synthesized materials. Biocompatibility is useful in novel applications like drug delivery of nanomaterials and biodegradability is useful to avoid environmental pollution.

Whenever two immiscible liquids are mechanically agitated or stirred together, they are known to form an 'emulsion'. The tendency of the liquids is such that the liquid in smaller quantity tries to form small droplets, coagulated droplets or layers so that they are all separated from the rest of the liquid in large quantity (for example droplets of fat in milk). The droplet sizes in emulsions are usually larger than 100 nm upto even few millimetres. Emulsions are usually turbid in appearance.

On the other hand there is another class of immiscible liquids, known as microemulsions which are transparent and the droplets are in the range of ~1–100 nm. This is the size needed for the synthesis of nanomaterials (Boxes 4.5 and 4.6).

Box 4.5: Amphiphilic Molecules in Liquids

If amphiphilic molecules are spread in an aqueous solution, they try to stay at air-solution interface with hydrophobic groups in air and hydrophilic groups in the solution (see Fig. 4.24a). Such molecules are known as surfactants (surface active agents). This is similar to what was discussed in L-B film synthesis.

Consider now a situation in which a hydrocarbon molecule solution is put in an aqueous medium. As shown in Fig. 4.24b, the hydrocarbon solution itself would be separated from aqueous solution and float on it. When surfactant molecules are mixed in large quantity in aqueous solution, they would try to form what are known as 'micelles' and 'inverse-micelles' when aqueous solution is mixed in oil. In micelles, the head groups float in water and tails are inside, whereas tails point outwards in case of inverse micelles.

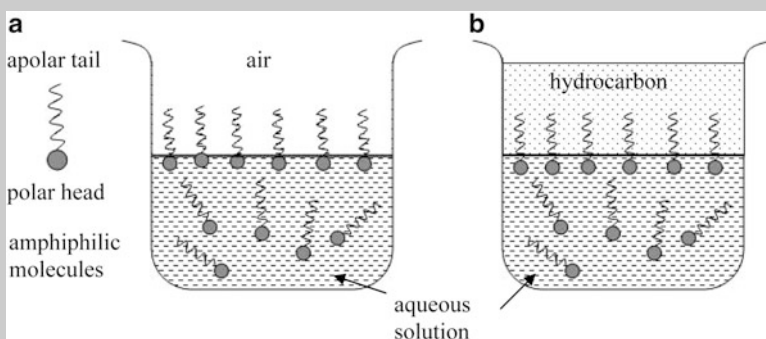
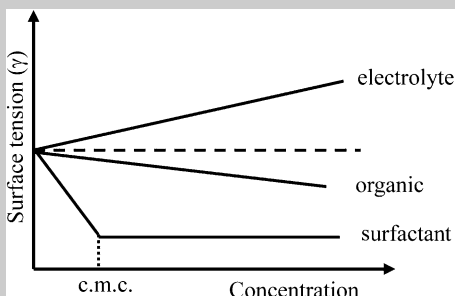


Fig. 4.24 Amphiphilic molecules in aqueous solutions

Box 4.6: Surface Tension of Liquids

Surface tension of a liquid can change if some electrolyte, organic or surfactant solutes are added. General behaviour is shown in Fig. 4.25.

Fig. 4.25 For surfactant molecules γ decreases rapidly upto certain concentration known as critical micelle concentration (c.m.c.)



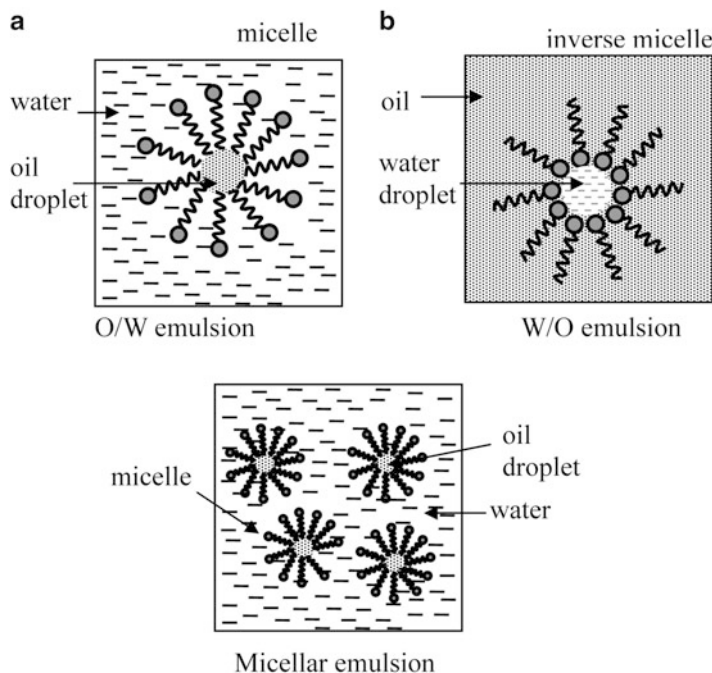
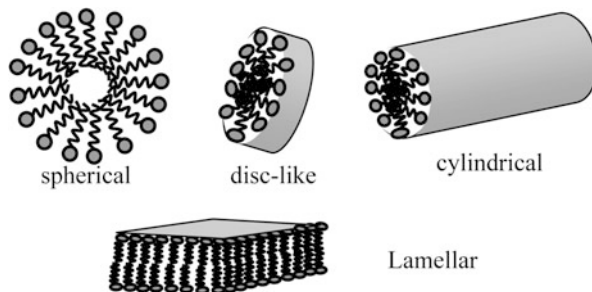


Fig. 4.26 Formation of micelles and inverse micelles

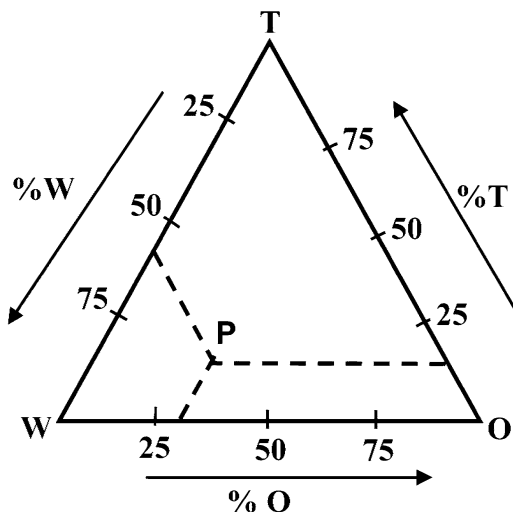
Fig. 4.27 Different shapes of micelles



Microemulsions are stabilized using surfactants (surface stabilized active agents). When an organic liquid or oil (O), water (W) and surfactant (T) are mixed together, under some critical concentration, ‘micelles’ or ‘inverse micelles’ are formed, depending upon the concentrations of water and organic liquid. As shown in Fig. 4.26, micelles are formed with excess water and inverse micelles are formed in excess of organic liquid or oil.

As shown in Fig. 4.26, micelles have the head groups floating in water, whereas tails and tail group filling the cavity along with organic liquid inside. Reverse is the case for inverse micelles. They can be formed in various shapes as well. In Fig. 4.27 different shapes taken by micelles under different synthesis conditions are illustrated.

Fig. 4.28 Ternary phase diagram of water (W), oil (O) and surfactant (T) mixture



The ratio of water (W), oil (O) and surfactant (T) is important to decide which type of micelle will be formed and can be represented in a ternary phase diagram, using a triangle as shown in Fig. 4.28.

Composition can be determined by drawing lines parallel to all three sides of the triangle as shown in Fig. 4.28. Here point *P* denotes 60 % water, 26 % oil and 14 % surfactant.

A modified phase diagram known as ‘Winsor Diagram’ (see Fig. 4.29a) further denotes the types of phases formed.

There is also another type of phase diagram which shows further details as illustrated in Fig. 4.29b.

The critical micelle concentration (CMC) depends upon all W, O and T concentrations as is evident from above diagram. The effect of T is to reduce the surface tension of water dramatically below CMC and remain constant above it, as the organic solvent concentration keeps on increasing. Organic solutes also reduce the surface tension to some small extent. If there are any electrolytes used, they slightly increase the surface tension.

There are four types of surfactants in general:

1. Cationic – For example CTAB, $\text{C}_{16}\text{H}_{33}\text{N}(\text{CH}_3)_3^+\text{Br}^-$
2. Anionic – For example sulphonated compounds with general formula $\text{R}-\text{SO}_3^-\text{Na}^+$
where R is $\text{C}_n\text{H}_{2n+1}$.
3. Nonionic – For example $\text{R}-(\text{CH}_2-\text{CH}_2-\text{O})_{20}-\text{H}$
4. Amphoteric – Some properties are similar to ionic and some to nonionic surfactants as in betaines.

A large number of nanoparticles (metals, semiconductors and insulators) of cobalt, copper, CaCO_3 , BaSO_4 , CdS or ZnS have been synthesized using

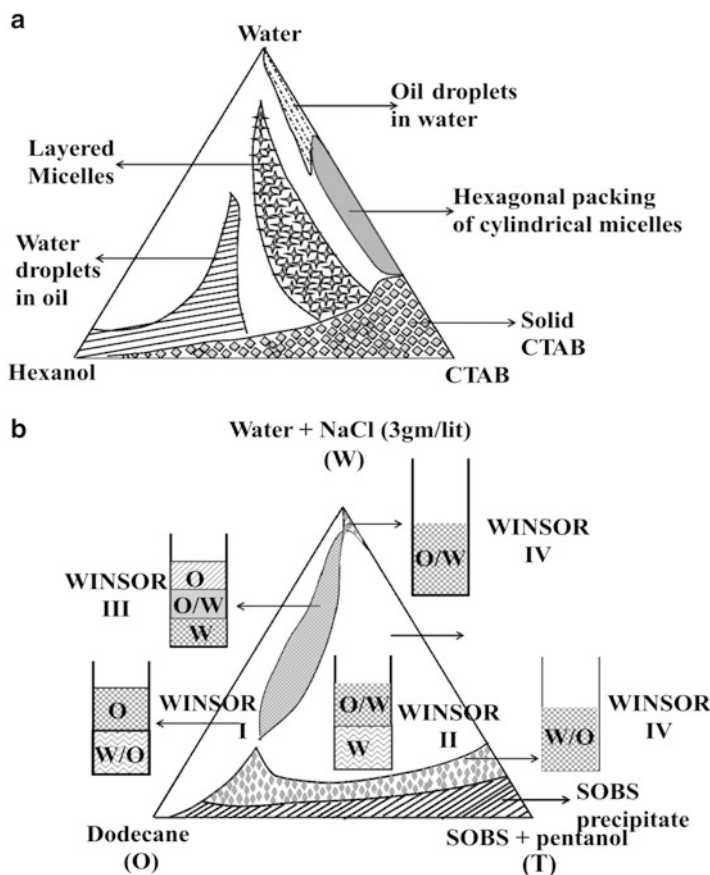


Fig. 4.29 Winsor phase diagram with different phases (a) water-hexanol-centrilytrimethylammoniumbromide (CTAB) and (b) water and NaCl-dodecane-paraocetylbenzene sodium sulphate (SOBS) and pentanol

microemulsions or inverse micelles. As an example, consider the synthesis of cobalt nanoparticles. A reverse micelle solution of water and oil can be stabilized using a monolayer of surfactant like sodium bis (2-ethylhexyl)sulfosuccinate or Na(AOT). The droplet diameter is controlled simply by controlling the amount of water. Two micelle solutions having same diameter of droplets can be formed. Thus one solution should have Co(AOT)_2 i.e. cobalt bis(2-ethylhexyl)sulfosuccinate and the other should have sodium tetrahydroborate (NaBH_4 i.e. sodium borohydride). When two solutions are mixed together the resultant solution appears clear but the colour changes from pink to black. One can find by electron microscopy or some other analysis that cobalt nanoparticles are formed.

4.8 Sol-Gel Method

As the name suggests sol gel involves two types of materials or components, 'sol' and 'gel'. Sol gels are known since the time when M. Ebelman synthesized them in 1845. However it is only since the last one or two decades that considerable interest in it, both in scientific and industrial field, has generated due to realization of the several advantages one gets as compared to some other techniques. First of all sol gel formation is usually a low temperature process. This means less energy consumption and less pollution too. It is therefore not surprising that in the nuclear fuel synthesis it is a desired process. Although sol-gel process generates highly pure, well controlled ceramics it competes with other processes like CVD or metallo-organic vapours derived ceramics. The choice of course depends upon the product of interest, its size, instrumentation available and ease of processing. In some cases sol-gel can be an economical route, provided precursors are not very expensive. Some of the benefits like getting unique materials such as aerogels, zeolites, and ordered porous solids by organic-inorganic hybridization are unique to sol-gel process. It is also possible to synthesize nanoparticles, nanorods or nanotubes using sol-gel technique.

Sols are solid particles in a liquid (see Fig. 4.30). They are thus a subclass of colloids. Gels are nothing but a continuous network of particles with pores filled with liquid (or polymers containing liquid). A sol gel process involves formation of 'sols' in a liquid and then connecting the sol particles (or some subunits capable of forming a porous network) to form a network. By evaporating the liquid, it is possible to obtain powders, thin films or even monolithic solid. Sol gel method is particularly useful to synthesize ceramics or metal oxides although sulphides, borides and nitrides also are possible.

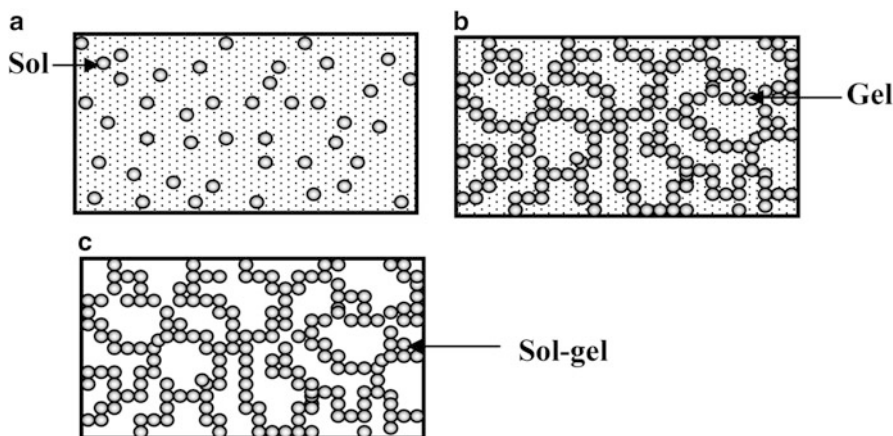


Fig. 4.30 Sol (a), gel (b) and sol-gel (c) monolithic solid

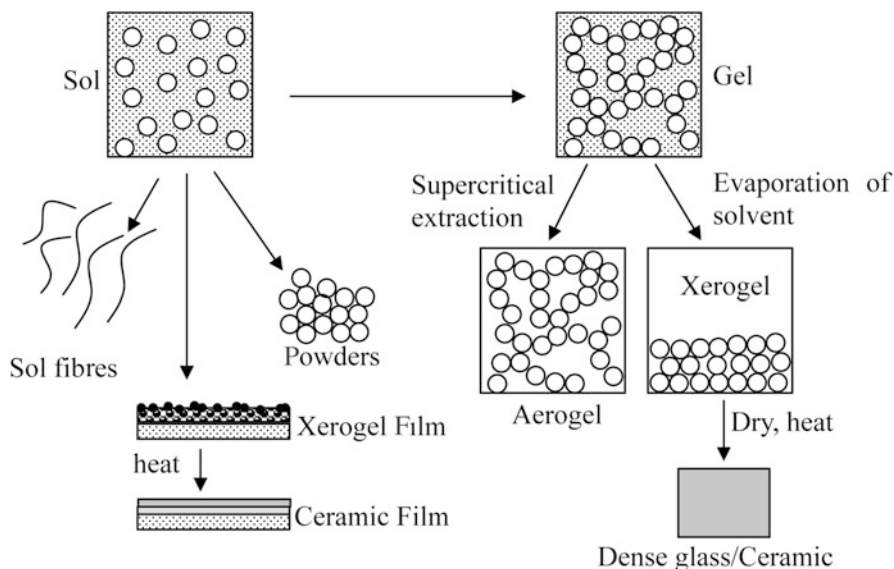


Fig. 4.31 Sol-gel options

Synthesis of sol-gel in general involves hydrolysis of precursors, condensation followed by polycondensation to form particles, gelation and drying process by various routes as shown in Fig. 4.31. Precursors (starting chemicals) are to be chosen so that they have a tendency to form gels. Both alkoxides or metal salts can be used. Alkoxides have a general formula $M(\text{ROH})_n$, where M is a cation and R is an alcohol group, n is the number of (ROH) groups with each cation. For example (ROH) can be methanol (CH_3OH), ethanol ($\text{C}_2\text{H}_5\text{OH}$), propanol ($\text{C}_3\text{H}_7\text{OH}$) etc. bonded to a cation like Al or Si . Salts are denoted as MX , in which M is a cation and X is an anion like in CdCl_2 , Cd^+ is a cation and Cl^- is an anion.

Although it is not mandatory that only oxides be formed by a sol-gel process, often oxide ceramics are best synthesized by a sol-gel route. For example in silica, SiO_4 group with Si at the centre and four oxygen atoms at the apexes of tetrahedron are very ideal for forming sols with interconnectivity through the corners of tetrahedrons, creating some cavities or pores (See Fig. 3.4 in Chap. 3).

Due to its higher electronegativity as compared to metal cations, Si is less susceptible to nucleophilic attacks.

By polycondensation process (i.e. many hydrolyzed units coming together by removal of some atoms from small molecules like OH), sols are nucleated and ultimately sol-gel is formed.

We shall discuss two special types of sol-gel materials viz. zeolites and aerogels in Chap. 11. We shall also see that there are recent methods to combine microemulsion method with sol-gel method to produce some novel materials.

4.9 Hydrothermal Synthesis

This synthesis method is useful to make a large scale production of nano to micro size particles. In this technique adequate chemical precursors are dissolved in water and placed in vessel made of steel or any other suitable metal which can withstand high temperature typically upto 300 °C and high pressure above 100 bars. The vessel, known as *autoclave*, is usually provided with temperature and pressure control as well as measuring gauges as illustrated in Fig. 4.32.

It is a very old technique, probably first used by the German scientist Robert Bunsen, way back in 1839 to synthesize crystals of strontium and barium carbonates. He used a thick glass tube and used temperature above 200 °C and pressure more than 100 bars. The technique was later used mostly by geologists and has become popular amongst nanotechnologists due to the advantages like large yield and novel shapes and sizes that can be obtained using this technique.

The technique becomes useful when it is difficult to dissolve the precursors at low temperatures or room temperature. It is also advantageous to use the technique to grow nanoparticles if the material has a high vapour pressure near its melting point or crystalline phases are not stable at melting point. The uniformity of shapes and sizes of the nanoparticles also can be achieved by this technique. Various oxide, sulphide, carbonate and tungstate nanoparticles have been synthesized by the hydrothermal synthesis.

Another variation of hydrothermal synthesis technique is known as *forced hydrolysis*. In this case usually dilute solutions (10^{-2} to 10^{-4} M) of inorganic metal salts are used and hydrolysis is carried out at rather higher temperatures than 150 °C.

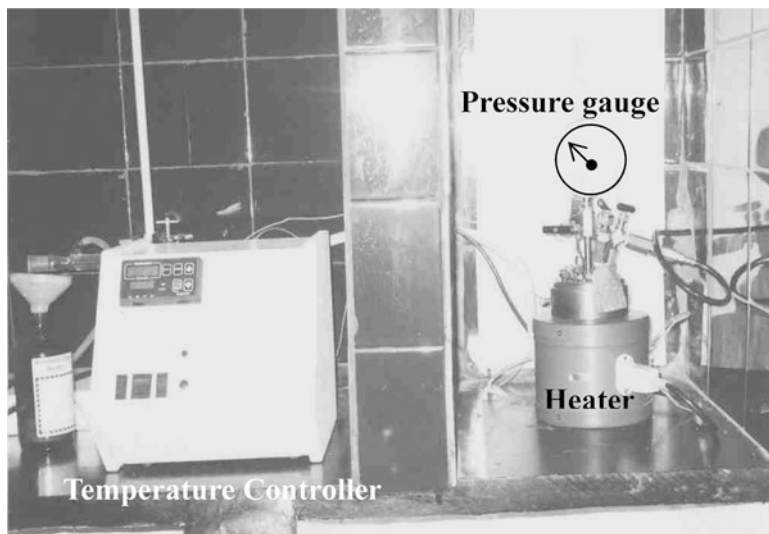


Fig. 4.32 Photograph of an autoclave set up

4.10 Sonochemical Synthesis

In this technique the reactivity of the precursors is enhanced by taking the advantage that large amount of energy can be released when bubbles burst in a liquid. Bubbles are formed (see Fig. 4.33) by using ultrasonic waves in a frequency range of ~ 20 kHz–2 MHz. It can be considered as an alternative method to enhance the chemical reactions in liquids by heating and/or pressurizing.

Although it is not well understood yet as to how nanoparticles can be synthesized using sonochemical method, it is well agreed upon that creation, growth and collapse of bubbles in liquids is most important pathway of causing the reactions. The ultrasonic waves while passing through the liquids create very small bubbles which keep on growing until they reach a critical size and then burst, releasing very high energy to locally reach a temperature of $\sim 5,000$ °C and a pressure of few hundred times that of atmospheric pressure. For the reaction to occur in the gas phase the solute in the liquid should diffuse to the growing bubble. The reaction can also occur in the liquid phase at exploding bubble where in the interfacial region surrounding the bubble (~ 200 nm distance) the temperature as high as $\sim 1,600$ °C can be reached. Typically the size of a bubble can be from ten to few tens of microns. Careful use of the solvents and solutes are very important. Non-volatile liquids would prevent formation of bubbles, which is desired, as only reactants should find their place inside the bubble in the form of vapour. Solvents should be inert and stable to ultrasonic irradiation. Interestingly, the cooling rates also can be as high as 10^{11} °C or more. Such a high cooling rate gives rise to amorphous nanoparticles as the atoms do not have sufficient time to reorganize. However amorphous particles can be more active than the crystalline particles of the same size and same material. This can be useful in some fields like catalysis. Various nanoparticles like ZnS, CeO_2 and WO_3 have been synthesized using sonochemical method.

The bubbles after their initiation, grow as shown in the schematic diagram (Fig. 4.33). Bubbles expand progressively in the region of rarefaction and contract

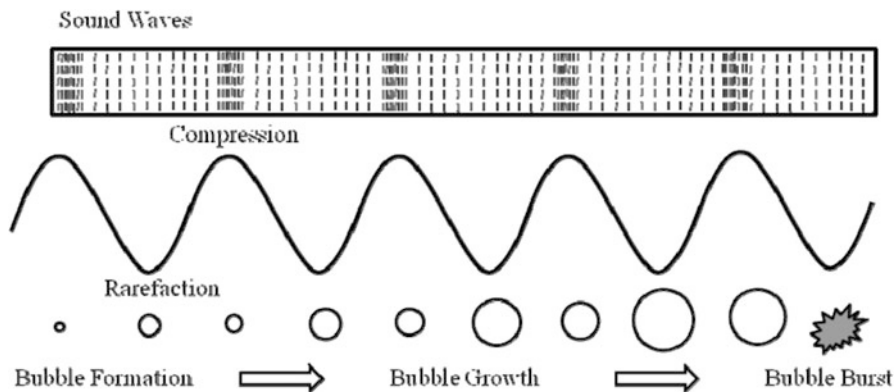


Fig. 4.33 Ultra-sound waves form compressions and rarefactions, shown as a sine wave

at compressions but quickly attain a micrometre size and ultimately burst releasing huge energy. Temperature rises to $\sim 5,000^\circ\text{C}$ and pressure is more than 100 times atmospheric pressure.

4.11 Microwave Synthesis

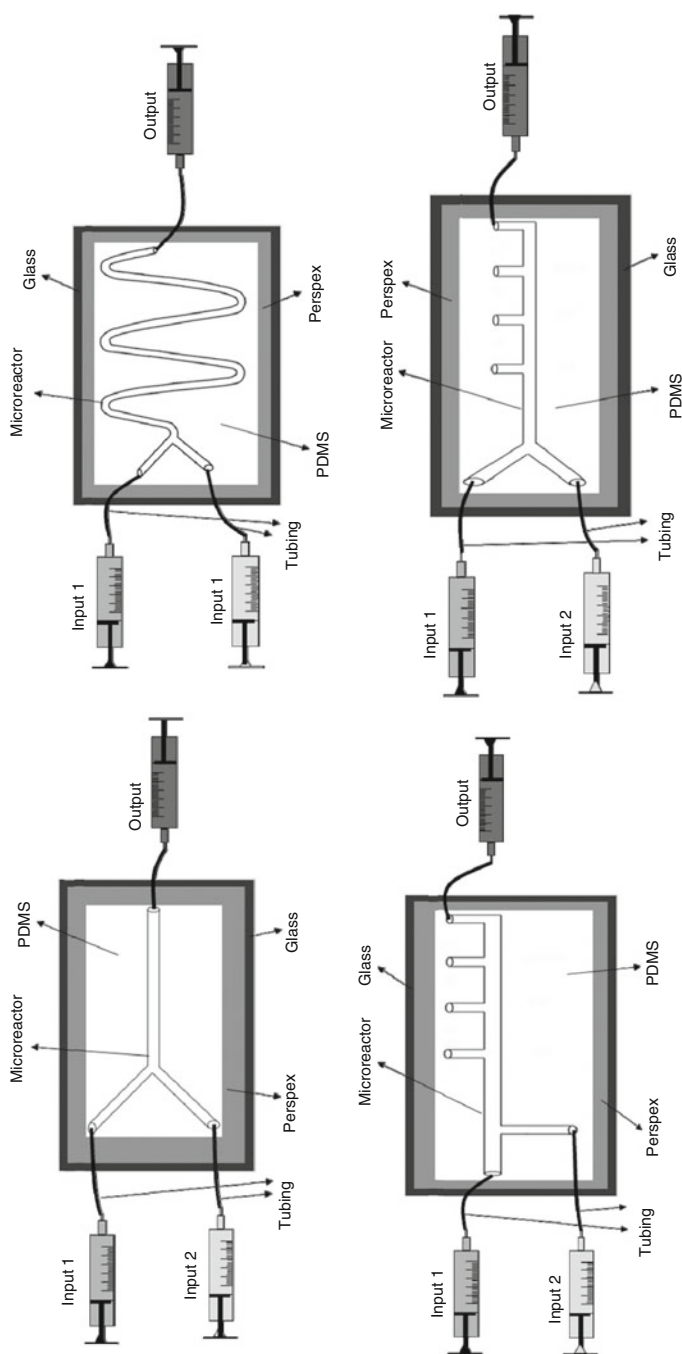
Use of microwave ovens for heating or cooking food is very common. Its entry in the scientific laboratories took place around 1986 when it was demonstrated by some scientists that rapid, large scale and uniform synthesis of materials is possible even using a domestic microwave oven. Kitchen microwave ovens, however, are not any more considered for the controlled chemical synthesis as there is no good control on stirring, temperature or power which is necessary for scientific equipment. However, due to its several advantages microwave synthesis has been used in research laboratories with equipment capable of controlling various parameters.

Microwaves are a part of electromagnetic spectrum with very long wavelength and frequencies in the range of $\sim 300\text{--}300,000$ MHz. However, only certain frequencies are reserved for domestic and other equipment, rest being used for communication purposes. A microwave has oscillating electric and magnetic fields associated with it which produces nodes and antinodes and correspondingly hot and cold spots in a vessel. This would lead to non-uniformities. Therefore sometimes only a 'single mode' is used in which length of the cavity (or reaction vessel) is equal to a single wave only. In a microwave apparatus heating is caused only due to a process in which molecules in a solution try to orient their dipoles appropriately to align themselves in the direction of the electric field. In the process they too oscillate, generating heat in the medium. Advantage in this case is that the external energy is not wasted in heating of the vessel.

Several types of oxide, sulphide and other nanoparticles have been synthesized to obtain various shapes and sizes. The reaction time is greatly reduced and the products are uniform in size and shape.

4.12 Synthesis Using Micro-reactor or Lab-On-Chip

Micro-reactor or lab-on-chip is a relatively new method of synthesizing nanoparticles in small quantities. Basically, very narrow channels (less than about $100\text{ }\mu\text{m}$ upto few tens of nm in width and depth) are made in some suitable substrates like glass, silicon or polymers like poly-dimethylsiloxane (PDMS) using lithography techniques. Similar to an electronic circuit in a semiconductor chip, these channels make some circuit where the fluids can mix. There can be some mixer regions where stirring takes place with the help of magnetic or some other actuation. There can be some valves to control the flow of liquids. The channels can be of short or long length depending upon the requirements and are designed to suit a particular

**Fig. 4.34** Schematic of microreactor set up

requirement. The size of the whole reactor can be as small as $\sim 10 \text{ cm}^2$. The liquids may be aqueous or non-aqueous and suitable reactor will have to be chosen according to the reactions to be carried out. It is also possible to heat some of the reactors to enhance the rates of reaction. The liquids should not react with the reactor material or percolate inside its body. The liquids are injected inside the channels using syringe pumps. Figure 4.34 illustrates few simple designs of a microreactor set up without complicated valves, mixers and other components.

Advantage of synthesizing nanomaterials in microreactors is that reactions can be carried out in a very short time using small amounts of reactants. This is advantageous when expensive or toxic reactants are to be used. Due to small amounts involved, the risk of pollution is minimized. Short synthesis time enables many reactions to be carried out in a short time. Therefore, optimization of reaction parameters can be done very quickly. However, due to small channel size, the fluid flow in channels and in large glass flask reactor may differ. However, using parallel processing one can increase the quantity of the product to be obtained in microreactors. One, however, has to be careful that if the particles grow to large size they may clog the channels. Cleaning of the channels also can pose problems. However disposable polymer-based inexpensive microreactors can be a good solution to avoid the cross contaminations.

There are many reports now which show that TiO_2 , ZnS, CdSe, Au, Ag etc. nanoparticles with narrow size distribution have been achieved using microreactors. Moreover doping of nanoparticles also is possible in microreactors. In short, whatever the synthesis is carried out in a chemistry laboratory can be carried out in the lab-on-chip.

Further Reading

- D.H. Everett, *Basic Principles of Colloid Science* (Royal Society of Chemistry, London, 1988)
C. Jeffrey Brinker, J.W. Scherer, *Sol-Gel Science: The Physics and Chemistry of Sol-Gel Processing* (Academic, Boston, 1990)
R.W. Jones, *Fundamental Principles of Sol-Gel Technology* (The Institute of Metals, Brookfield, 1989)
S.K. Kulkarni, Doped II-VI semiconductor nanoparticles. *Encycl. Nanosci. Nanostruct. Mater.* **2**, 537 (2004)
E.J. Matijevic, Monodispersed metal (hydrous) oxides—A fascinating field of colloid science. *Acc. Chem. Res.* **14**, 22–29 (1981)
A.C. Pierre, *Introduction to Sol-Gel Processing* (Springer, 1998)
U. Schubert, N. Hüsing, *Synthesis of Inorganic Materials*, 2nd edn. (Wiley-VCH, 2005)

Chapter 5

Synthesis of Nanomaterials—III

(Biological Methods)

5.1 Introduction

In his, very famous speech delivered in 1959, before the scientists of American Physical Society, Nobel Laureate Richard Feynman asked the scientists to derive the inspiration from *Mother Nature* to make the things smaller and see the advantages of making things smaller. Indeed the biological world, animal kingdom and plants make optimum use of materials and space. Nature indeed makes use of small spaces and corresponding confinement to synthesize inorganic materials or minerals abundantly found in earth's crust. It uses insoluble, complex organic molecules as a reactor in which nucleation and growth of complex and hierarchical structures of inorganic materials takes place by reactions of organic soluble molecules. When we think of biological world, we normally think of delicate, temperature sensitive, carbonaceous or organic materials like leaves, roots, stems, cells, tissues and skin. Inorganic materials are also produced in biological systems. Bones, teeth, shells, nanomagnets inside the bodies of some bacteria and birds are some examples. Inorganic materials inside organic matter or organisms are known as biocomposites or biominerals. A variety of mechanically strong or weak, rigid or flexible, porous or nonporous, thick or thin materials either organic or inorganic in small or large quantities are abundantly produced in contact with live cells. These materials exhibit a wide variety in their functions like providing support to body, allow body movements and in general carry out various essential body functions. Table 5.1 gives a list of some biominerals produced or observed in the biological systems. The list is not at all complete and many more minerals are observed to be synthesized in nature under different environmental conditions.

Many of the materials synthesized by microorganisms, animals and plants in nature can indeed be synthesized using them in laboratories even on large scale. This is considered to be a very attractive possibility so as to have *eco-friendly* or so-called *green synthesis*. Further the ability of bio-world to synthesize materials of variety of shapes and sizes is quite unique and often difficult to mimic, though not impossible.

Table 5.1 Some of the biominerals produced by the biological world

Minerals	Formula
Arsanate	
Orpiment	As_2S_3
Carbonates	
Calcite	CaCO_3
Mg-Calcite	$(\text{Mg}_x\text{Ca}_{1-x})\text{CO}_3$
Monohydrocalcite	$\text{CaCO}_3 \cdot \text{H}_2\text{O}$
Hydrocerrussite	$\text{Pb}_3(\text{CO}_3)_2(\text{OH})_2$
Amorphous calcium carbonate (various forms)	$\text{CaCO}_3 \cdot \text{H}_2\text{O}$
Chlorides	
Athcamite	$\text{Cu}_2\text{Cl}(\text{OH})_3$
Fluorides	
Fluorite	CaF_2
Hieratite	K_2SiF_6
Hydrated silica	
Amorphous silica	$\text{SiO}_2 \cdot n\text{H}_2\text{O}$
Hydroxides and hydrous oxides	
Goethite	$\alpha\text{-FeOOH}$
Lepidocrocite	$\gamma\text{-FeOOH}$
Ferrihydrite	$5\text{Fe}_2\text{O}_3 \cdot 9\text{H}_2\text{O}$
Brinessite	$\text{Na}_4\text{Mn}_{14}\text{O}_{27} \cdot 9\text{H}_2\text{O}$
Metals	
Sulfur	S
Organic crystals	
Uric acid	$\text{C}_5\text{H}_4\text{N}_4\text{O}_3$
Cu tartarate	$\text{C}_4\text{H}_4\text{CaO}_6$
Guanine	$\text{C}_5\text{H}_3(\text{NH}_2)\text{N}_4\text{O}$
Oxides	
Magnetite	Fe_3O_4
Amorphous iron oxide	Fe_2O_3
Amorphous manganese oxide	Mn_3O_4
Phosphates	
Octacalcium phosphate	$\text{Cu}_8\text{H}_2(\text{PO}_4)_6$
Francolite	$\text{Ca}_{10}(\text{PO}_4)\text{F}_2$
Virianite	$\text{Fe}^{2+}_3(\text{PO}_4)_2 \cdot 8\text{H}_2\text{O}$
Amorphous calcium pyrophosphate	$\text{Ca}_2\text{P}_2\text{O}_7 \cdot 2\text{H}_2\text{O}$
Sulphates	
Gypsum	$\text{CaSO}_4 \cdot 2\text{H}_2\text{O}$
Barite	BaSO_4
Jarosite	$\text{KFe}_3^{+3}(\text{SO}_4)_2(\text{OH})_6$

(continued)

Table 5.1 (continued)

Minerals	Formula
Sulphides	
Pyrite	FeS_2
Hydrotroilite	$\text{FeS} \cdot n\text{H}_2\text{O}$
Sphalerite	ZnS
Wurtzite	ZnS
Galena	PbS
Acanthite	As_2S

Many processes taking place inside the living cells are not understood well yet. Interestingly shape, texture and variety of structures observed in the biological systems, taking place by atomic (ionic) molecular interactions and diffusion are not controlled by thermodynamics but are kinetically driven. Mineral nucleation is controlled by interfacial energies and growth depends upon the control through passivation of the material surface. Solution composition, ionic activities, their transport, availability of suitable growth sites, ion accumulation and transportation are all governing factors in the nucleation and growth of particles. Starting with nanoparticles, hierarchical structures are produced by self assembly which will be discussed in the next chapter. Such structures using nanomaterials also is a topic of great interest and will be discussed under applications of nanoscience to technology. In this chapter we shall restrict only to synthesis of nanomaterials.

Driven by the motivation of understanding biological systems as well as mimicking the nanosynthesis by nature's way, scientists have been using the methods by which inorganic materials are synthesized by using biomaterials like enzymes, DNA and membranes. A variety of metal, semiconductor and insulator nanoparticles or their assemblies (one dimensional, two dimensional and even three dimensional) have been made. There are also attempts to make biocompatible, bioactive or biopassive materials, specially for different medical applications like body implants, drug delivery, cancer therapy and so on. Here we shall discuss only the synthesis part using biological methods.

Synthesis of nanomaterials using biological ingredients can be roughly divided into following three types.

1. Use of microorganisms like fungi, yeasts (eukaryotes) or bacteria, actinomycetes (prokaryotes). Eukaryotes are higher organisms with proper nucleus in their cells and prokaryotes are lower organisms without any nucleus.
2. Use of plant extracts or enzymes.
3. Use of templates like DNA, membranes, viruses and diatoms.

In the following sections we shall outline these methods with some examples (Boxes 5.1, 5.2, and 5.3).

Box 5.1: Biomaterials

Carbon, hydrogen, oxygen and nitrogen are the major constituents (~95 % by weight) of any living cell, animal or plant. Some other elements like iron, copper, manganese, zinc, selenium etc. are also found to be present in varying amounts (sometimes only as trace elements) in biosystems. All the elements together form various body parts and have to carry out different functions to execute life cycle activities. This involves a rich variety of molecules or materials from nano to macro scale. They are all synthesized in bioenvironment and can self assemble or dissociate; they can recognize some specific molecules or bodies due to their charges or shapes, bind weakly or strongly with other molecules.

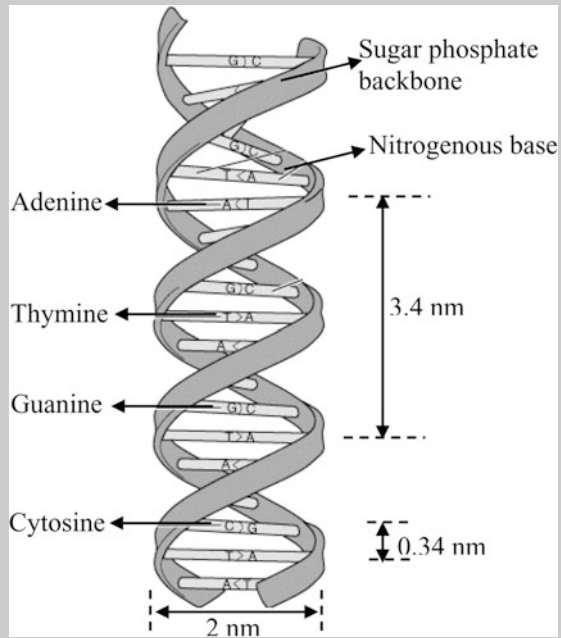
Building blocks of living systems can be divided into four classes viz. very small molecules, proteins, nucleic acids and carbohydrates. Small molecules like water (H_2O), oxygen (O_2), nitric oxide (NO) or even slightly bigger molecules like sugars, enzymes, acids are quite important in biosystems. They may be present as energy sources, purifiers or excretions in various parts.

There is a large number of different types of proteins. Proteins form a major part of biological systems. Nails, hair etc. are made up of proteins. They are responsible for numerous life activities.

DNA and RNA form nucleic acids. They are major constituents of proteins. They are genetically coded and are, therefore, most important as nature's messengers of life from one generation to the next. DNA is a long chain molecule composed of a backbone of sugars and phosphates. There are four different small planer molecules viz. adenine (A), thymine (T), guanine (G) and cytosine (C). They are usually denoted by the capital letters in the brackets. These molecules have a peculiarity that A can pair only with C and G can pair only with T. Any strand can therefore be coded with a long or short sequence of molecules like ACGGT, AGCTT and even keep on repeating the sequence. This gives a unique identification to DNA. Two complementary strands of DNA bind to form a double helix as illustrated in Fig. 5.1. Length of a DNA can be as short as few nanometres and as long as few micrometres. There is a large variety of shapes too. There are circular, branched, T-shaped, Y-shaped DNA molecules.

Carbohydrates are macromolecules as important as proteins for living organisms. They are basically sugars or polysaccharides of long chain lengths.

(continued)

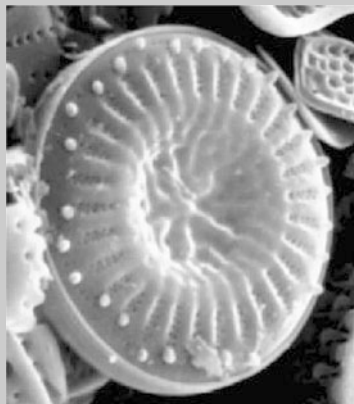
Box 5.1 (continued)**Fig. 5.1** Double helix of DNA**Box 5.2: Diatoms**

Diatoms are beautiful algae found in marine or any moist environment. They are unicellular, forming amorphous silica shell with living part inside. Diatoms can be $\sim 20\text{--}200\ \mu\text{m}$ long. Silicic acid $[\text{Si}(\text{OH})_4]$ is absorbed inside the cell, which is bonded to cofactor so that polycondensation is avoided. Golgi bodies in cell store silicic acid. Silica transport vesicles. Vesicles are phospholipid bilayers which can trap some solution or water and combine to form silica deposition vesicles (SDV), mineralized part of organism. Silicic acid condense in SDV and grow rapidly. It forms bottom part of new cell wall. Forms of silica shells are decided by genetic factors and are influenced by the organic patterns inside. A large variety of diatoms exist. Dead diatoms are deposited on ocean beds (Fig. 5.2).

(continued)

Box 5.2 (continued)

Fig. 5.2 Scanning Electron Microscope (SEM) image of a diatom

**Box 5.3: Molecular Recognition**

Response of living systems to surrounding depends upon molecular recognition. For example, smell, allergy, diseases etc. are due to molecular recognition. Molecules recognize each other due to their opposite charges or fitting shapes. Insects like ants even attract others by excretion of some molecules like pheromones.

5.2 Synthesis Using Microorganisms

Microorganisms are the organisms that can be observed under a microscope, such as bacteria, fungi or yeasts. Some bacteria are quite useful and are used in the processing of cheese, curds, bread, alcohol and vaccines. Some are harmful and responsible for spoiling food or causing diseases.

Microorganisms are capable of interacting with metals coming in contact with them through their cells and form nanoparticles. In Fig. 5.3a prokaryotic and (Fig. 5.3b) eukaryotic cells are illustrated respectively. The cell-metal interactions are in general difficult to understand due to complexity of cells themselves. It is, however, well known that certain microorganisms are capable of separating metal ions. This is widely used to either recover precious metals or detoxify water. However use of microorganisms in deliberate synthesis of metal, semiconductor or insulator nanoparticles is relatively new area and will be discussed now with some examples.

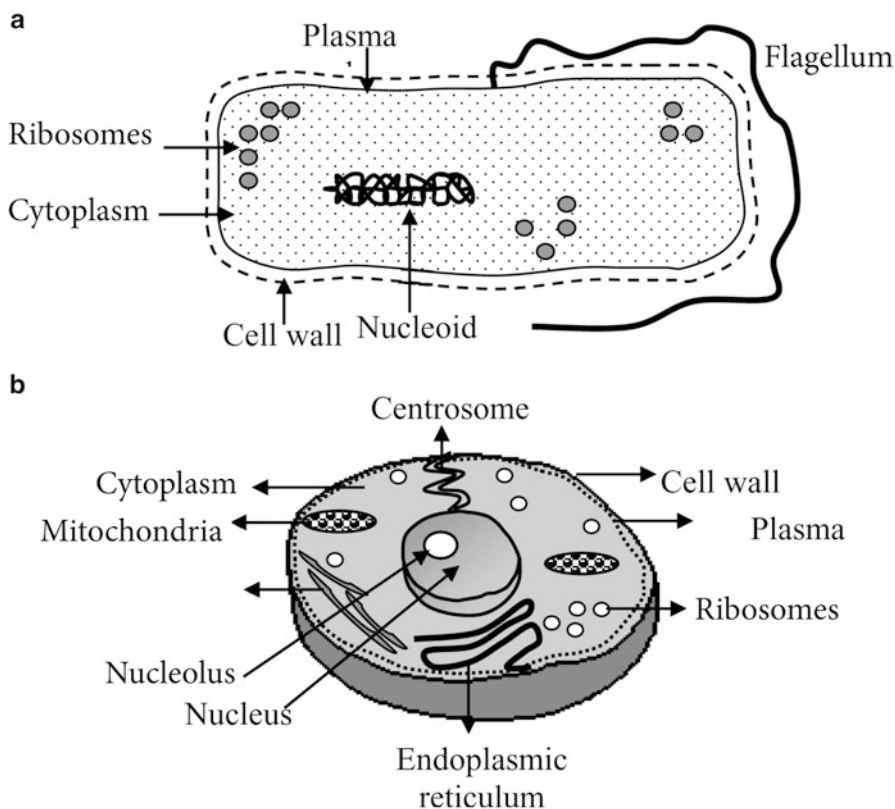


Fig. 5.3 (a) Prokaryotic and (b) Eukaryotic cells

Some microorganisms produce hydrogen sulphide (H_2S). It can oxidize organic matter forming sulphate, which in turn acts like an electron acceptor for metabolism. This H_2S can, in presence of metal salt, convert metal ions into metal sulphide, which deposits extracellularly.

In some cases, metal ions from a metal salt enter the cell body. The metal ions are then converted into a non-toxic or less toxic form and covered with certain proteins in order to protect the remainder of the cell from toxic environment.

Certain microorganisms are capable of secreting some polymeric materials like polysaccharides. They have some phosphate, hydroxyl and carboxyl anionic groups which complex with metal ions and bind extracellularly.

Cells are also capable of reacting with metals or ions by processes like oxidation, reduction, methylation, demethylation etc.

Different processes of metal-microorganism interactions are schematically shown in Fig. 5.4.

We shall have now some examples of metal-microorganism interactions.

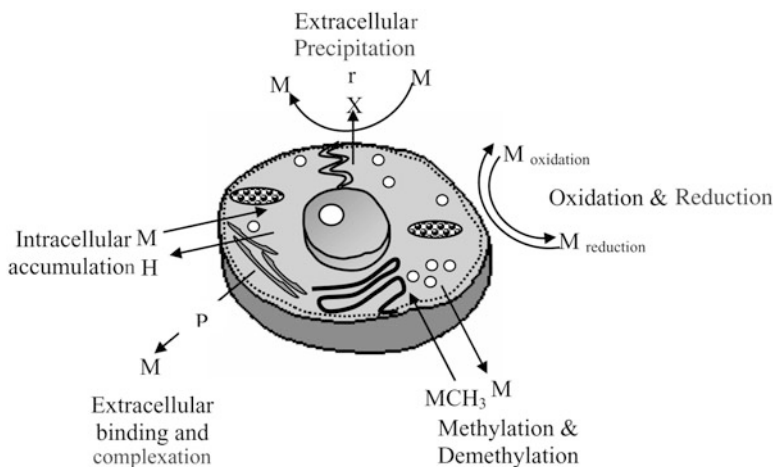


Fig. 5.4 Metal microorganisms interactions

Pseudomonas stutzeri Ag259 bacteria are usually found in silver mines and are capable of accumulating silver inside or outside of their cell walls. Using this fact, these bacterial strains can be challenged with high concentration of silver salt like AgNO_3 . Numerous silver nanoparticles of different shapes can be intracellularly produced having size <200 nm. Although detailed analysis shows that silver nanoparticles are quite abundant per cell, some silver sulphide particles also are found to be present.

Low concentrations of metal ions (Au^+ , Ag^+ etc.) can be converted to metal nanoparticles by *Lactobacillus* strain present in butter milk. By exposing the mixture of two different metal salts to bacteria, it is indeed possible to obtain alloys under certain conditions.

Use of various fungi also has been made to obtain large quantities of metal nanoparticles. For example *Fusarium Oxysporum* challenged with gold or silver salt for approximately three days produces gold or silver particles extracellularly. Extremophilic actinomycete *Thermomonospora* sp. produces gold nanoparticles extracellularly.

When silver metal salt is treated with another fungus *Verticillium* sp., the nanoparticles can be produced intracellularly. The procedure for the synthesis can be briefly described as follows. *Verticillium* sp. extracted from *Taxus* plant should be placed at 25°C in potato–dextrose agar slant. Fungus should be grown in Erlenmeyer flask containing MGY medium composed of 3 % malt extract, 1 % glucose, 0.3 % yeast extract and 0.5 % peptone maintained at $\sim 25^\circ\text{C}$. The flask needs to be shaken for 4 days. After fermentation mycelia separated from culture broth by centrifugation at 10°C , for ~ 20 min need to be washed in water. Harvested mycelial mass should be exposed to AgNO_3 solution of 5.5–6.0 pH in Erlenmeyer flask for 3 days while shaking and maintaining the 28°C temperature. *Verticillium* biomass can be fixed in glutaraldehyde in distilled water by maintaining it at room temperature for 2 h.

After fixation, the centrifugation leads to cell sedimentation rich in silver nanoparticles. Changes in biomass colour from initial yellow to final brown, after exposure to silver salt, is a visual indication of silver nanoparticles formation. Particles can be recovered by washing with some suitable detergent or ultrasonication. Interestingly, *Verticillium* does not die if exposed to AgNO_3 solution. Small specks of Ag nanoparticles containing biomass, placed in an agar plate, clearly indicate the growth of *Verticillium* in about eight days. The possible mechanism of nanoparticles formation can be considered as follows. Silver ions from a silver salt possibly get trapped on the surface of the fungal cells perhaps due to electrostatic interaction between positively charged silver ions and negatively charged carboxylic groups in the enzymes present in cell walls of mycelia. The ions after nucleation can grow by further accumulation of silver ions to form nanoparticles.

In a similar way gold nanoparticles can be produced using *Verticillium* sp. However the colour of biomass is from pink to blue depending upon the particle size.

Semiconductor nanoparticles like CdS, ZnS, PbS and many others can be produced using different microbial routes. There are some sulphate reducing bacteria of the family *Desulfobacteriaceae* which can form 2–5 nm ZnS nanoparticles. There are some reports in the literature which discuss such ZnS nanoparticle formation observed in nature.

Bacteria *Klebsilla pneumoniae*, which is a pathogen, can be used to synthesize CdS nanoparticles. When $[\text{Cd}(\text{NO}_3)_2]$ salt is mixed in a solution containing bacteria and solution is shaken for about one day at $\sim 38^\circ\text{C}$, the CdS nanoparticles in the size range ~ 5 –200 nm can be formed.

CdS nanoparticles with narrow size distribution can be synthesized using the yeasts like *Candida glabrata* and *Schizosaccharomyces pombe*. A nitrogen-rich medium containing 2 % tryptone, 1 % yeast extract and 2 % glucose (pH 5.6) is used to grow the *Schizosaccharomyces pombe*. When challenged with cadmium sulphate solution after about 12–13 h of growth, after 36 h cells rich with CdS get formed. Solution can be centrifuged and frozen at -20°C . The frozen cells can be thawed at 4°C for 2–4 h. The thawed solution can be centrifuged again so that the cell debris precipitates and supernatant contains CdS nanoparticles. The particles are covered by proteins which can be removed by heating at 80°C for few minutes. Process of CdS nanoparticles formation can be understood as follows. When the cells are exposed to cadmium salt, a series of reactions takes place. Cadmium ions are toxic to the cells and in order to reduce this effect they need to coat them with some protective coating of proteins available to them. To begin, an enzyme phytochelatin synthase gets released to synthesize phytochelatin with structure $(\text{Glu-Cys})_n\text{-Gly}$ with $n \sim 2$ –6. A low molecular weight phytochelatin–Cd complex is formed. An ATP binding cassette (ABC) type vacuolar membrane protein HMT1 transports Cd complex across the membrane. When inside the vacuole, Cd ions receive sulphur ions forming large molecular weight phytochelatin–CdS complex. These are nothing but the CdS nanoparticles.

Similarly it is possible to synthesize PbS by challenging *Torulopsis species* with lead salt like PbNO_3 .

Silver nanoparticles can be synthesized using MKY3 yeast strain which is isolated from garden soil. This is capable of tolerating ~ 0.8 mM AgNO_3 . Thus it produces silver nanoparticles extracellularly with negligible amount of intracellular particles. MKY3 can be inoculated in Erlenmeyer flask in growth medium having yeast extract (1 %), tryptone (2 %) and glucose (2 %). Solution pH should be ~ 5.6 . Flask should be shaken at 30°C for eight hours and then silver nitrate solution can be poured into it. Flask shaking should be continued for one day in dark. Cells can be then recovered by centrifugation. The cell-free medium contains silver nanoparticles which can be recovered by freeze-thaw technique in which advantage is taken of different thawing temperatures of different ingredients. Thus the solution separated from cells is poured upto the brim in a polycarbonate bottle and frozen to -20°C . It is then thawed to 0°C . At -8°C , there is a swelling of the medium. At this temperature silver particles get pushed upwards and can be collected in a sample cup. The procedure results into silver nanoparticles formation with size less than ~ 20 nm. A possible mechanism for silver nanoparticles formation may be as follows. When highly reactive silver ions from silver nitrate solution come in contact with yeast cells, the cells secrete some chemicals donating electron to metal ion. Electron donor group along with biomolecules control the accumulation of silver leading to nanoparticles formation.

5.3 Synthesis Using Plant Extracts

Use of plants in synthesis of nanoparticles is quite novel leading to truly *green* chemistry that technologists are looking for. However, compared to the use of microorganisms to produce nanoparticles, use of plant extracts is relatively less investigated. There are few examples which suggest that plant extracts can be used in nanoparticles synthesis. For example it has been reported that live alfalfa plants are found to produce gold nanoparticles from solids.

Leaves of geranium plant (*pelargonium graveolens*) have also been used to synthesize nanoparticles of gold. It should be mentioned that there is also a plant associated fungus which can produce compounds such as taxol and gibberellins. There is an exchange of intergenic genetics between fungus and plant. However, the nanoparticles produced by fungus and leaves have quite different shapes and size distributions. Nanoparticles obtained using *Colletotrichum* sp. fungus related to geranium plant has a wide distribution of sizes and particles are mostly spherical. On the other hand geranium leaves produce rod and disk shaped nanoparticles.

Synthesis procedure to obtain gold nanoparticles from geranium plant extract is as follows. Finely crushed leaves are put in Erlenmeyer flask and boiled in water just for a minute. Leaves get ruptured and cells release intracellular material. Solution is cooled and decanted. This solution is added to HAuCl_4 aqueous solution, when nanoparticles of gold start forming within a minute.

5.4 Use of Proteins, Templates Like DNA, S-Layers etc.

As discussed earlier, various inorganic materials like carbonates ($-\text{CO}_3$), phosphates ($-\text{PO}_4$) and silicates ($-\text{SiO}_4$) are found to be parts of bones, teeth, shells etc. Thus biological systems are capable of integrating with inorganic materials. This has been widely used not only to synthesize nanoparticles but obtain organized arrays, superlattices or hierarchical structures of inorganic materials using biological systems. DNA, S-layers (i.e. surface layers of cell walls of some bacteria) or some membranes have long range periodic order in terms of some molecular groups of their constituents. Therefore on some periodic active sites preformed nanoparticles can be anchored. Alternatively, using certain protocols nanoparticles can be synthesized using DNA and membranes as templates. Such ordered arrays are formed as a result of various interactions that take place between the templates and the particles. Such arrays are known as self assembly and will be discussed in Chap. 6. Here we shall discuss only the nanoparticle synthesis using ferritin (a protein) and DNA to produce nanoparticles which are not self assembled.

Ferritin is a colloidal protein of nanosize. It stores iron in metabolic process and is abundant in animals. It is also capable of forming uniform three dimensional hierarchical architecture. Ferritin proteins in different animals have only upto 14 % difference in the amino acid sequences. There are 24 protein (peptides) subunits in a ferritin, which are arranged in such a way (see Fig. 5.5a) that they create a central cavity (see Fig. 5.5b) of ~ 6 nm. Diameter of polypeptide shell is 12 nm. Ferritin can accommodate 4,500 Fe iron atoms. They are in Fe^{3+} state as hydrated iron oxide mineral, ferrihydrite. The protein subunits are composed of light as well as heavy chains having dinuclear ferroxide centres. These centres are catalysts for in vitro oxidation of Fe^{2+} ions.

Methods to extract ferritins are quite standard in biology. The ferritin without inorganic matter in its cavity is known as apoferritin and can be used to entrap desired nanomaterial inside the protein cage. Therefore first step is to remove

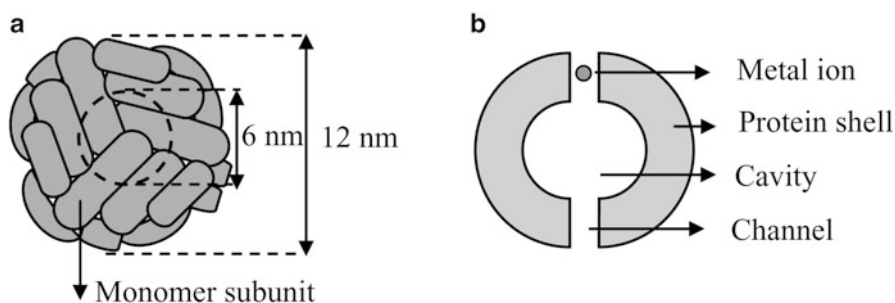


Fig. 5.5 Schematic representation of (a) ferritin molecule and (b) cavity formed by polypeptide units

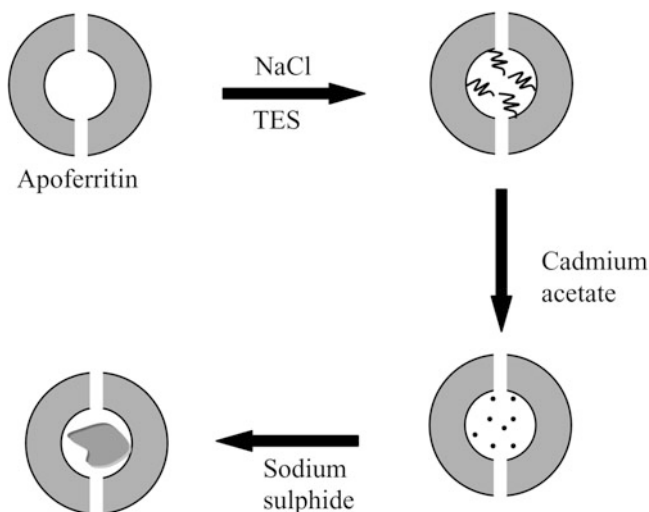


Fig. 5.6 Synthesis of CdS nanoparticles using ferritin

iron from ferritin to form apoferritin and then introduce metal ions to form metal nanoparticles inside the cavity or carry out some controlled reaction with metal ions to make a compound inside the cavity. In any case, ions can be removed or introduced inside the ferritin, through some available channels (see Fig. 5.5).

We shall discuss now the procedure to convert ferritin into apoferritin and how to use it in the synthesis of CdS nanoparticles. Spleen ferritin diluted with sodium acetate buffer should be placed in dialysis bag. A solution of sodium acetate and thioglycolic acid is made in which dialysis bag is kept under nitrogen gas flow for 2–3 h. Solution needs to be replaced from time to time for total 4–5 h. Further dialysis of apoferritin solution should be done against saline for one hour and in refreshed saline for ~15–20 h. Apoferritin should then be mixed with solution having sodium chloride (NaCl) and N-tris (hydroxymethyl) methyl-2-aminoethanosulphonic acid (TES). Aqueous cadmium acetate is added to this solution and stirred continuously with constant N_2 gas purging. Aqueous solution of sodium sulphide (Na_2S) is added twice with one hour interval.

Process of CdS formation is stepwise (see Fig. 5.6) with Cd loading of 55 atoms per apoferritin colloid taking place in each step. Higher loadings like 110, 165, 220 ... are possible. Due to remarkably constant size of ferritin colloids and apoferritin derived from them, it is possible to obtain nanoparticles of very uniform size.

Besides CdS there are several other examples like controlled iron oxide, manganese, uranyl oxide, cobalt, cobalt-platinum alloy being synthesized inside ferritins.

5.5 Synthesis of Nanoparticles Using DNA

CdS (or other sulphide nanoparticles) nanoparticles can be synthesized using DNA. We have seen in Chap. 4 that organic molecules can cap the surfaces of nanoparticles growing in solutions. Similarly one can use DNA to bind with surface of growing nanoparticles. For example double stranded Salmon Sperm DNA can be sheared to an average size of 500 bp. Cadmium acetate can be added to desired medium like water, dimethylformamide, ethanol, propanol etc. and reaction is carried out in a glass flask with facility to purge the solution and flow with an inert gas like nitrogen. Addition of DNA should be made and then Na_2S can be added dropwise. Depending upon the concentrations of cadmium acetate, sodium chloride and DNA nanoparticles of CdS with sizes less than ~ 10 nm can be obtained. It is possible to prove the presence of DNA on CdS nanoparticles. It is found that CdS nanoparticles synthesized by this route have cadmium-rich surface. DNA probably bonds through its negatively charged phosphate group to positively charged (Cd^+) nanoparticle surface. The other end of DNA is in fact free to interact with suitable proteins; such particles can be used as sensor of proteins.

We shall not go into more discussion of use of DNA here, but preformed charged nanoparticles can get bonded with phosphate group of DNA and even form organized arrays of nanoparticles.

Further Reading

- G.L. Hornyak, H.F. Tibbals, J. Dutta, J.J. Moore (eds.), *Introduction to Nanoscience & Nanotechnology* (CRS Press, Boca Raton, 2009)
- S.K. Kulkarni (ed.), Special issue on nanomaterials. *Phys. Educ.* **19**, (2002)
- C.M. Lukehart, R.A. Scott (eds.), *Nanomaterials: Inorganic and Bioinorganic Perspectives* (Wiley, New York, 2008)
- C.M. Niemeyer, C.A. Mikić (eds.), *Nanobiotechnology: Concepts, Applications and Perspectives* (Wiley, Weinheim, 2004)

Chapter 6

Self Assembly

6.1 Introduction

The term ‘Self Assembly’ itself indicates its meaning. It is a gathering or collection of certain entities without any external influence. Although the term ‘self assembly’ has received a scientific acceptance over the last 3–4 decades and such assemblies are recognized in biology, chemistry or physics, some philosophers like Kanad, Democritus and Descartes several centuries back had imagined that everything in the world, from small objects around us to solar system, galaxies and universe is a result of tiny, non-divisible units (or atoms). The word *átomos* (Greek: ἄτομος) was first used by Democritus to mean uncuttable. Out of a chaotic situation, the various forms of matter get ordered following some laws of nature. Today we know that atoms can be further smashed with very high energy into still tinier particles viz. electrons, protons and neutrons. Protons and neutrons also are composed of fundamental particles like quarks. Yet, the concept that matter is composed of tiny units like atoms or molecules is still very useful for understanding most of the phenomena around us.

One may tend to think that the term ‘self assembly’ can be used for any matter in which atoms or molecules stick together, like in a solid. However this is not the case. As will be discussed in this chapter the term self assembly applies for spontaneously formed, reversible, locally ordered, thermodynamically stable assemblies. Self assemblies are very sensitive and can transform back to the state of disorder. In disorder state there are some building blocks or motifs which are uniform in the shape and size which assemble to form ordered structure by some weak interactions spontaneously. The building blocks themselves is an assemblage of strong interacting particles (atoms or molecules).

Self assembly initially, during the twentieth century, was considered to be limited to biological world where origin of beautiful colours of peacock feathers,

butterfly wings and many birds or insects were understood as a result of ordered structures. Microscopy analysis reveals that often there are micro plus nano or just nanostructures involved. Phenomenon of self assembly occurs at many length scales and is not only limited to tiny objects but even gigantic astronomical objects like stars, planets, galaxies and entire universe which have been considered as a kind of self assembly.

Observations of self assembly in naturally occurring living and non-living world have led scientists to understand what leads to the ordered or random assemblies of some smaller entities, blocks, motifs or units. We know that the major bonds prevailing in the inorganic solids are ionic, covalent or metallic. They have rather large energies of formation (or dissociation) typically from ~ 0.5 to a few electron volts. Self assemblies are formed spontaneously by weak interactions like π - π , van der Waals, colloidal, electrical, optical, shear or capillary forces.

In nanotechnology, self assembly plays an important role. Close packed arrangement of organic molecules and nanoparticles is useful in novel devices. The technique of organic thin films deposition by Langmuir-Blodgett (L-B) technique discussed in Chap. 4 (Sect. 4.6) is a very nice example of self assembly. However, there was no specific recognition of the term 'self assembly' associated with L-B technique. The term 'self assembly' became popular when Nuzzo and Allara in 1983 used this terminology when some disulphide alkyl molecules were deposited on gold surfaces.

Recognition of importance of self assembly and its origin has led scientists to make deliberate attempts to fabricate self assembled organic, inorganic or materials useful to obtain novel, electrical, mechanical, magnetic or optical materials with unprecedented properties. Langmuir-Blodgett films, micelles, liquid crystals, layers obtained by dip-pen lithography, deposition of materials in the voids of the self assembled spheres, and anodization of alumina templates with ordered pores are some of the methods to realize self assembled structures with desired materials. Such structures are useful to create photonic band gap materials, novel sensors, lasers, Bragg mirrors, electroluminescent devices, photovoltaic solar cells to name a few. Nanofabrications using DNA have potential applications in nanoelectronics, nanomechanical devices as well as computers.

A very important branch developed in chemistry, known as "Supramolecular chemistry" is solely the manifestation of 'self assembly'. The term "Supramolecular chemistry" was coined by Nobel laureate Jean-Marie Lehn to mean the *chemistry beyond molecules*. It is essentially an assemblage of molecules of one or few more types to make aggregates or larger crystals through non-covalent interactions. The 'molecular recognition' (like lock and key) helps build larger assemblies as in two strands of DNA winding around each other. Three dimensional ordered arrangement of such molecular assemblies can lead to build up 'superlattices' or large single crystals of self assembled molecules.

6.2 Mechanism of Self Assembly

It is important to realize that the self assembly involves weak to strong forces and nanoscopic to gigantic structures in one, two or three dimensions. As mentioned in the introduction, the self assembly can be a result of very weak forces like van der Waals interaction, hydrogen bonds, static charges, magnetic interactions and so on. The driving force behind the self assembly has been recognized as due to attempt of any system to go to the lowest energy state. Ability of a system to go to a well ordered low energy state depends upon the availability of units of same size and shape. Molecules with definite shape, number of atoms and therefore size, already in the state of low energy are good candidates for the self assembly. It should be remembered that atoms in such molecules or building blocks of self assembly themselves are held together by stronger forces.

When the building blocks (one, two or more types of building blocks also can form self assemblies) are available for self assembly, state of minimum energy may be attained spontaneously in the absence of any external force. However self assembly may occur in the presence of an external driving force like temperature, pressure, magnetic field and so on. These two types viz. assembly in the absence and presence of an external driving force are known as *static* and *dynamic* assemblies respectively (see Fig. 6.1). Static assembly is realized when system achieves minimum energy state and can stay there unless it is subjected to strong external forces. On the other hand, dynamic self assembly involves constant influence of external force from the ambient. If the energy intake from the ambient stops, the self assembly can leave the state of organized structure and de-assemble. Formation of ordered crystalline structure from a melt can be considered as an example of static self assembly. Thin films formed by L-B technique (Chap. 4) or dip pen lithography

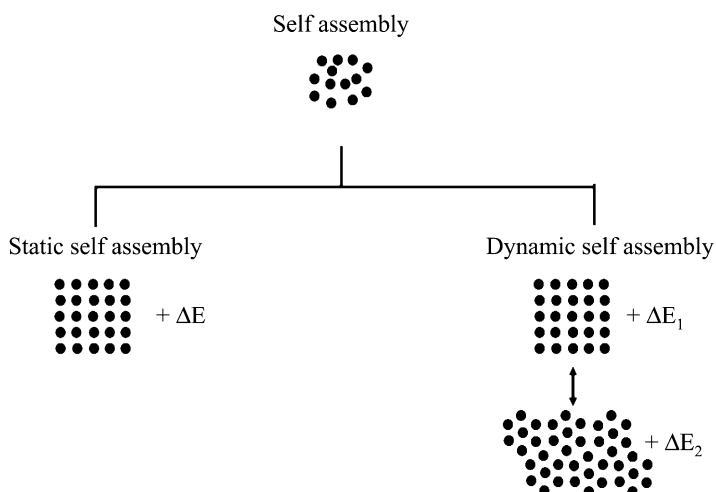


Fig. 6.1 Two types of self assemblies

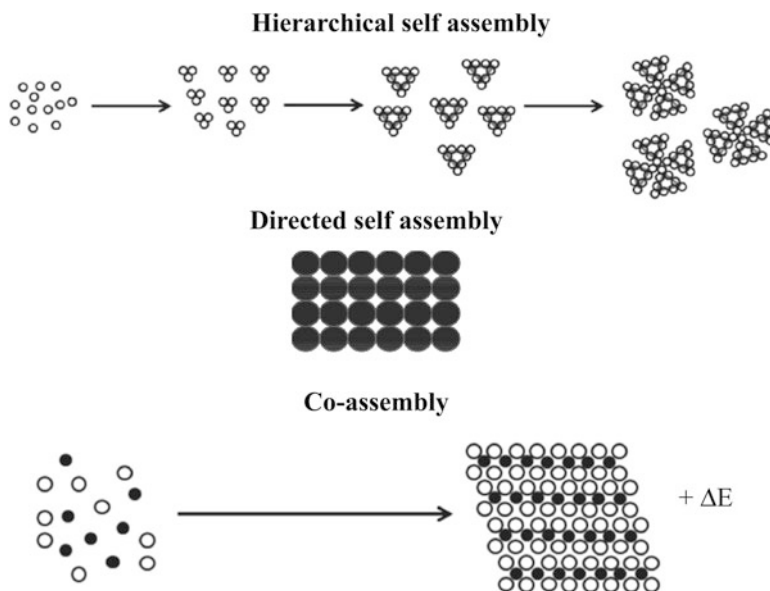


Fig. 6.2 Hierarchical and co-assembly

technique (to be discussed in Chap. 9) also are examples of static self assembly. The living animals or plants form good examples of dynamic self assembly. As soon as the supply of food, proper temperature and air pressure discontinue, the animals and plants disintegrate.

The static and dynamic assemblies can be further divided into ‘*hierarchical self assembly*’, ‘*directed self assembly*’ and ‘*co-assembly*’ as illustrated schematically in Fig. 6.2.

Hierarchical self assembly is characterized by small range, medium range and long range interactions of one type of a building block.

Directed self assembly occurs when the building blocks occupy the pre-designed places like some portions of a lithographically patterned substrate, pores in membranes or spaces between ordered particles.

Co-assembly, as the name suggests, can be formed with two or more types of blocks which can fit into each other.

We shall discuss in the following section some examples of self assembly. Evaporation or biological templates like DNA or S-layers are useful for self assembly. In most of these cases, it is possible to place the nanoparticles at some well determined location of polymer chain, organic molecule or template through specific bonding. When well structured templates like S-layers or DNA are used, direct patterning of nanoparticles becomes feasible. These are the examples of inorganic (nanoparticles)-organic molecules or inorganic-biomolecules assemblies. However it is also possible to have assemblies of purely inorganic particles. The driving force for such assemblies is often quite different. We shall discuss now little more details of various assemblies.

6.3 Some Examples of Self Assembly

6.3.1 Self Assembly of Nanoparticles Using Organic Molecules

Preformed inorganic nanoparticles can be assembled on solid substrates through some organic molecules adsorbed on their surfaces. For example CdS nanoparticles functionalized with carboxylic ($-\text{COO}^-$) group can be transferred (Fig. 6.3a) to aluminium thin films. Dithiols adsorbed on metal surfaces also could adsorb CdS nanoparticles (Fig. 6.3b) to form layers of them. Silver particles (Fig. 6.3c) have been adsorbed on oxidized aluminium layers using bi-functional molecule such as 4-aluminium layers carboxylthiophenol. Molecules bind to aluminium oxide layer through carboxylic group and thiol attaches to silver particles.

Using a two-phase reaction alkanethiol or alkylamine capped gold, silver, palladium etc. nanoparticles have been self assembled. Here chemical reaction takes place in an aqueous medium. The particles are then transferred into an organic solvent and drop casted on an appropriate solid substrate. Solvent is allowed to evaporate which leaves self assembled layer.

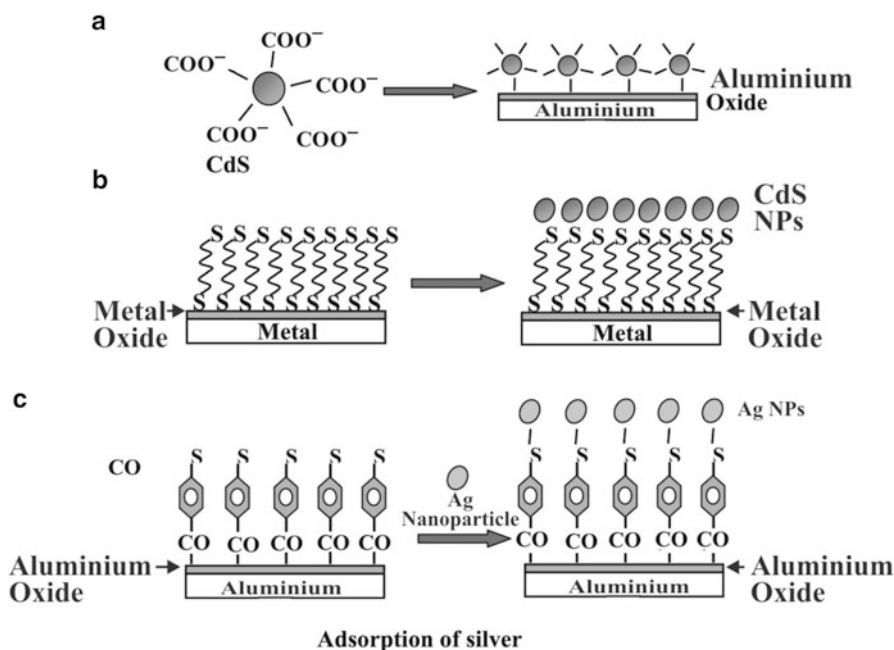


Fig. 6.3 Self assembly of nanoparticles

6.3.2 Self Assembly in Biological Systems

There are many examples of self assembly in the biological systems like S-layers, proteins or DNA. When organized arrays of inorganic crystals are embedded in biological systems they are often referred to as biomineralized systems. Here we shall discuss only the ordered self assembled layers as examples of biomineralization.

Magnetotactic bacteria are small bacteria, $\sim 35\text{--}120$ nm sized with permanent magnets inside them. The magnets are of either iron sulphide (Fe_3S_4 -greigite) or iron oxide (Fe_3O_4 -magnetite). Such magnets make a chain of nanomagnets. The size $35\text{--}120$ nm is quite critical and smaller or bigger magnets would not have served the purpose for which these magnets are used. The magnets smaller than 35 nm cannot have permanent magnetism at ambient temperature and those larger than 120 nm would have reduced magnetism due to multidomain formation above this size. Only those between 35 and 120 nm size are single domain permanent magnets, the chain of which is useful for navigation of bacteria. These bacteria live in mud, marshy areas, ponds or sea. and prefer anaerobic condition. If by any chance they come to the surface of water or soil, they navigate downwards making use of their magnets aligning in the earth's magnetic field direction. Earth's magnetic field has a dip in the north and south hemispheres which helps bacteria to seek downwards direction. Even some birds have ordered nanomagnets which they sometimes use for navigation purpose.

A very common self assembly in biological systems is S-layers. They are one of the simplest kind of biomembrane evolved during billions of years and is invariably a part of cell envelope of prokaryotic organisms, with just a few exceptions. They are two dimensional, crystalline single proteins or glycoprotein monomers organized in hexagonal, oblique or square lattices, as illustrated schematically in Fig. 6.4.

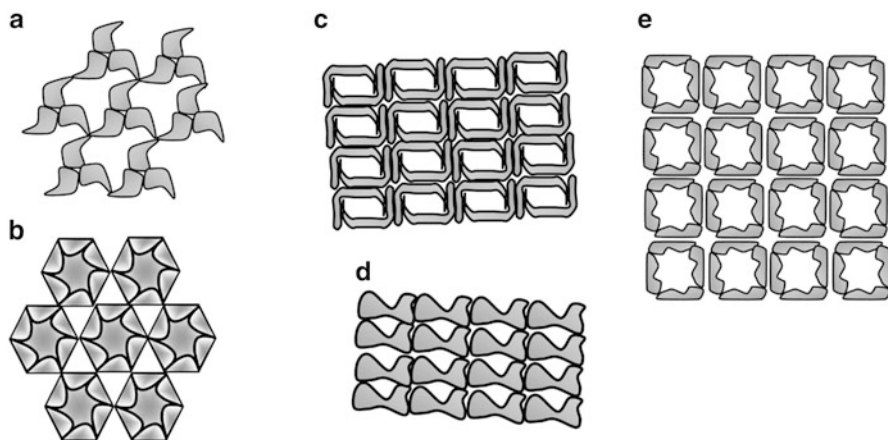


Fig. 6.4 Different types of S-layer lattice

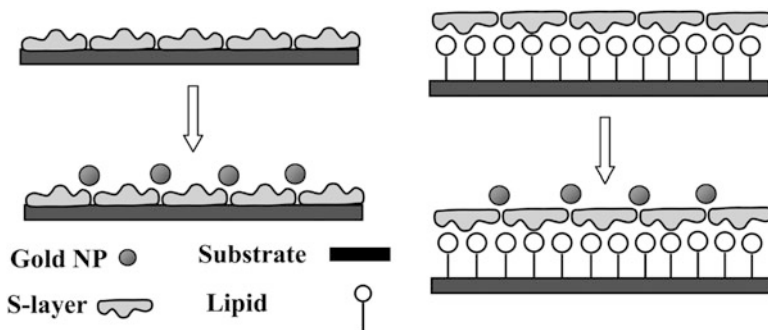


Fig. 6.5 Assembly of CdS nanoparticles using S-layer

These lattices have ordered pores. The periodicity of pores can vary, depending upon the protein, from 3 to 35 nm. In some cases pores $\sim 2\text{--}8$ nm have also been observed.

Such S-layers after extraction from bacterial cells have been transferred on some metallic substrates (or grids). When treated with cadmium salt and subsequently with Na_2S , ordered arrays of CdS nanoparticles could be formed. This is confirmed using electron microscopy. S-layers are reported to have been used to assemble Au, Pt, Fe and Ni metal nanoparticles. In general, S-layers extracted from the biological cells can be directly used to deposit nanoparticles from liquid or vapour phase as shown in Fig. 6.5.

It was discussed in Chap. 5 that ferritins are protein colloids of 12 nm size found in all animals. Ferritins have cavities $\sim 6\text{--}8$ nm in size filled with iron oxide. It is possible to remove iron oxide and replace it with metal or other nanoparticles. Further it is possible to make a two dimensional array of ferritins in solution. For example ferritin solution in NaCl and phosphate at ~ 5.8 pH can be filled in a trough. Chloroform containing dichloroacetic acid can be used to dissolve poly-1-benzal-L-histidine (PBLH) and spread over ferritin solution in the trough. After about two hours the solution can be heated at 38°C for one hour and cooled back to room temperature. This produces ordered layer of ferritin at liquid-air interface. The layer can be transferred on silicon substrate by dipping it in the solution. An ordered layer of ferritins is transferred on the surface. By heating the silicon substrate for one hour in nitrogen ambient at 500°C , protein from ferritin is completely removed, leaving an ordered layer of iron oxide or whatever nanoparticles were filled in ferritin. There are several reports of two dimensional ordered layers being transferred by this route of Co, CoPt, FePt. The major application of this type of layers is in high density magnetic data storage systems. It is speculated that CoPt particles in this manner will be able to have as large as $1,550\text{ Gbit/cm}^2$ per particle. This is quite large compared to current $\sim 11\text{ Gbit/cm}^2$.

DNA (deoxyribonucleic acid) is a long helical molecule and was discussed in Chap. 5. It has a large aspect ratio (i.e. ratio of length to diameter is large) and acts like a long one dimensional template in its simplest form. Its four nucleotide bases

viz. quinine (Q), cytosine (C), adenine (A) and thymine (T) can form a rich variety of sequences and structures. Thus circular, square, branched and many more long (few micrometre) or short (few nanometre) DNA templates are possible. Besides planar geometry, they can adopt even three dimensional structures. As DNA has alternate sugar and phosphate groups on its strands, it is possible to anchor metal, semiconductor or oxide particles by different bonding on DNA to have assembly of particles.

6.3.3 Self Assembly in Inorganic Materials

It is possible to spontaneously create the quantum dots for example of germanium (Ge) on silicon (Si) or indium arsenide (InAs) on gallium arsenide (GaAs). The origin of self assembly is strain induced. Germanium and silicon have only 4 % lattice mismatch. Therefore Ge can be deposited epitaxially on Si single crystal upto 3–4 monolayers. Although grown (hetero)epitaxially, the layers of deposited Ge are highly strained (coherently i.e. without any defects or dislocations). When further deposition takes place, the lattice strain caused by depositing Ge on Si with different lattice constants cannot be accommodated.

This results in spontaneous formation of nanosized islands or quantum dots. However the temperature of the substrate has to be $>350^\circ\text{C}$ during deposition or post-deposition annealing is required. Figure 6.6 schematically illustrates the growth mechanism as well as an electron microscopy image of germanium islands on Si (111) surface. The size of the islands depends upon the growth temperature as well as the substrate plane on which it grows.

Preformed inorganic particles of materials like silica (SiO_2), titania (TiO_2), polymer beads or latexes are able to organize themselves just by sedimentation also. But they need to have very uniform size. As illustrated in Fig. 6.7, organization of SiO_2 particles is quite clear. The silica particles can be formed by a sol-gel route. The particles synthesized in aqueous medium are simply allowed to evaporate from

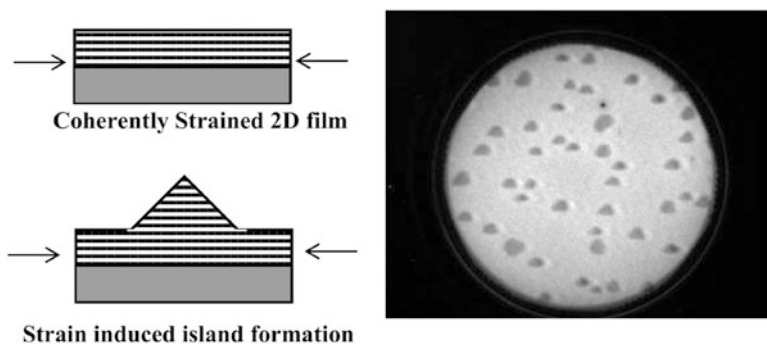
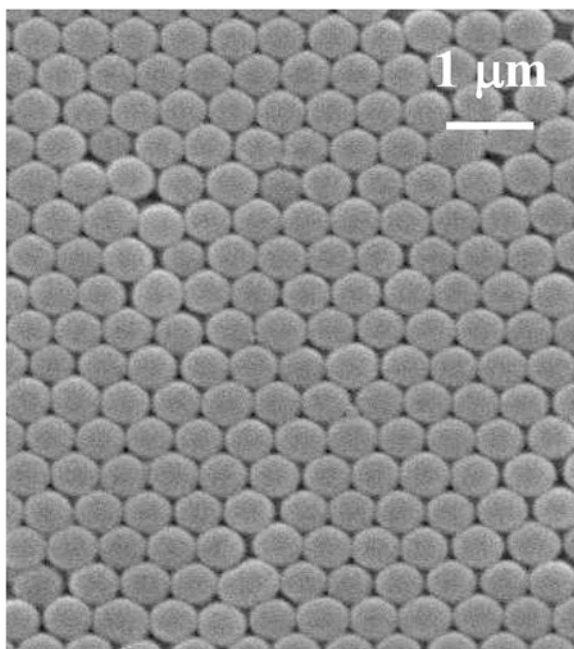


Fig. 6.6 Growth mechanism of Ge on Si and photograph showing island formation (Field of view $10\ \mu\text{m}$)

Fig. 6.7 SEM image of self assembled silica (SiO_2 particles) on glass slide



the solution drop placed on a glass substrate. After some time the particles self assemble due to weak van der Waals interaction amongst the particles. The driving force is the capillary force. Minimization of surface energy takes place by forming a hexagonal network. The uniform size of the particles is helpful to make an ordered two dimensional network of particles.

Self assembly thus has a very rich variety of practical examples. Many other self assemblies can also be designed and fabricated. Such assemblies have a great potential in molecular information technology as it is possible to store information and process information. It also is useful for nanofabrication. If complex arrangements of self assembly can be made like in biological processes, compact devices can also be realized using self assembly. One such example given by Lehn is the brain which is self assembled and wired by self organization to different organs. It is able to store information, process and also has control over different body parts, all through self assembly or organization.

Further Reading

H. Dodziuk, *Introduction to Supramolecular Chemistry* (Springer, Dordrecht, 2002)

J.-M. Lehn, *Supramolecular Chemistry: Concepts and Perspectives* (VCH, Weinheim, 1995)

Materials Today, special issue on Nanofabrication by Self Assembly, **12**, May 2009

J.W. Steed, D.R. Turner, K.J. Wallace, *Core Concepts in Supramolecular Chemistry and Nanochemistry* (Wiley, New York, 2007)

Chapter 7

Analysis Techniques

7.1 Introduction

Nanomaterials, dispersed in the form of colloids in solutions, particles (dry powders) or thin films, are characterized by various techniques. Although the techniques to be used would depend upon the type of material and information one needs to know, usually one is interested in first knowing the size, crystalline type, composition, thermal, chemical state, and properties like optical or magnetic properties. A list of various commonly used techniques and their utility can be found in Box 7.1.

In this chapter, we shall briefly outline some of the techniques from those mentioned in Box 7.1.

7.2 Microscopes

Low dimensional materials such as quantum dots, quantum wires, quantum wells, self-assembled materials, interactions of small molecules with surfaces, multilayers etc. need special microscopes like Transmission Electron Microscope (TEM), Scanning Tunnelling Microscope (STM), Atomic Force Microscope (AFM) and Scanning Near-Field Optical Microscope. However, in many instances, one inspects the sample, specially thin films, with an optical microscope, to check the quality of samples like presence of cracks, agglomeration at large scale, etc. Therefore we shall begin with a simple optical microscope, define some common terminologies used in microscopy analysis work and then discuss other microscopes.

7.2.1 *Optical Microscopes*

Human eye perceives an object when visible light reflected from an object enters the eye. Size of an object observed by an eye depends upon the arc subtended by

the object at the lens and image on the retina of the eye. As illustrated in Fig. 7.1a, smaller the distance from the eye, bigger is the image of the object in the eye. There are, however, two limitations. An object kept at a distance smaller than ~ 25 cm (this distance is known as the distance of distinct vision) from the eye cannot produce a sharp image of the object and other is that a human eye cannot detect an object smaller than $\sim 100\text{ }\mu\text{m}$ as a distinct object if placed close to another object. However, as shown in Fig. 7.1b, by placing a convex lens close to an eye, a magnified virtual image can be formed at a larger distance so as to form an image with larger angle θ' . Such a magnifying lens forms the simplest kind of microscope.

Box 7.1: Commonly Used Techniques in Materials Analysis

Microscopes

Optical microscope, Confocal microscope, Scanning Electron

Microscope (SEM), Transmission Electron Microscope (TEM), Scanning Tunnelling Microscope (STM), Atomic Force Microscope (AFM), Scanning Near-Field Optical Microscope (SNOM).

Microscopes are useful to investigate morphology, size, structure and even composition of solids depending upon the type of microscope. Some of the powerful microscopes are able to resolve structures up to atomic resolution.

Diffraction Techniques

X-ray Diffraction (XRD), Electron Diffraction, Neutron Diffraction, Small Angle X-ray Scattering (SAXS), Small Angle Neutron Scattering (SANS) and Dynamic Light Scattering (DLS).

Scattering or diffraction techniques are often used in particle shape and average particle size analysis as well as structural determination.

Spectroscopies

UV-Vis-IR absorption (transmission and reflection modes), Fourier Transform Infra Red (FTIR), Atomic Absorption Spectroscopy (AAS), Electron Spin (or Paramagnetic) Resonance (ESR or EPR), Nuclear Magnetic Resonance (NMR), Raman Spectroscopy, various luminescence spectroscopies, Electron Spectroscopy for Chemical Analysis (ESCA) or X-ray Photoelectron Spectroscopy (XPS), Auger Electron Spectroscopy (AES).

Spectroscopies are useful for chemical state analysis (bonding or charge transfer amongst the atoms), electronic structure (energy gaps, impurity levels, band formation and transition probabilities) and other properties of materials.

Electric and Magnetic Measurements

Two or four probe measurements, Magnetoresistivity, Vibrating Sample Magnetometer (VSM), Superconducting Quantum Interference Device (SQUID), Magneto-Optical Measurements (Kerr and Faraday rotations).

(continued)

Box 7.1 (continued)

Measurement of resistivity is necessary for many applications. Magnetic and magneto-optical measurements throw light on the behaviour of the materials in presence of external magnetic fields.

Mechanical Measurements

Hardness, strength (Elastic moduli), Nanoindentation.

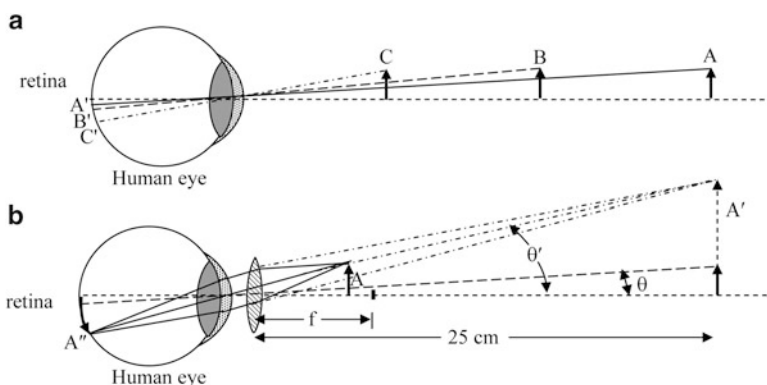


Fig. 7.1 (a) Size of the image depends on the distance from the eye. (b) By keeping a convex lens close to the eye, image of an object can be magnified

Magnification is defined as the ratio of angles subtended by the image (θ') to that (θ) by the object

$$M = \frac{\theta'}{\theta} \quad (7.1)$$

Magnification is approximately related to focal length f of the lens as

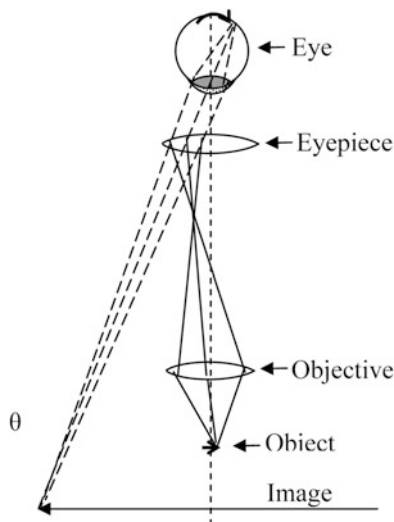
$$M \simeq \frac{25}{f} \quad (7.2)$$

where 25 in the numerator is the distance of distinct vision in cm and ' f ' is the focal length of magnifying lens in cm.

An object would appear ten times larger with lens having $f = 2.5$ cm ($M = 25/2.5$) and is written as '10 $\times$ '. When an image is 100 times bigger than the object, it is written as '100 $\times$ ' and so on.

With the magnifying ability of lenses in mind, Galileo invented in the year 1610 the simplest optical microscope. Currently used microscopes make use of at least two lenses viz. objective and eyepiece. Objective lens is the lens close to the object and eyepiece is close to the eye, as shown in Fig. 7.2.

Fig. 7.2 Ray diagram of the simplest kind of an optical microscope using just two convex lenses



The objective lens (from a distance larger than its focal distance) forms the real image of an object, which in turn gets magnified as a virtual image due to eyepiece.

$$\text{Overall magnification } M = \frac{-x}{f_1} \cdot \frac{25}{f_2} \quad (7.3)$$

where $-x/f_1$ is magnification due to objective lens and $\frac{25}{f_2}$ is the magnification due to eyepiece. Here 'x' is the distance of image from the focal point of the lens. Negative sign indicates that the image is inverted (Box 7.2).

Box 7.2: Resolution, Magnification and Depth of Focus of a Microscope

While using any microscope we should know its capabilities in terms of its resolution, magnification and depth of focus, which determine the extent to which we can get information about the samples under investigation. We shall discuss these below.

Resolution

It is the ability to produce two separate images of two closely spaced objects. The limit up to which the objects are resolved is determined by the diffraction of the waves. *Rayleigh criterion*: Maximum intensity of the peak due to one object should fall at the minimum intensity due to the other object.

(continued)

Box 7.2 (continued)**Magnification**

It is the ratio of the size of the image to size of the object. Overall magnification is given by the product of magnification produced by various components of a microscope and camera factor.

$$M = M_o \times M_e \times M_c \times C$$

where M is overall magnification of the microscope, M_o – magnification due to objective lens, M_e – magnification due to the eyepiece, M_c is magnification due to change (if used) and C is camera factor.

Depth of Focus

It decreases with increased magnification. Sharpness of an image depends upon the numerical aperture 'NA' and the refractive index n of the lens material.

It should be remembered that having large resolution is not sufficient. Magnification is also very important. A good microscope should have high resolution, sufficient magnification and adequate depth of focus.

In a commercial optical microscope more than two lenses, apertures, sample stage and light source are present in order to improve the quality of the observed image (lenses produce various defects in images like distortion, astigmatism etc. which can be partially corrected). Magnification of an optical microscope cannot be increased indefinitely as all the microscopes are limited by their ability to resolve the images to some limit. This is limited by the diffraction of the scattered light from an object. Two close-by areas on a sample can be considered as two apertures, the light passing through which can interfere to form a combined image. The closest distance between two points (or two areas), which can be seen as separate or resolved is given by

$$R = \frac{\lambda}{2n \sin \theta} \approx \frac{\lambda}{2NA}, \text{ for small values of } \theta \quad (7.4)$$

where λ is wavelength of light and θ is the semi cone angle of light entering the objective lens from the sample (or object) (see Fig. 7.2).

$NA = n \sin \theta$ is the numerical aperture of a lens with n as the refractive index of the lens. Approximating $NA = 1$, $R = \lambda/2$ which is often referred to as $\lambda/2$ limit or diffraction limited resolution. $\lambda/2$ limit is common to all microscopes based on the principle of scattering of waves, may be electromagnetic or those associated with particles.

Optical microscopes in general can resolve up to $\sim 0.2 \mu\text{m}$ as visible light ranges from 400 to 700 nm and the smallest wavelength, which can be used is 400 nm.

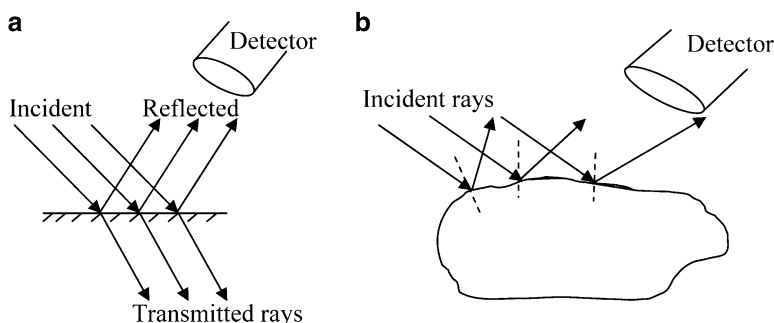


Fig. 7.3 (a) On a flat surface intense reflected beam will pass in the detector. (b) On a rough surface, due to scattering in different directions, the intensity of reflected beam passed in the detector would be less

In order to observe the images of objects, additionally, it is necessary to obtain sufficient contrast between the image of interest and its surrounding. This depends upon the method of illumination, absorption of light due to sample and some other factors like polarization of light etc.

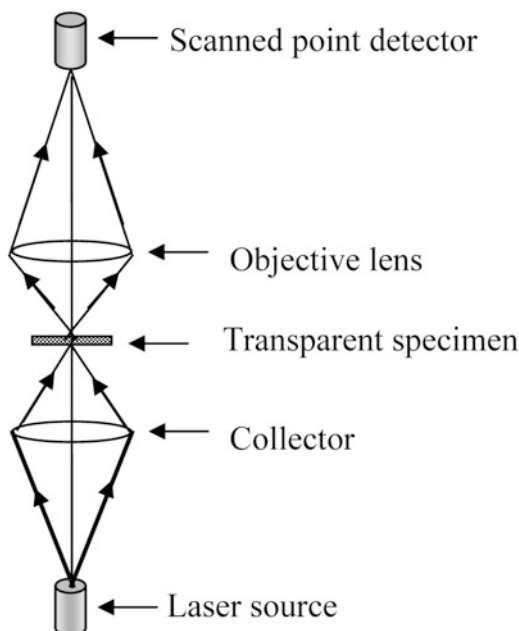
When a beam of light is incident on a perfectly flat solid surface, it is known, following Snell's law, that it can get partly reflected and partly refracted (see Fig. 7.3a). However, if we consider a beam of light falling on a rough surface, as illustrated in Fig. 7.3b, then depending upon the surface roughness or morphology the intensity of the reflected beam can vary in different directions (Snell's law is obeyed but one needs to consider local normal to the surface) or the reflected beam would diverge.

This would mean that there would be an intensity variation from the sample surface. If the sample is having some grains or different optical properties (refractive index or reflecting power) for different parts of the sample surface, then intensity variations would occur and can be detected. The optical microscopes in the current use are equipped with a set of lenses to improve the sample illumination, sample movement stage, lenses with different magnification, camera with various apertures and other facilities to obtain high quality images as well as ease of operation.

7.2.2 Confocal Microscope

Resolution of an optical microscope can be improved by limiting the field of view. This is the principle used in a confocal microscope. A confocal microscope can be of transmission or reflection type. In Fig. 7.4, a ray diagram of a confocal microscope in transmission mode is illustrated. A point source of light and a small area detector are used in this microscope, which restrict the field of view. The light from the point source is focussed on the specimen to cover only a point (or very small area) on the sample. The objective lens in turn forms a small image of this illuminated portion on the point detector. It is possible to raster (move in x-y directions or in

Fig. 7.4 Schematic diagram of a confocal microscope in transmission mode. Either specimen or beam falling on the specimen are synchronously rastered and image is reconstructed on the computer



a plane) the sample so that light falls on each part of the sample and the signal gathered by the detector is used to construct the image. Alternatively, the point source and detector are synchronously rastered to view the entire specimen and construct the corresponding image. Use of point source and detector improves the depth resolution of a confocal microscope. It is capable of optically sectioning a three dimensional thick object to a resolution determined by the detected sample volume. The detector uses an aperture which eliminates the light not coming from the focus on the sample and scans the sample in one plane point by point. Image of one plane is stored in a computer. By adjusting the focus of light in a different plane, point by point the entire plane is scanned and image of that plane is stored. Like this the sample is scanned to achieve high resolution 3-D image of the sample. Confocal microscope is therefore widely used in biology to study objects like cells.

7.3 Electron Microscopes

Electron microscopes bear similarity with optical microscopes. In optical microscopes, electromagnetic waves of appropriate wavelength scattered from the specimen are detected using a system of focussing lenses. In electron microscopes, electrons are used in place of electromagnetic radiation and electrostatic or magnetic lenses are used instead of glass lenses. According to wave-particle duality, electrons can sometimes behave as particles and sometimes as waves. Therefore, just like electromagnetic radiation, which can be used to image the objects, electron waves

Table 7.1 Wavelengths of electrons at some particular accelerating voltages

Applied voltage, kV	Wavelength in nm
20	0.0085
50	0.0053
80	0.0041
100	0.0037

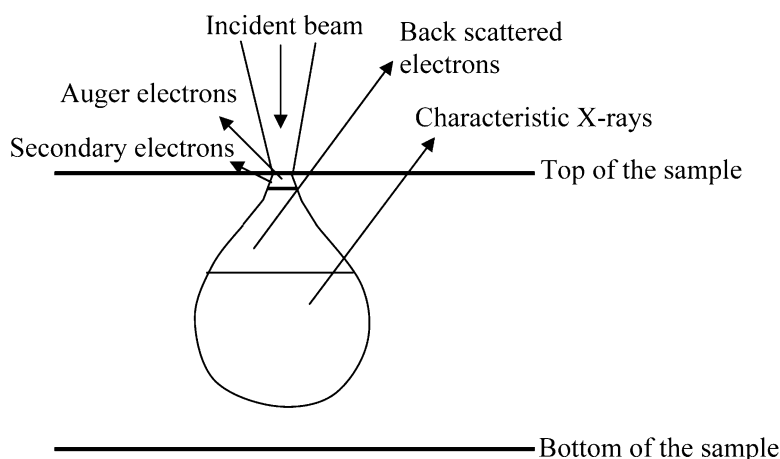


Fig. 7.5 Interaction of high energy electrons with solid

can be used to image the objects. Advantage of using electrons is that their wavelength can be tuned to a very small value, just by changing their energies so that the resolution can be increased. In Table 7.1 wavelengths of electrons at different energies are shown.

Wavelength can be simply calculated using de Broglie relation $\lambda = h/mv$ where h is Planck constant, m – mass of electron and v is the velocity of electrons. Velocity can be found out using $eV = \frac{1}{2}mv^2$, where V is the applied voltage.

Although the wavelengths shown in Table 7.1 appear to be very small and one would have expected extremely high resolution, but in general the interactions between electrons and solid are quite complicated due to charge on incident electrons and subsequent interaction with electrons and ions in solids. As shown in Fig. 7.5, this interaction results into back scattering of electrons, production of Auger electrons, visible light, UV light and even X-rays depending upon the energy of the incident electrons, type of sample and thickness of sample. In any case, even a parallel beam of electrons after interaction with solid becomes defocussed as illustrated and in general, a focussed beam therefore forms a ‘tear’ shape volume of beam interaction zone.

There are two types of electron microscopes viz. *Scanning Electron Microscope (SEM)* and *Transmission Electron Microscope (TEM)*. Scanning electron microscope uses backscattered electrons from a sample for imaging and transmission electron microscope utilizes electrons transmitted through a sample. Obviously SEM can be used to image the surface of a thick sample but TEM needs to

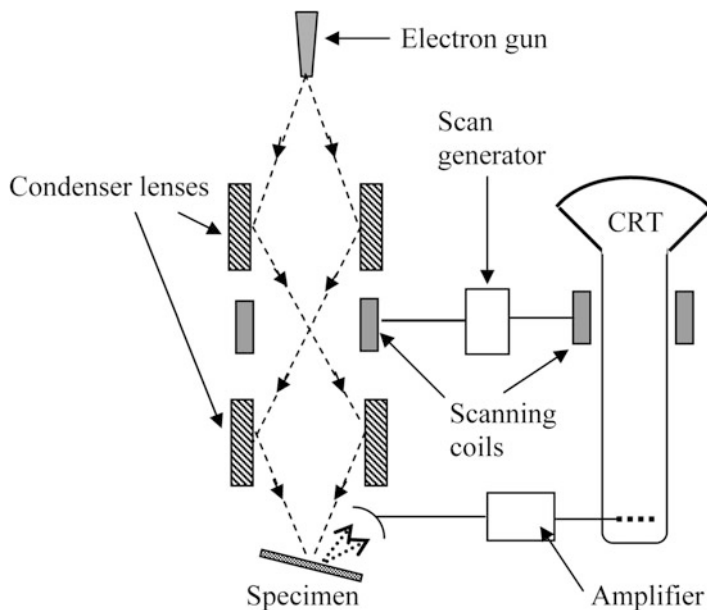


Fig. 7.6 A typical sketch of a scanning electron microscope. Electron gun, specimen and various electrodes etc. need to be mounted in a vacuum chamber

have a thin (maximum thickness ~ 300 nm) sample so that high energy electrons can transmit through the sample. Both the microscopes use electrons which need to reach the sample without getting scattered by air. Therefore, the electron microscopes need vacuum for their operation.

In the following sections, we shall briefly study both SEM and TEM.

7.3.1 Scanning Electron Microscope

In Fig. 7.6, essential parts of a scanning electron microscope are shown. In an electron microscope, electrons emitted from a hot filament are usually used. However, sometimes cold cathode (a cathode that emits electrons without heating it) is also used. A cold cathode emits electrons under the application of a very high electric field. It is also known as a field emitter. Such SEMs are known as FE-SEM and are able to give better images than hot filament SEM. However, such FE-SEM is less common than hot cathode SEM.

In a scanning electron microscope, backscattered electrons or secondary electrons are detected (in some cases it is also possible to use sample current). Due to interaction of focussed beam with solid, the backscattered electrons are somewhat defocussed resulting into lowered resolution than one would expect.

In an electron microscope, the electron beam can be focussed to a very small spot size using electrostatic or magnetic lenses. Usually the electrostatic lenses are used

for a SEM. The fine beam is scanned or rastered on the sample surface using a scan generator and back scattered electrons are collected by an appropriate detector.

Signal from scan generator along with amplified signal from the electron collector generates the image of sample surface. In order to avoid the oxidation and contamination of filament as well as reduce the collisions between air molecules and electrons, filament and sample have to be housed in a vacuum chamber. Usually vacuum $\sim 10^{-2}$ – 10^{-3} Pa or better is necessary for a normal operation of scanning electron microscope. This makes electron microscopes rather inconvenient. However some manufacturers have been successful in marketing electron microscopes known as *environmental microscopes*, in which samples can be at rather high pressure of few hundreds of Pa (100–500 Pa). Sample preparation is therefore minimized and sample in biological conditions can be investigated. For this the electrons are accelerated as usual in a high vacuum system but they enter the sample chamber through a thin foil or aperture so that a large pressure difference can be maintained (Box 7.3).

Box 7.3: Ernst Ruska (1906–1988)

Ernst Ruska was born on 25th December 1906 in Heidelberg, Germany. He had his higher education at the Technical University of Munich. Around 1927 it was getting accepted that electrons have associated waves. Simple, back of the envelope calculation shows that high velocity electrons have waves having wavelengths which are four or five orders of magnitude smaller than that of the visible light. This led Ruska around this time to think that electrons can give much higher resolution than the usual optical microscopes. He therefore built electromagnetic lenses to focus electrons and used them to develop first electron microscope in 1933. In his microscope he used a thin specimen and electrons passed through towards a photographic plate producing a high resolution image.



(continued)

Box 7.3 (continued)

In 1937 Ruska joined Siemens-Reiniger-Werke AG and marketed in 1939 the first electron microscope. He remained with Siemens company until 1955 when he became the Director for Electron Microscopy Institute in Fritz Haber Institute of Max Planck Society in Berlin. He remained there as a Director until 1972. He shared with G. Binnig and H. Rohrer the Nobel Prize for Physics in 1986. E. Ruska died on 25th May 1988.

One disadvantage of electron microscopes is that insulating samples cannot be analyzed directly as they get charged due to incident electrons and images become blurred/faulty. Therefore insulating solids are coated with a very thin metal film like gold or platinum (<10 nm) making them conducting without altering any essential details of the sample. The metal film is usually sputter coated on the sample to be investigated prior to the introduction into the electron microscope. This enables even biological samples to be analyzed using an electron microscope. Additionally, some microscopes provide with a low energy electron flood gun to reduce the sample charging effect by providing more electrons to an insulating sample.

Electron microscopes can also be used to obtain the composition of sample using a technique known as Energy Dispersive Analysis of X-rays (EDAX). The high energy electrons striking the sample produce characteristic X-rays of atoms with which they interact. When analysis of the energies and intensities of such characteristic X-rays are compared one can obtain the composition analysis of the sample under investigation (Fig. 7.7).

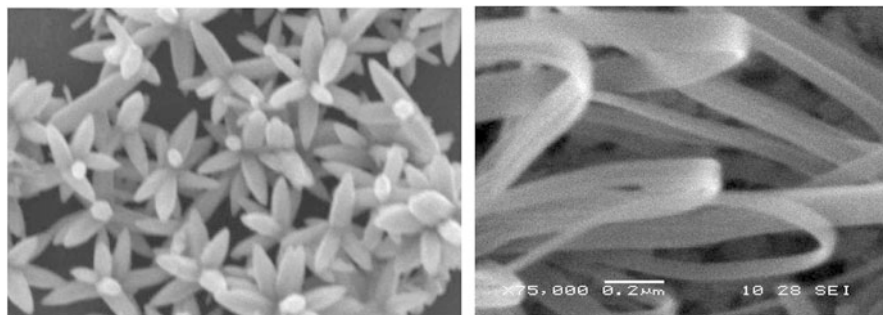


Fig. 7.7 Scanning electron microscopy images of ZnO with flower-like morphology and belt-like morphology. The chemical co-precipitation synthesis method has been used to obtain various morphologies by varying the synthesis parameters

7.3.2 Transmission Electron Microscope (TEM)

Transmission electron microscope is ideal for investigating the nanomaterials, as very high resolution is possible (better than ~ 0.5 nm) using it. As the name suggests the electrons are transmitted through the specimen in this microscope. Electrons of very high energy (typically >50 keV) are used which pass through a series of magnetic lenses, as in an optical or SEM discussed earlier. Interaction of electrons with matter was also illustrated in Fig. 7.5. The basic components of TEM are electron source, condenser lens, specimen, objective lens, diffraction lens, intermediate lens, projector lens and a fluorescent screen in the given order. There may be some additional lenses in different microscopes in order to improve the image quality and resolution. The lenses are electromagnetic whose focal lengths are varied to obtain optimized images rather than moving the lenses themselves as is done in an optical microscope. Similar to SEM, the components (and specimen) of a TEM also have to be housed in a chamber having high vacuum $\sim 10^{-3}$ – 10^{-4} Pa for its proper functioning.

As illustrated in Fig. 7.8, TEM has the advantage that one can not only obtain the images of the specimen but also diffraction patterns, which enable to understand the detailed crystal structure analysis of the sample. Using diffraction analysis one can find out size dependent changes in the lattice parameters as well as defects in the sample. Figure 7.9 illustrates an example of SiO_2 nanoparticles and their corresponding diffraction pattern. Moreover it is also possible to analyze single

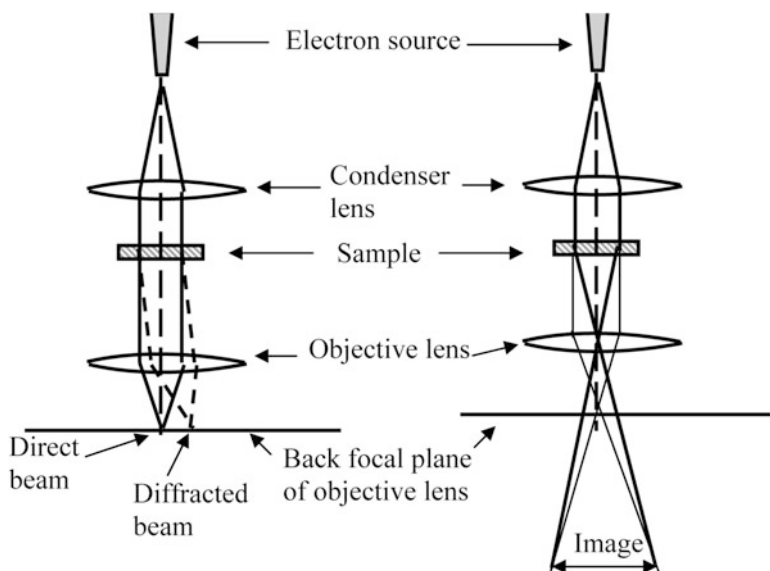


Fig. 7.8 Basic components of TEM

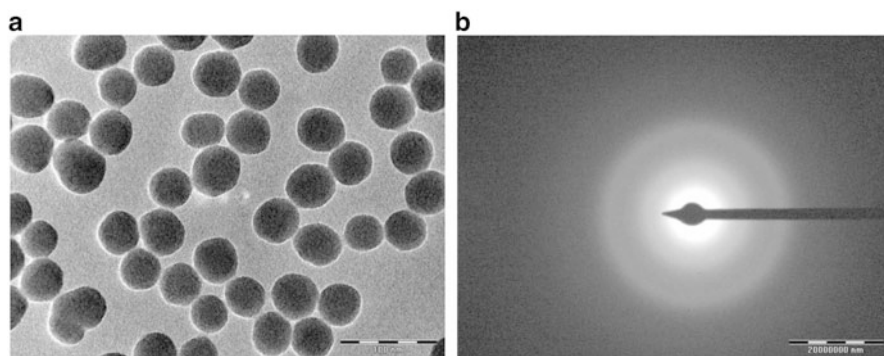


Fig. 7.9 SiO₂ nanoparticles image and diffraction pattern

particles of very small (nanometre) dimension. In some microscopes, it is possible to vary the sample temperature. This enables to investigate the problems such as size dependent melting point variation of nanoparticles (Box 7.4).

Box 7.4: Sample Preparation for Electron Microscopy

Sample preparation for electron microscopy can be quite a skillful and tedious job. For SEM, mostly sample thickness is not a problem as electrons are collected in the backscattering mode. In SEM, usually much lower electron energy is used as compared to that in TEM. Sample charging effect can be more severe in SEM and requires that insulating samples are coated with thin layer of some noble metal (Au or Pt, <10 nm) before the sample is introduced in the microscope chamber for analysis.

In case of TEM, as sufficient electrons need to be transmitted through the sample, thickness of the sample has a limitation. It becomes necessary to thin down the films (or multilayer samples) by proper grinding and/or ion etching as they might be deposited on some solid substrates. A technique like jet thinning is used for metallic samples and ion milling for non-conducting samples. Often the thin films are deposited on solid substrates. It becomes necessary to investigate then the cross sections. For biological samples a technique known as ultramicrotomy is used.

Powder samples (particle size <300 nm) can be held on some metal (usually copper) grid coated with a thin (~5 nm) carbon film or some polymer.

As both SEM and TEM require vacuum environment for their operation, sample degassing has to be avoided. Biological samples, therefore, have to be 'freeze dried' by some special process.

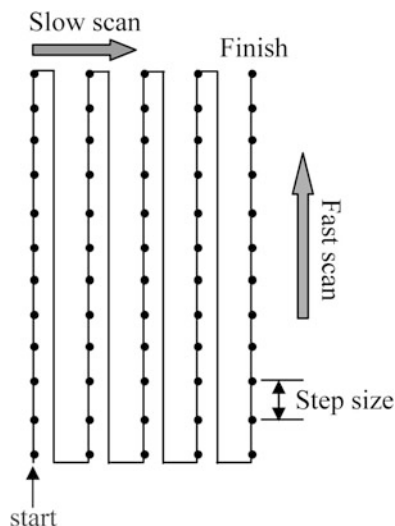
7.4 Scanning Probe Microscopes (SPM)

Scanning Probe Microscopes (SPM) is a generic name given to a family of microscopes in which a sharp tip of a metal is scanned across a sample surface in a raster mode to produce the images of samples even at subatomic resolution in some instances. The first SPM known as Scanning Tunnelling Microscope (STM) was developed around 1982 by G. Binnig and H. Rohrer for which they received Nobel Prize in 1986 along with Ernst Ruska. Subsequently many other SPMs like Atomic Force Microscope (AFM), Magnetic Force Microscope (MFM) and Scanning Near Field Microscope (SNOM) were developed to overcome some of the limitations as well as carry out 'spectromicroscopies' i.e. microscopy as well as spectroscopy combined in same instrument so that spectroscopy of same sample area is performed of which microscopy is performed. With spectromicroscopy, one is able to get not only the details of morphology and structure of a material but also know the chemical nature or electronic structure of the material and study mechanical, thermal, optical and magnetic properties too. Besides these benefits, great advantage is that unlike in electron microscopes, no special sample preparation is necessary nor vacuum is necessary unless one wants to carry out analysis of a clean surface. One can even use liquid environment for these microscopes. Thus it has been possible to investigate not only insulating but live or biological samples. Some of these microscopes and their modified versions can also be used for doing lithography and will be discussed in Chap. 9.

Here we shall discuss STM, AFM and SNOM, quite commonly used microscopes for nanotechnology work. In all these microscopes, scanning probe and raster principle are common. Probe is a fine metal tip of ~ 10 nm diameter. Tip materials are Si, Pt-Ir or Pt-Rh. Even diamond film coated tips are used. Tips are obtained by etching a fine metal wire in some suitable chemicals or prepared lithographically. In case of microscopes like STM the tip is directly mounted on some specially designed piezo drive (or piezo tube). For AFM investigations, tip is mounted on a cantilever which is then mounted on a piezo drive. Function of a piezo drive is to scan the sample surface to be imaged. Materials like lead zirconium titanate (PZT) are known to be piezo crystals.

SPM scans are made over few nm to $100\text{ }\mu\text{m}$ in the horizontal plane (x-y) and about few nm to $10\text{ }\mu\text{m}$ in vertical plane (z). As shown in Fig. 7.10, starting at one point the scans are made by moving the tip on a line and then moving it to the next line. Digital images are collected at several points (pixels) like 64, 512 or 1,024 per line. Distance between two pixels determines the step size. Scanning is usually fast in one line going slowly to the other line. An image of the predetermined surface area is usually acquired in few minutes. Piezo drives along with the tip are very crucial in determining the resolution of the acquired images using SPM.

In order to achieve very high resolution like atomic level resolutions using SPM, it is essential that the microscopes be shielded from mechanical vibrations. As the

Fig. 7.10 Raster scan

currents (or forces) involved are very small, influence of external magnetic fields as well as electrical noise need to be avoided. Considerable efforts are usually made to avoid external disturbances by isolating the SPM using anti-vibration platform as well as using properly shielded electronics.

7.4.1 Scanning Tunnelling Microscope

As the name suggests, the scanning tunnelling microscope is based on the tunnelling principle (see Chap. 1). When two metals say M_1 and M_2 are brought at small distance (but larger than 10 nm) as depicted in Fig. 7.11a, even though their Fermi levels do not coincide, transfer of electrons from one metal to the other is not possible. To transfer electrons from one metal to the other, it is necessary for the electrons in the vicinity of the Fermi level to overcome the potential barrier known as the work function of the material. Typically, the work functions of metals are few electron volts (2–5 eV) and transfer of electrons at room temperature is forbidden. However, the metals brought in extremely close distance of the order of a few nanometres (usually less than 10 nm) behave differently. Electrons as shown in Fig. 7.11b can be transferred from one metal to the other to establish a common Fermi level without going over the potential barrier, set by the work function. At short distance of few nanometres, the wavefunctions of electrons from either side decay into the other metal. In other words, electrons can ‘tunnel’ from one metal to the other to occupy state of lower energy. This causes Fermi levels of the two metals to coincide with a small ‘contact potential’. This reduces the barrier heights but changes are still small and barriers are sufficiently large for

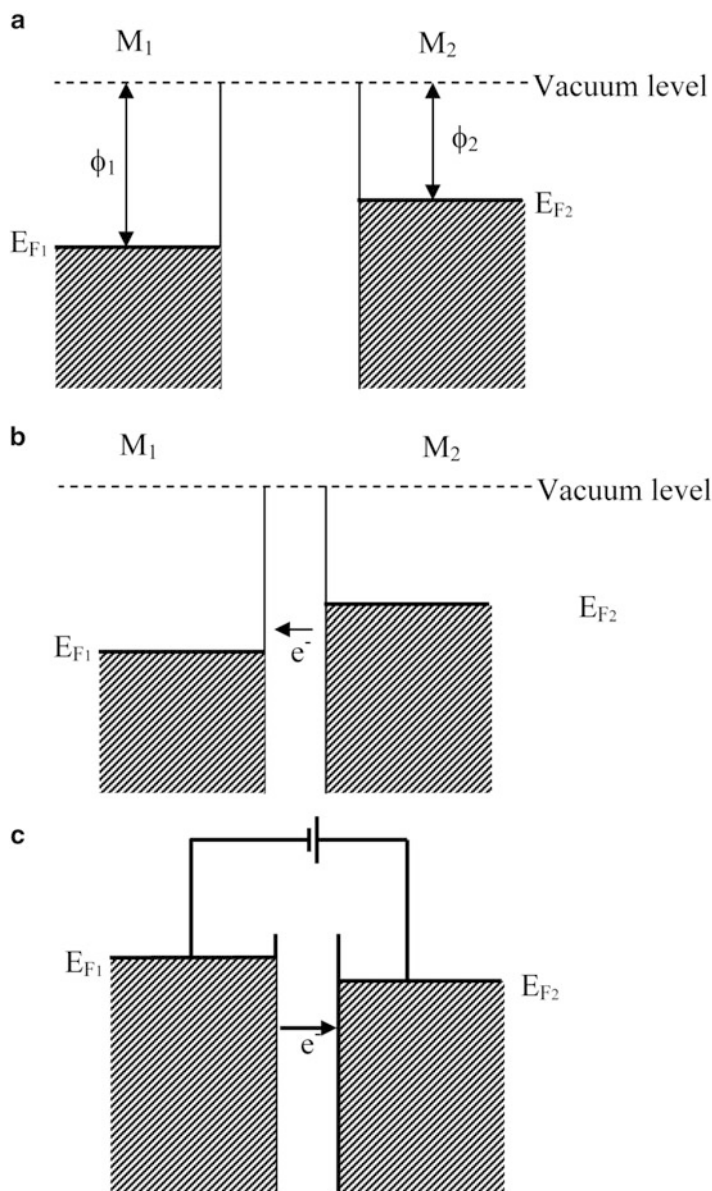


Fig. 7.11 Tunnelling of electrons from one metal to other. (a) Metals are at small distance, but not less than 10 nm. (b) Metals are in close contact with each other, at a distance less than 10 nm. (c) Potential is applied between two metals

electrons to overcome them. Once the Fermi levels coincide, the electrons cannot flow from one metal to the other. However, by raising the Fermi level of one metal with respect to the other, electrons can tunnel from one metal to the other, as shown in Fig. 7.11c.

The energy required by electrons to overcome the energy barrier is still very high and not obtained by applying the potential, but electrons can tunnel. The tunnelling probability or current depends upon the availability of the empty states in metal in which electrons flow (density of empty states) and distance between the two metals. Fermi level positions can be altered by applying a small voltage ($V < \phi$) between the two metals. The metal (M_1) which is connected to the negative terminal of the power supply has raised Fermi level with respect to the other metal (M_2) whose Fermi level is lowered. This is made use of in an STM. As illustrated in Fig. 7.11c, the tip potential is made negative, therefore its Fermi level is raised and current flows from tip to the sample. Indeed it is possible to raise Fermi level of sample higher than the tip, so that electrons flow from sample to the tip. It is then quite obvious that by lowering the sample with respect to the tip and measuring the current flowing towards the sample, we are able to probe unoccupied states or empty energy levels of the sample. If the sample Fermi level is at higher level, electrons below Fermi level flow to the tip. Therefore, one can know about occupied states in the sample. Thus, STM is capable of performing even spectroscopy of occupied and unoccupied levels.

The tunnelling current is given as

$$I = C \exp(-kd) \quad (7.5)$$

where C is proportionality constant, d – the distance between two metals and k is known as decay constant. It is obtained as

$$k^2 = 8\pi^2 m \frac{(V_b - E)}{h} \quad (7.6)$$

where V_b is the difference between the Fermi level positions of the tip and the sample. E is the energy of the level to which electron tunnels.

Usually the distance between tip and the sample is between 0.5 nm to 1 nm and current of few picoamperes (pA) to nanoamperes (nA) is expected. It is quite clear from the Eq. (7.5) that current is very sensitive to the distance between tip and sample.

An STM can be operated in two different modes viz.

1. Constant current mode
2. Constant height mode

Constant current mode: Probe in the form of a sharp metal tip is moved slowly on the sample surface so that the current between the tip and the sample remains constant. In order to maintain the constant current between the tip and the sample, distance between the tip and the atomic corrugations also needs to be kept

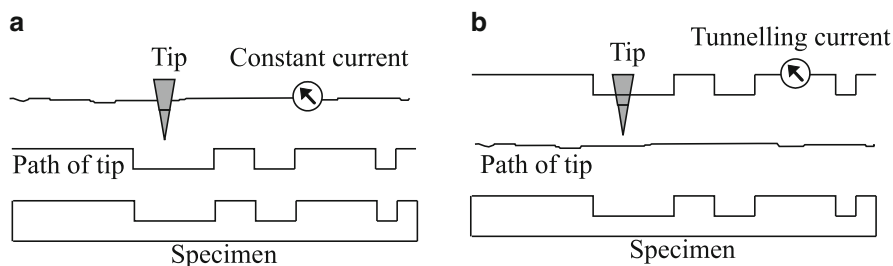


Fig. 7.12 (a) Constant current and (b) constant height modes

constant (see Fig. 7.12a). Thus the tip will have to follow the atom contours. By successively scanning the desired sample area in a raster mode, profile of surface atoms can be generated as an image, which is really the movement of the tip or the probe in an attempt to keep constant current between the sample and the tip, controlled by a proper feed-back loop. This is known as constant current mode.

Constant height mode: Alternatively, the tip can be moved on the sample surface at a constant height (typically >0.5 nm) as illustrated in Fig. 7.12b. As there is a relation given by Eq. (7.5), between current and the distance, a surface profile can be generated from the variations observed in the tunnel current. Thus the image is the replica of the variation of current as the tip scans the desired area of the sample surface. Advantage of the constant height mode as compared to the constant current mode is that the tip can be moved faster on the sample surface, as there is no necessity of the feed-back circuit. Besides it is dangerous to move the tip close to the sample in constant current mode, as that can occasionally hit some rough hillocks of the sample and get destroyed. This is avoided in the constant height mode and tip can be moved faster. However this would be at the cost of better sensitivity in the constant current mode.

Major limitation of STM is that the tunnelling current has to flow between the sample and the probe. Although the current is very small (of pico ampere order), it can be detected. However, in case of insulating samples, even this much current is not possible. Therefore, realizing this problem other scanning probe microscopes were developed (Fig. 7.13).

7.4.2 Atomic Force Microscope

Limitation of an STM is that it requires the sample to be a conductor or at least a semiconductor. This limitation was immediately realized and overcome by introducing Atomic Force Microscope (AFM) by G. Binnig and C. Gerber in 1985, quite immediately after the development of first STM. As discussed in Chap. 2, when two



Fig. 7.13 Photograph of a JSPM-5200 scanning probe microscope

atoms are close to each other, there are attractive and repulsive forces which depend upon the distance of their separation. Combined force is given by the equation

$$F = \frac{A}{R^{12}} - \frac{B}{R^6} \quad (7.7)$$

where F is resultant force between two atoms, A and B – constants and R is distance between two atoms.

The first term is the repulsive force and the second term is attractive force between two atoms. It can be seen that repulsive force is more effective at very short distance and changes rapidly with distance. This is due to repulsive interaction between the electron clouds at a short distance (Pauli exclusion principle).

AFM has a flexible cantilever $\sim 100 \mu\text{m}$ long, $10 \mu\text{m}$ wide and $1 \mu\text{m}$ in height attached to a piezo drive. A tip is mounted on cantilever as shown in Fig. 7.14 which can be brought close to sample surface. For cantilever

$$F = K \cdot \delta z \quad (7.8)$$

where F is force experienced by the cantilever, K is related to the natural resonance frequency of the cantilever (spring constant) and δz is the displacement of the cantilever.

Resonant frequency of cantilever

$$(\omega_r) = \sqrt{\frac{K}{m}} \quad (7.9)$$

where ω_r is resonant frequency and m is mass of the cantilever.

If there is a gradient in force with distance,

$$F = F_0 + \left(\frac{\delta F}{\delta z} \right) \delta z \quad (7.10)$$

and

$$F_0 = \left(K - \frac{\delta F}{\delta z} \right) \delta z \quad (7.11)$$

Thus, the effective spring constant ' K ' changes in the presence of gradient. Resonant frequency also correspondingly changes as

$$(\omega_r) = \sqrt{\frac{\left(K - \frac{\delta F}{\delta z} \right)}{m}} \quad (7.12)$$

The resonant frequency is used to control the tip-sample interaction.

AFM tip, in close vicinity of the sample surface, experiences a repulsive force which results into minute amount of bending of the cantilever. A laser beam is directed on back of the cantilever (see Fig. 7.14). Small deflections caused by the tip-sample interaction are recorded by a position sensitive photodiode. By rastering the probe on sample surface and measuring the cantilever deflections, surface image is obtained.

An AFM can be operated in three different modes viz. (1) Contact mode, (2) Non-contact mode and (3) Tapping mode.

Contact mode: In this case, the tip is in contact with the sample surface and is almost forced into it. However due to repulsive interaction between electron charge

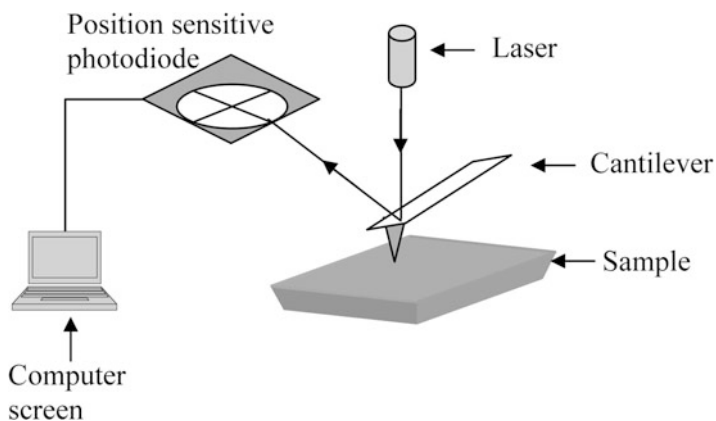


Fig. 7.14 Schematic of an Atomic Force Microscope

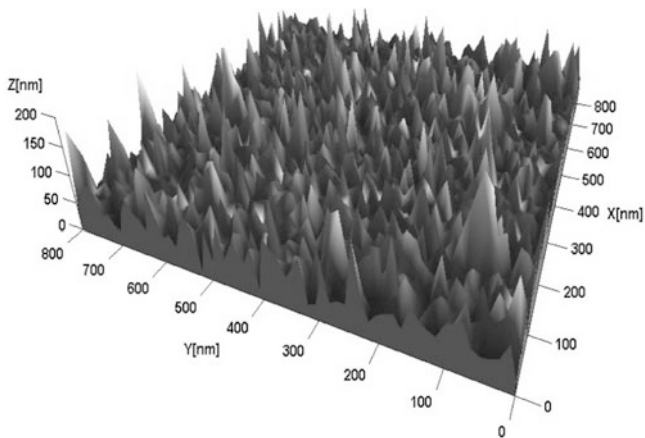


Fig. 7.15 An atomic force microscopy image of TiO_2 thin film deposited on silicon substrate

cloud of the tip atom and that of the surface atom, the tip is repelled back which bends the cantilever and deviates the direction of the laser beam. In this mode the interaction due to the first term on right hand side of Eq. (7.7) is dominant due to very small value of ' R ', the distance between two atoms. The main disadvantage of this mode is that the tip or sample can get damaged due to forcing of the tip into sample, especially, polymers or other organic materials like biological samples which can get damaged by this method.

Non-contact mode: In non-contact mode, the tip or the probe moves at some small distance away from the sample surface. Therefore, it cannot damage the sample. In this mode the second term on right hand side of the Eq. (7.7) is the dominant term. This term arises due to polarization of interacting atoms and is due to dipole-dipole interaction of two atoms (Fig. 7.15).

Tapping mode: Tapping mode is a combination of contact and non-contact modes. The resolution in contact mode is higher than that due to non-contact mode, because in contact mode the interaction between tip and surface atoms is much more sensitive to the distance as compared to that in non-contact mode. With tapping mode, high resolution advantage of contact mode and non-destructiveness of non-contact mode are achieved. The tip is oscillated in the vicinity of the surface at a distance of ~ 50 nm in such a way that it nearly touches the sample during its cycle of oscillation. Tapping mode is simple and robust to use.

7.4.3 Scanning Near-Field Optical Microscope (SNOM)

The imaging in conventional optical microscope is based on the principle of interference of light waves. For this the object is first illuminated and then the scattered light by the object is collected so as to form the magnified image. The image of an object is reconstructed in the image plane.

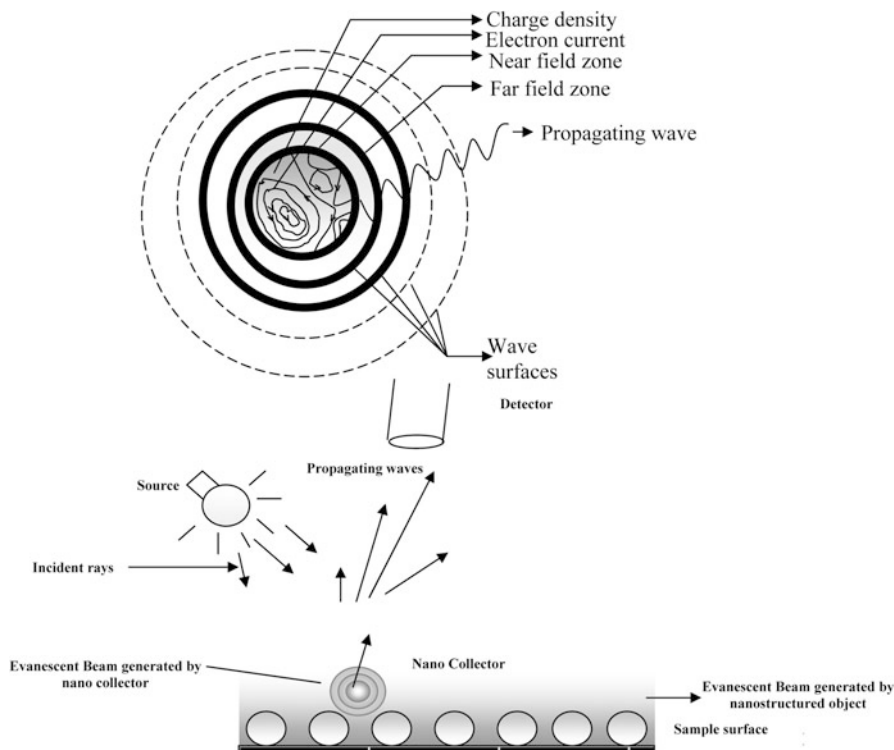


Fig. 7.16 Field of an object is shown. The electron currents and charge densities inside the object induce an electromagnetic field radiating from the surface. Far away from the surface, the field has the structure of propagating waves. Very close to the object, the field has a complex structure as it is composed of propagating and non – radiating components. Using a nanocollector i.e. an object like fibre with small diameter, changes in near field can be detected

In an optical microscope, the detector is kept at a very large distance of at least a few mm to cm from the sample. This is very large compared to wavelength of visible light normally used. For example if the wavelength of light used is ~ 500 nm even the distance of 5 mm would be 10,000 times the wavelength of light.

It is known for a long time that the reflected light or the light leaving from a luminous object has two components viz. *evanescent* beam and *propagating* beam. This is shown schematically in Fig. 7.16. Evanescent beam is related to what is known as *near-field* and propagating beam is related to *far-field*. Near-field extends to very small distance and evanescent beam does not have a propagating wave nature. However photons cannot be trapped and must escape as propagating waves. This in fact is quite advantageous. If the near-field zone is disturbed, it also affects the far-field and propagating waves. This idea is used to overcome the diffraction limit and obtain a high resolution using scanning near-field optical microscope.

As shown in Fig. 7.17, a special probe with very small aperture of few nm diameter is brought very close to the sample surface.

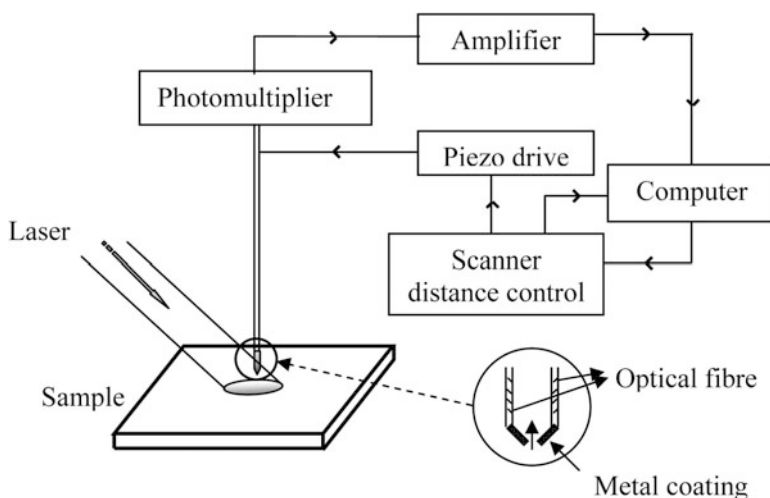


Fig. 7.17 SNOM probe and schematic of surface scan

The diameter of the aperture in the fibre as well as distance between the aperture and sample has to be smaller than the wavelength of light. Under such conditions, light leaves the sample before diffracting.

It can be readily appreciated now that the resolution obtained in a SNOM will depend upon the size of the aperture and the distance at which the probe can be placed. Thanks to the technological developments like availability of intense source like laser and fine movement of the probe using piezoelectric drive that high resolving power microscopy using SNOM is a reality. Very fine optical fibres are tapered to the diameter of less than 100 nm routinely and coated with some metal like aluminium. Metal coating of the fibre aperture is necessary because narrower is the aperture greater are the chances that the light can escape from the sides of the aperture wall. With current technology, apertures smaller than 100 nm can be routinely made out of glass fibre, achieving resolution even as good as $\sim 20\text{--}40$ nm. In few cases resolution better than even 20 nm also has been achieved. Attempts are made to even overcome the limit due to aperture dimension. An Apertureless Near-Field Scanning Optical Microscope (ANSOM) has also been proposed. It is experimentally shown that ANSOM can resolve objects as small as 1 nm (Box 7.5).

Box 7.5: Photon Tunnelling

A scanning near-field optical microscope (SNOM) is sometimes also referred to as a photon tunnelling microscope. This idea is similar to that in STM. In an STM, electrons tunnel between tip and the sample under investigation. In order to observe the tunnelling current and use for imaging the distance

(continued)

Box 7.5 (continued)

between the probe and the sample has to be very small ($<1\text{--}2\text{ nm}$). By analogy, one can consider that if the distance between the sample and the probe is very small the photons tunnel. In an STM exponentially decaying electron tunnelling is sensed. In SNOM, photons leaving the exponentially decaying near-field are detected. Therefore, it is quite reasonable to think of SNOM as ‘Photon Scanning Tunnelling Microscope (PSTM)’.

In Table 7.2, a comparison is made of various microscopes described in this section. Magnification and resolution numbers are given to indicate typical expected values and can vary from model to model (Boxes 7.6 and 7.7).

Table 7.2 Comparison of various microscopes

Microscope (radiation/interaction used)	Magnification	Resolution
Human eye (visible light)	–	100 μm
Optical microscope (UV-Vis-IR)	10^3	0.1 μm
Scanning electron microscope (electrons)	$10^5\text{--}10^6$	3 nm
Transmission electron microscope (electrons)	$>10^6$	0.1 nm
Scanning tunnelling microscope (electron current)	$>10^6$	0.1 nm
Atomic force microscope (e-e repulsion)	$>10^6$	0.1 nm
Scanning near-field optical microscope (vis light)	10^5	20–50 nm

Box 7.6: History of SNOM

As early as in 1928, E.H. Synge wrote a short note to the *Philosophical Magazine* (Volume 6, Dec 1928, p. 356) suggesting a method for extending the resolution of an optical microscope to about $0.005\text{--}0.004\text{ }\mu\text{m}$. He mentioned that the method was suggested to him by a distinguished physicist (without mentioning the name). He, however, wrote about the difficulties in having intense source of light, adjusting distance between the sample and the source close to $\sim 10^{-6}\text{ cm}$ and creating the aperture for light source as small as $\sim 10^{-6}\text{ cm}$ diameter. It can be appreciated that in the absence of lasers and piezoelectric devices at that time, the practical realization of SNOM was not possible.

Ash and Nicholls (*Nature* 237 (1972), 510) demonstrated the resolution which was not limited by incident wavelength. However, they used microwaves to demonstrate the principle of near-field microscopy.

(continued)

Box 7.6 (continued)

Success of the near-field concept in microwave region led Pohl and others (*Appl. Phys. Lett.* 44 (1984), 651) to demonstrate SNOM in the visible light range. They could produce the resolution as high as 25 nm with a wavelength of 488 nm.

Box 7.7: Optical Stethoscope

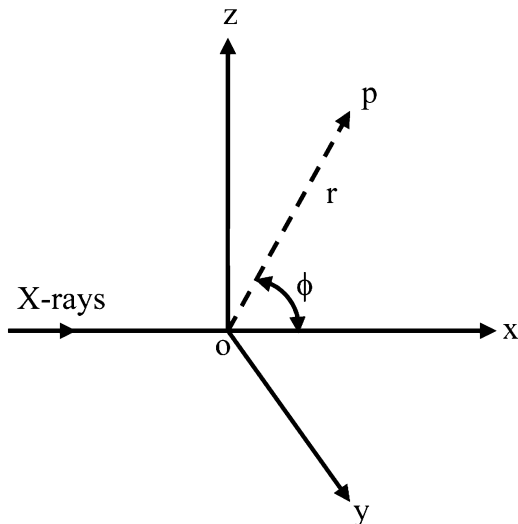
Pohl et al. (*Appl. Phys. Lett.* 44 (1984), 651) have compared the SNOM with the doctor's stethoscope. A doctor can locate the position of the heart to within less than 10 cm by moving the stethoscope over the patient's heart and listening to the sound of heart beat. If the frequency of heart beating can be considered to be in the range of 30–100 Hz, its wavelength would be in the range of almost 100 m. Thus the stethoscope provides a resolving power of $\sim \lambda/1000$. Such a high resolution has become possible because of the smallness of the probe (the end of stethoscope) as well as its placement at a distance much smaller than the wavelength (100 m) from the object (heart) to be examined.

The concept of stethoscope can be applied for other wavelengths and other types of waves. An optical stethoscope which allows image recording with sub wavelength resolution was further demonstrated by Pohl et al. in their paper.

7.5 Diffraction Techniques

We discussed in the previous section some microscopy techniques. These techniques are useful to understand the surface morphology of the samples under investigation. In cases where atomic resolution is achieved one can know about surface structures too. However, knowledge of bulk structure is normally not possible using these techniques. Diffraction techniques using electrons, X-rays or neutrons produce information about crystal structure, and are used to understand structure (Bravais lattice) of bulk materials and can be extended to investigate nanomaterials. However the usual diffraction analysis relies on the long range periodic arrangement of atoms/molecules. For larger nanoparticles (>20 nm or so, sufficient long range order is established). While investigating very small particles (<20 nm) special precautions need to be taken as will be discussed from time to time. We shall discuss here only the X-ray diffraction technique as electron or neutron diffraction are similar to some extent and can be very well understood once X-ray diffraction is studied.

Fig. 7.18 Thomson's explanation of scattering of X-rays by electrons located at origin. Intensity is measured at point P



7.5.1 X-Ray Diffraction (XRD)

We are able to see objects around us as light is scattered by the objects and enters our eyes. Similarly, X-rays scattered by atoms enable us to understand about arrangement of atoms in solids. As early as in 1912 von Laue postulated that if X-rays are waves and distances between atoms in solids are comparable to wavelength of X-ray, then they should be diffracted by atoms in solid.

When we say that X-rays are scattered by atoms we really mean that they are scattered by electrons of atoms. It was shown by J.J. Thomson that if unpolarized X-rays are incident on free electrons, located at the origin O (see Fig. 7.18), the scattered intensity at point P is given as

$$I_e = \frac{nIe^4}{2r^2m^2c^4} (1 + \cos^2\phi) \quad (7.13)$$

where I_e is intensity of X-rays at point P, e – charge of an electron, I – intensity of the X-ray beam falling on a free electron at the origin, r – distance of point P from the origin, n – number of electrons, m – mass of an electron, ϕ – angle between incident beam and scattered beam and c is the velocity of light.

Above equation indicates that I_e is proportional to $1/m^2$. It is easy to convince therefore that nucleus, which is much heavier than an electron will not be able to scatter X-rays. It is also clear that most of the beam will be scattered in the forward direction and the intensity in other directions would be very small. Intensity also reduces drastically away from the scattering electrons as $1/r^2$.

7.5.2 Atomic Scattering Factor

It is also evident from Eq. (7.13) that more is the number of electrons (n), more would be the scattered X-ray intensity. An atom with atomic number Z has Z electrons in it. However in atoms we do not have the situation assumed by Thomson viz. free electrons. The scattered intensity, however, is proportional to Z , X-ray wavelength and angle of scattering. This is known as atomic scattering power.

Scattering power of an atom is referred to as atomic scattering factor f and is defined as

$$f = \frac{\text{Amplitudte of wave scattered by an atom}}{\text{Amplitudte of wave scattered by a single electron}} \quad (7.14)$$

It can be shown that

$$f = \frac{z \sin \theta}{\lambda} \quad (7.15)$$

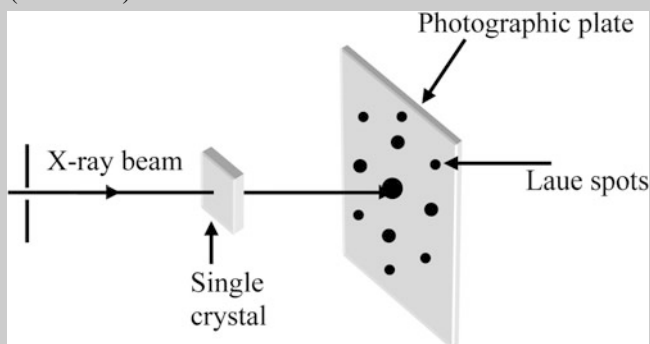
where Z is atomic number, θ – half the angle between incident beam and scattered beam direction and λ is wavelength of X-rays.

Further, it is the interference of scattered rays that is important. If the scattered rays are either in or out of phase, their amplitudes will either add or cancel each other. If, however, there is no phase relation between the atoms then intensities would simply add. Consider, for example, an ensemble of N atoms of a gas in a closed vessel. As the atoms of a gas are in continuous motion, without any regular arrangement of atoms, each atom scatters X-rays but there is no definite phase relation amongst them. Scattered intensity from an atom is proportional to square of amplitude of the rays. If A is the amplitude of the incident rays, intensity scattered by each ray would be A^2 . As there are N atoms, total intensity scattered by gaseous atoms would be NA^2 . On the other hand if there is a correlation between the phases of N atoms each scattering wave of amplitude A , then amplitudes would add up to give a total amplitude of value NA . The intensity would be given by $(NA)^2$. Thus the intensity would be N times larger than that from uncorrelated atoms (Box 7.8).

Box 7.8: Laue Pattern

Laue with his collaborators placed in front of a single crystal of CuSO_4 a photographic plate and allowed a narrow beam of X-rays to pass through the crystals as illustrated in Fig. 7.19. He observed well defined diffraction spots on the photographic plate. Today this kind of arrangement is still used, replacing the photographic plate with X-ray detector, to study the crystals and is known as ‘Laue Method’.

(continued)

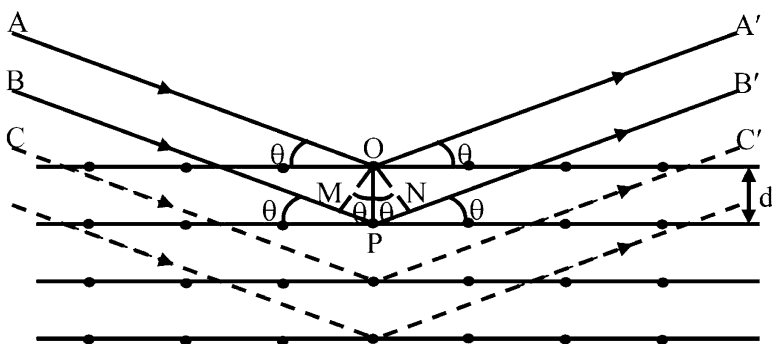
Box 7.8 (continued)**Fig. 7.19** Arrangement to obtain Laue photograph

In the same year (1912), as Laue performed diffraction experiments, W.H. Bragg and his son W.L. Bragg studied structures of a number of crystals like NaCl, KCl, KBr etc. to verify the Laue method.

W.L. Bragg also offered a simple geometrical explanation of observed diffraction pattern. This is now known as 'Bragg condition' or 'Bragg Law' and is discussed here.

7.5.3 Bragg's Law of Diffraction

Bragg considered that a beam of X-rays falls on crystal planes at some grazing incidence, θ , as shown in Fig. 7.20.

**Fig. 7.20** Geometrical explanation of X-ray diffraction

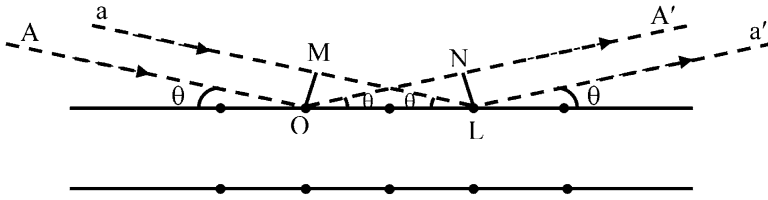


Fig. 7.21 Diffraction from the surface atoms

Beam of parallel rays is assumed. Consider that parallel rays AO and BP are incident on atoms at O and P respectively making an angle θ . Rays are scattered as OA' and PB'. Distance between O and P is ' d ' (distance between the consecutive planes).

The path difference between AOA' and BPB' is

$$\text{Path difference} = MP + PN = 2 OP \sin \theta = 2 d \sin \theta \quad (7.16)$$

If λ or multiple of λ is the path difference between ray AOA' and BPB', they will reinforce or interfere constructively.

$$\text{Therefore, } n\lambda = 2d \sin \theta, \text{ where } n = 1, 2, 3 \dots \quad (7.17)$$

n is known as order of diffraction. Equation (7.17) is Bragg diffraction condition. In most of the cases only $n = 1$ is considered.

Consider now, parallel rays AO and aL falling on the first row of atoms at a grazing angle θ and scattered rays OA' and La' as shown in Fig. 7.21. The rays OA' and La' are also parallel to each other and make grazing angle with the surface. OM and LN are perpendiculars on aL and OA' respectively.

$$\text{Path difference} = ML - ON = OL \cos \theta - OL \cos \theta = 0 \quad (7.18)$$

Thus OA' and La' scattered from the first plane are able to interfere constructively as they are in phase. Thus all atoms at the surface receiving radiation at angle θ will interfere constructively.

Consider now the rays AOA' and CQC' diffracted from 1st and 3rd planes respectively as in Fig. 7.22. Between AOA' and CQC'

$$\text{Path difference} = M'Q + QN' = 2d \sin \theta + 2d \sin \theta = 4d \sin \theta$$

In order that rays reinforce constructively, they should be multiples of wavelength. If path difference is 2λ ,

$$2\lambda = 4d \sin \theta \Rightarrow \lambda = 2d \sin \theta$$

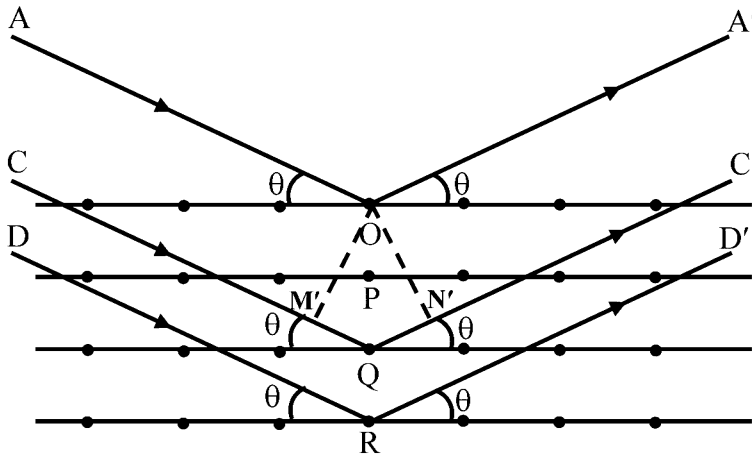


Fig. 7.22 Rays scattered from different planes, in same direction make constructive interference

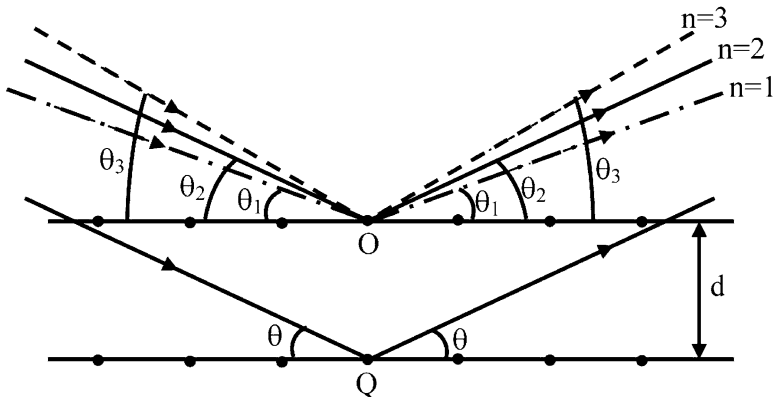


Fig. 7.23 Given set of planes can diffract for same wavelength of X-rays in different directions to satisfy Bragg law

For rays AOA' and DRD' have to interfere in same direction, path difference has to be 3λ .

i.e. $3\lambda = 6d \sin \theta \Rightarrow \lambda = 2d \sin \theta$ and so on.

Thus rays scattered by those in different planes in same direction will all reinforce to give constructive interference.

Consider now a situation as in Fig. 7.23, if λ and d are fixed, for angle θ_1 , $\lambda = 2d \sin \theta_1$, will be satisfied for $n = 1$ i.e. path difference is just λ .

But for another angle θ_2 , path difference can be 2λ . Still constructive interference can take place but in another direction. Thus

$$2\lambda = 2d \sin \theta_2 \quad (n = 2) \quad (7.19)$$

For yet another angle θ_3 the condition

$$3\lambda = 2d \sin \theta_3$$

must be satisfied.

Consider following numerical example

Let $\lambda = 1.5 \text{ \AA}$ and $d = 3 \text{ \AA}$.

For $n = 1$ and angle θ_1 , using Bragg's law, we get, $\theta_1 = 14.50^\circ$.

$$\text{For } n = 2, \theta_2 = 30^\circ$$

Thus keeping λ and d constant we can get diffraction in different orders (n) from a set of planes at angles $\theta_1, \theta_2, \theta_3 \dots$. Only condition will be $\theta < 90^\circ$.

7.5.4 Diffraction from Different Types of Samples

So far we considered the diffraction from a set of planes, assuming them to be infinite in number. We may have samples which are single crystalline, polycrystalline solids, liquids or even gases. Amorphous solids, liquids, gases and nanocrystalline materials do not have infinite ordered arrangement of atoms but have few atoms and few planes (small single crystals or crystalline nanomaterials). Nanoparticles with disordered structure can be treated like amorphous bulk solids. In Fig. 7.24 the differences in scattering of X-rays by different materials are shown.

It can be seen from Fig. 7.24 that in case of a monoatomic gas (like He, Ar, Kr etc.) forward scattering occurs without any diffraction peaks in any other direction. In case of liquids or amorphous solids one or two peaks at angles other than the angle of incidence appear due to short range order in these materials. For single crystals, diffraction peaks appear at various angles. Intensities of peaks depend upon atomic scattering factor as well as crystal structure (or form) factor which we shall consider in the following section. The diffraction peaks from ideal single crystals are sharp, and broadened to certain extent only by instrumentation factor. However in case of polycrystalline sample the peaks are broadened due to the size of the grains. All the grains in crystal may not be of the same size. Therefore the width of the diffraction peaks can be considered as the effect of convolution of different peaks giving the average grain size. The diffraction peaks are broadened in case of nanoparticles also due to small number of atoms and planes present in them. We shall discuss this in a following section and show how one can determine the average particle size using the widths of diffraction peaks.

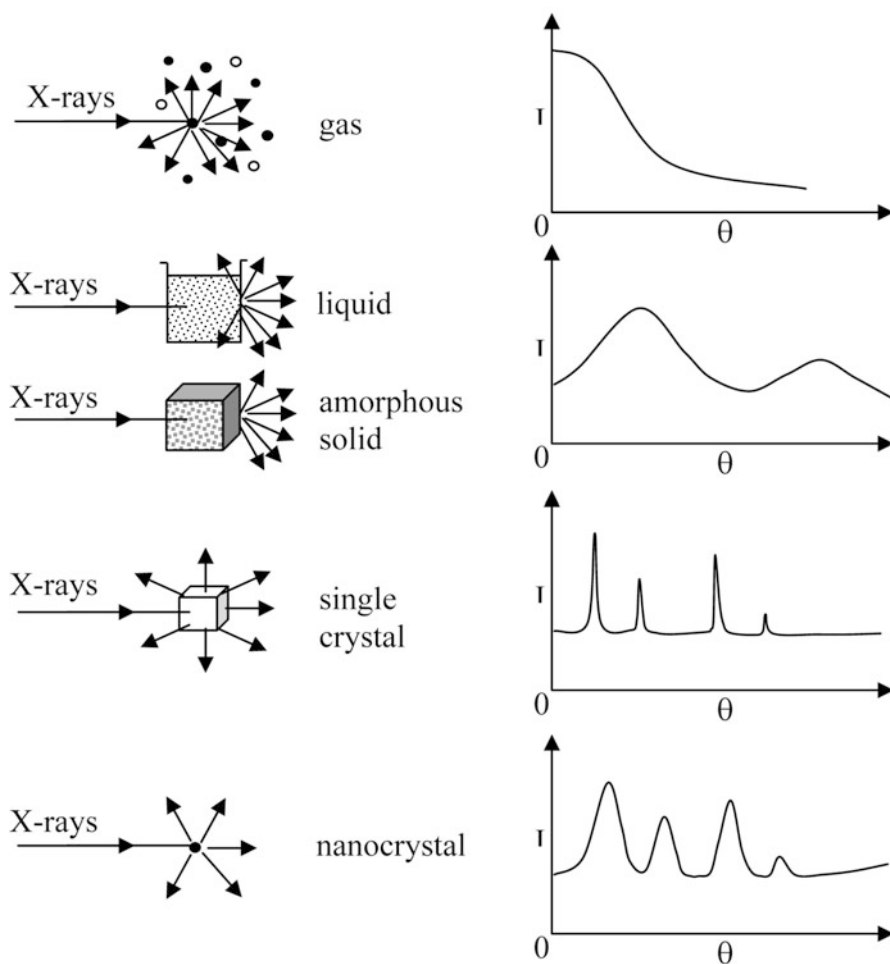


Fig. 7.24 Diffraction from gases, liquids, solids and nanocrystals

7.5.5 Crystal Structure Factor

The intensity of diffracted X-rays from different planes even from a perfect single crystal cannot be same. Different planes in unit cell of a crystal have different number of atoms. Due to specific arrangement of atoms in crystals (characteristic of material), it so happens that the X-rays scattered from different atoms may not be in phase. Therefore intensities of X-rays from different planes can be shown to be given by what is known as crystal structure or crystal form factor F . Crystal structure factor is also related to atomic scattering factor f . Crystal structure factor is defined as follows.

$$|F| = \frac{\text{Amplitude of waves scattered from all atoms in a unit cell}}{\text{Amplitude of waves scattered by an electron}} \quad (7.20)$$

$$= \sum_P^N f_P e^{2\pi i(hu + kv + lw)} \quad (7.21)$$

where sum is over all atoms, P to N , u , v and w denote the coordinates of atoms in unit cell and h , k and l are Miller indices.

Intensity of a diffraction peak is $|F|^2$.

It can be easily verified that in some unit cells systematic absence or enhancement of diffraction intensities from different planes can occur due to crystal structure factor.

7.5.6 Diffraction from Nanoparticles

We know that in a nanoparticle number of atoms is very small. Nanoparticles cannot be considered as an infinite arrangement of atoms as is usually assumed in Solid State Physics, in order to determine various properties of solids. In case of amorphous nanoparticles, similar to an amorphous bulk solid material, broad diffraction peaks are expected to occur. However in case of nanoparticles in which atoms do have ordered lattices some changes in diffraction are to be expected as compared to single crystal or polycrystalline bulk solid diffraction (often the nanoparticles are similar to single crystals and do not have grain boundaries). It has been found that the diffraction peaks in nanocrystalline particles are broadened compared to single or polycrystalline solid of the same material. This effect can be understood as follows.

Consider a limited number of planes say $0, 1, \dots, m$ or total $m + 1$ planes in a set. Note that here first plane is numbered as 0. Let d_{hkl} be the interplaner distance. If the thickness of the crystal is say t then $t = m d_{hkl}$. Let the X-rays of a single wavelength λ fall on the set of planes d_{hkl} . All the rays in a beam may not be exactly parallel to each other. As shown in Fig. 7.25, rays AO, BP, $\dots$, LS only are parallel and make angle θ_B with parallel planes (hkl). Let these rays scatter coherently and interfere constructively to satisfy Bragg's diffraction condition. This should produce an extremely sharp peak widened only due to instrument limitation. However ray like CO is at slightly larger angle than θ_B and so is the ray MS striking the m th plane. Both CO and MS are parallel to each other. It can be understood that the ray like CO cannot interfere constructively with rays parallel to it just in 1st, 2nd etc. planes. But there would be a plane say m at which it interferes constructively (corresponding to thickness of crystal). Therefore

$$(m + 1)\lambda = 2T \sin \theta_1 \quad (7.22)$$

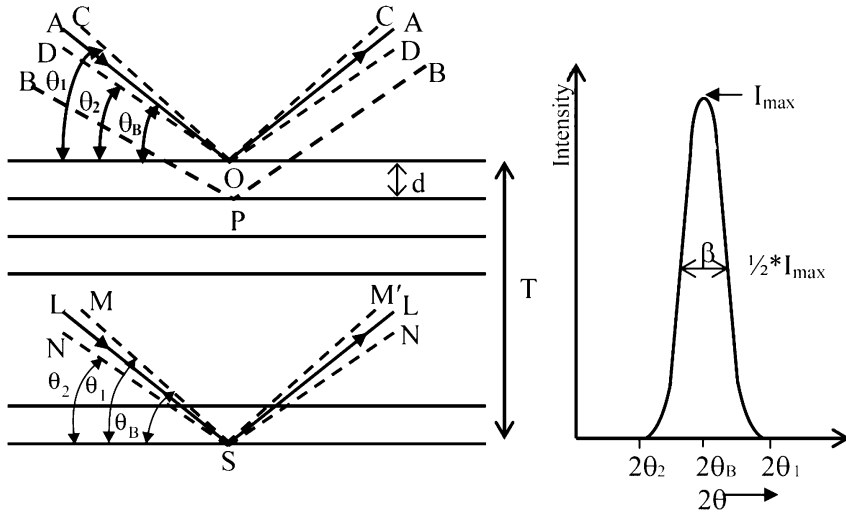


Fig. 7.25 Effect of crystal size on the diffraction

and for ray such as DO let the constructive interference occur just one plane before the m th plane. Therefore

$$(m - 1) \lambda = 2T \sin \theta_2 \quad (7.23)$$

Subtracting (7.23) from (7.22), we get

$$\lambda = T (\sin \theta_1 - \sin \theta_2) \quad (7.24)$$

Therefore

$$\lambda = 2T \cos \left(\frac{\theta_1 + \theta_2}{2} \right) \sin \left(\frac{\theta_1 - \theta_2}{2} \right) \quad (7.25)$$

$$\text{as } \left(\frac{\theta_1 + \theta_2}{2} \right) = \theta_B \text{ and } \sin \left(\frac{\theta_1 - \theta_2}{2} \right) = \frac{\theta_1 - \theta_2}{2} = \frac{\beta}{2} \quad (7.26)$$

Thus we get

$$\lambda = 2.T.\frac{\beta}{2}. \cos \theta_B \quad (7.27)$$

Therefore,

$$T = \frac{\lambda}{\beta \cos \theta_B} \quad (7.28)$$

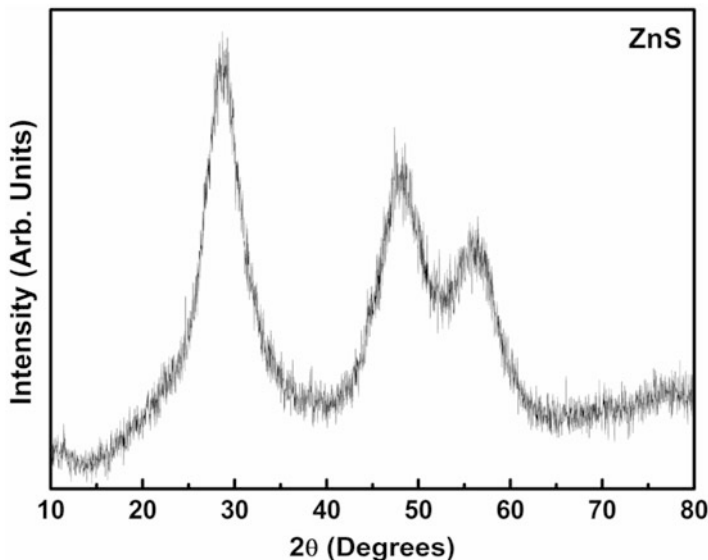


Fig. 7.26 X-ray diffraction pattern of ZnS nanoparticles

More accurate form of the equation for spherical nanoparticles is

$$T = \frac{0.9 \lambda}{\beta \cos \theta_B} \quad (7.29)$$

This is known as Scherrer equation. The width β is the broadening caused by nanoparticle size. Scherrer formula can be used to evaluate the average particle size smaller than ~ 100 nm which is just what is needed in nanoscience. However for extremely small particles (< 2 nm), the broadening of diffraction peaks may become very large so that they resemble the peaks in an amorphous solid. Under such conditions, the peaks may result from convolution of peaks due to different planes also and unambiguous size determination is prohibited. For example, consider Fig. 7.26 in which the diffraction peak at 30° due to ZnS nanoparticles is so much broadened that it is not possible to decide whether it is the convolution of three planes viz. of a hexagonal lattice structure or it is a size broadened peak of a zinc blende structure. In such situations, another method known as Debye Function Analysis is adopted. In order to determine the shapes of nanoparticles, study of self-assembly in particles, multilayers etc. additional techniques like Small Angle X-ray or Neutron Diffraction (SAXS or SANS), X-ray Reflectivity or Neutron Reflectivity are available. Such techniques require intense X-ray sources such as rotating anode or synchrotron radiation and will not be discussed here. For neutron scattering, neutrons available from a nuclear reactor or special spallation sources are used.

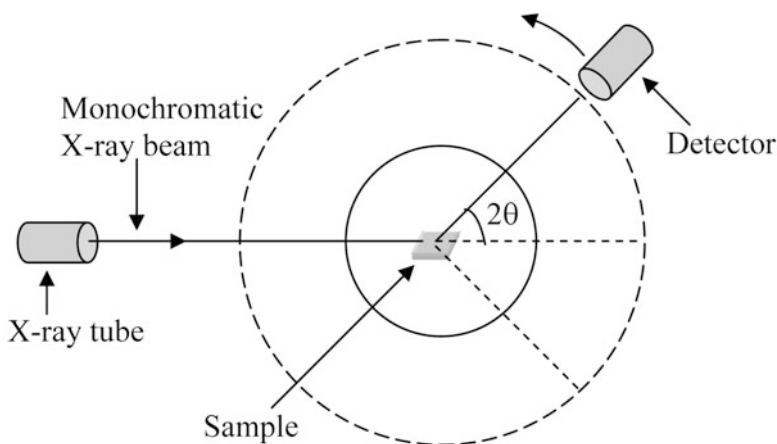


Fig. 7.27 Schematic of X-ray diffractometer

7.5.7 X-ray Diffractometer

There are different types of X-ray diffractometers available for crystal structure analysis. Same can be used for nanomaterials analysis. The most commonly used diffractometer is known as Powder Diffractometer or Debye-Scherrer Diffractometer after its inventors. This diffractometer is conceptually simple and allows quite an accurate determination of crystal structure of polycrystalline samples, thin films and nanoparticles.

As illustrated in Fig. 7.27 it consists of a monochromatic source of X-rays, usually from a copper target or anode giving CuK_α ($\lambda = 0.154 \text{ nm}$ after passing through nickel filter), sample holder and an X-ray detector. Both sample and the detector move around an axis passing through sample centre and normal plane of the paper.

Some times sample heating/cooling facilities are provided.

The diffracted rays make an angle 2θ at the detector with respect to incident beam direction. A plot of intensity (counts) as a function of angle 2θ (usually 20° – 160°) is a diffraction pattern ready for analysis. For most of the routine work, this is quite sufficient. Detector is a suitable photon counter like Geiger Muller tube, proportional counter, scintillation counter or solid state detector. Usually, due to finite size of X-ray beam $\sim 1\text{--}2 \text{ mm}^2$, smaller angles ($<20^\circ$) are not accessible using these diffractometers. However for some detailed analysis of thin films or nanoparticles where additional information can be obtained at as small as $\theta \sim 0.1^\circ$ – 0.2° , modifications of Debye Scherrer Diffractometer or another diffractometer is needed.

7.5.8 Dynamic Light Scattering

It is a technique capable of determining the ‘hydrodynamic’ size of the particles. Hydrodynamic size can be defined for a particle of irregular shape as the effective size of a particle when it is dispersed in a liquid. For spherical particles, the hydrodynamic size is same as the actual particle size with radius say R .

Dynamic Light Scattering (DLS) technique is also known by various names like Photon Correlation Spectroscopy (PCS), Quasi Elastic Scattering, Diffusive Light Scattering, 3-D Dynamic Light Scattering, Beating Spectroscopy, Homodyne Spectroscopy and Intensity Fluctuation Spectroscopy. However, DLS and PCS are the names which are more common in use.

The DLS technique, unlike microscopy techniques like SEM, TEM, STM and AFM as discussed previously in this chapter, is capable of determining the particle size or size distribution only when the particles are dispersed in some liquid. The technique depends on the intensity fluctuations of visible light scattered from the particles while they make random Brownian motion in the liquid. It is basically Rayleigh scattering and size of the particles has to be much smaller than the wavelength of light used as a source.

When a beam of intense and monochromatic coherent beam like laser light falls on a small volume of liquid, the scattered light intensity measured by a detector at certain angle θ with respect to the direction of the incident beam depends on the angle θ , wavelength of light λ and refractive index n of the medium. The scattering vector value is given as

$$q = \frac{4\pi n}{\lambda} \sin\left(\frac{\theta}{2}\right) \quad (7.30)$$

The scattering geometry is depicted in Fig. 7.28. If the scattering from each particle occurs only once, the analysis becomes straightforward. Therefore multiple scattering events are avoided using dilute samples or small concentration of particles

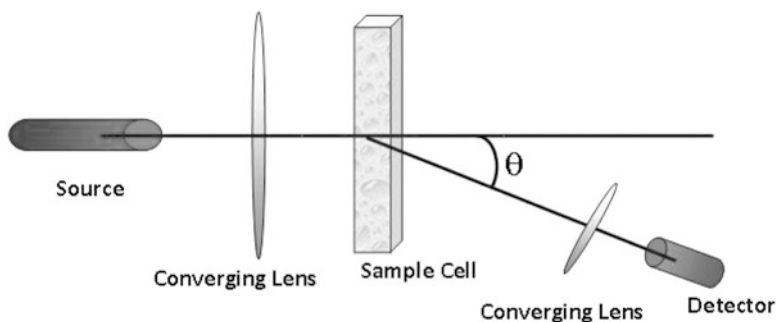


Fig. 7.28 Light scattering geometry

in the liquid. It is known from the Stokes-Einstein relation that the diffusion coefficient of particles with radius R is given by

$$D = \frac{kT}{6\pi\eta R} \quad (7.31)$$

where k is Boltzmann constant, T – temperature of the liquid and η is viscosity of the liquid.

As the particles move randomly in the liquid, their scattered intensities at the detector also change randomly or fluctuate. This is because the observed intensity of light from a given volume of liquid depends upon the interference of light being scattered from randomly distributed particle of the irradiated volume. The intensity at any instance, therefore, depends upon the relative positions of the particles or phase differences. If the intensity fluctuations due to random motion of particles are analyzed on a time scale smaller than the time taken by the particles to move the distance of the incident wavelength λ , then one can find that movements of the particles are still correlated. Obviously if the intensities are measured after a very long time the positions would not be correlated. When the measurements are made on a very short time scale, the particle position would not change drastically. As more and more time elapses the correlation would become less and less. Correlation decays exponentially. The exponential decay is related to the diffusion of particles and through Eq. (7.31) to the particle size. If the particles are monodispersed, a single exponential function is expected. Statistical analysis of the intensity fluctuation can be made using the following equation.

$$g^1(t) = \exp(-q^2 D t) \quad (7.32)$$

where $g^1(t)$ is the first order normalized autocorrelation function, q is the scattering vector length defined by Eq. (7.30) and D is the diffusion coefficient in the equation. By measuring $g^1(t)$ at one scattering angle, diffusion coefficient (hence particle radius R) of particles can be obtained. If the measurements can be done over a range of scattering angles, plot of corresponding $g^1(t)$ against q^2 gives a straight line, slope of which yields D , the diffusion coefficient.

If the particles are not monodispersed, the first order correlation function $g^1(t)$ would depend upon more than one diffusion coefficients. If D_i is the diffusion coefficient of the i th particle, $g^1(t)$ can be written as

$$g^1(t) = \sum_{i=1}^n \exp(-q^2 D_i) A_i \quad (7.33)$$

where A_i denotes the weight or amount of each component. A common method used to analyze the polydispersed particles is ‘cumulants method’. Here $g^1(t)$ is expanded and written as

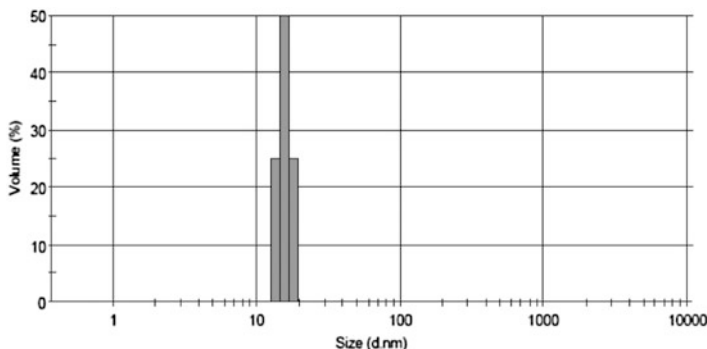


Fig. 7.29 Size determination of PbS nanoparticles using a Dynamic Light Scattering equipment

$$g^1(t) = \exp(-q^2 D_t) (1 + K_1 + K_2 + \dots) \quad (7.34)$$

$K_1, K_2 \dots$ are the first, second etc. cumulants.

The experimental set up is very simple as schematically shown in Fig. 7.28 and the commercial instruments are equipped with determining the autocorrelation function including cumulant method. They are able to give information of diffusion coefficients as well as particle size distribution over a wide range of ~ 1 nm to $10 \mu\text{m}$ size. Usually a beam of laser light is necessary, as concentration of particles in the solution has to be kept low in order to avoid the multiple scattering effects. One also tries to use efficient detectors like avalanche photodiode so that good signal is obtained. Figure 7.29 depicts a typical size distribution obtained for PbS nanoparticles obtained from such a DLS set up.

7.6 Spectroscopies

7.6.1 Optical (Ultraviolet-Visible-Near Infra Red) Absorption Spectrometer

Optical absorption spectroscopy is a very useful technique to study metals, semiconductors and insulators in bulk, colloidal, thin film and nanostructure forms. Semiconducting as well as some insulating materials have an optical energy gap. When the energy of photons is insufficient to excite electrons (see Fig. 7.30) from valence band to conduction band, no absorption takes place. At some critical photon energy, a sudden rise in absorption occurs as energy of photons is just sufficient to excite the electron to conduction band minimum. At still shorter wavelengths or higher energy photons continue to get absorbed. The absorbed (or reflected) intensity as a function of wavelength from ultraviolet (~ 200 nm) to near infra red ($\sim 3,000$ nm or many a times only up to $1,000$ nm) is useful to understand electronic structure and transitions between valence and conduction band of materials.

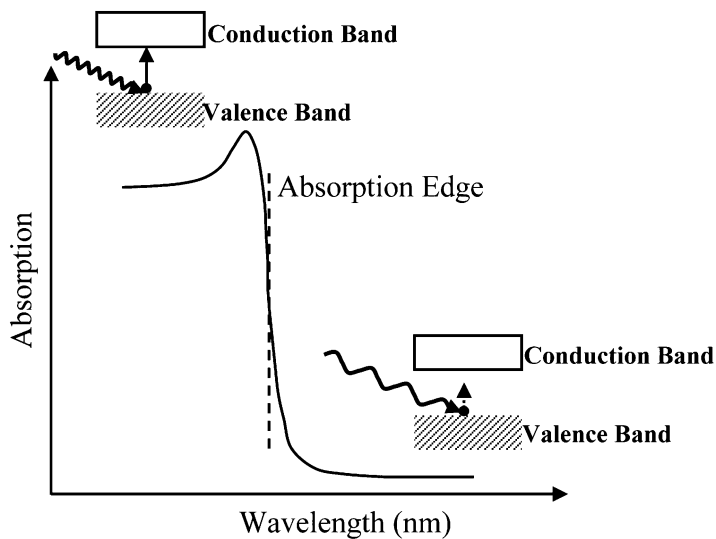


Fig. 7.30 Schematic optical absorption spectrum (semiconductor)

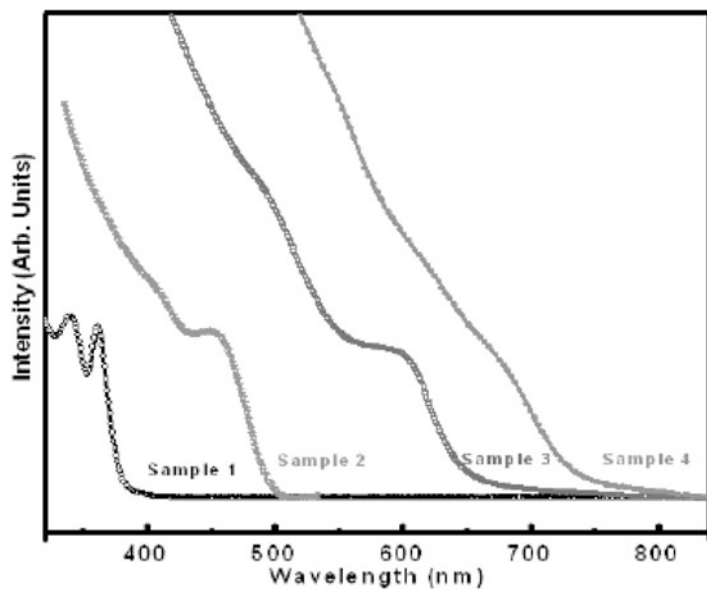
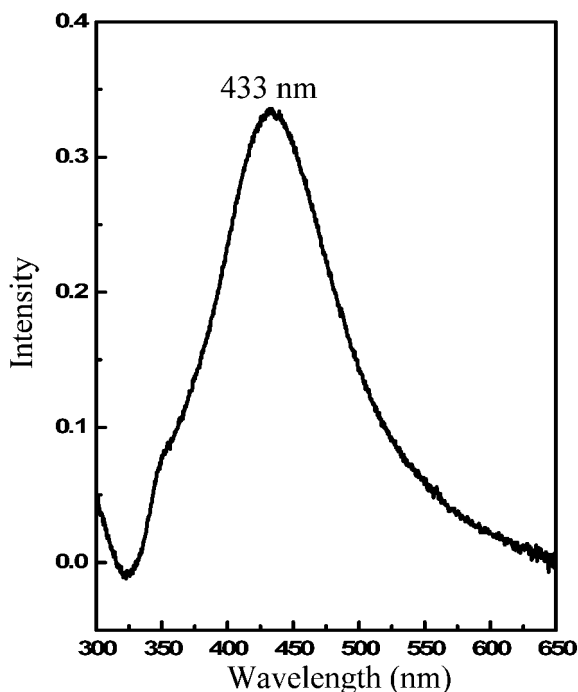


Fig. 7.31 Absorption spectra of CdSe nanoparticles

For nanomaterials with reduction in the particle size characteristic of materials a shift in the absorption edge can occur. The shift is usually to shorter wavelength and therefore known as *blue shift*. The theory of size dependent energy gap changes is discussed in Chap. 8. Figure 7.31 shows size dependence of absorption spectra

Fig. 7.32 Optical absorption due to silver nanoparticles



for CdSe nanoparticles. Note that absorption of bulk CdSe should have started at ~ 730 nm as bulk CdSe has an energy gap of 1.7 eV at room temperature.

In case of metal thin films or particles like gold or silver, one may have strong surface plasmon resonance (more details in Chap. 8) due to resonant absorption of photons. Peaks are also expected in the range of UV-Vis-NIR range. The peaks are broad and their positions are size dependent. They too show blue shift with reduction in the particle size. Peak widths for both metal or semiconductor nanoparticles depend upon the size as well as size distribution of particles. See Fig. 7.32 for silver nanoparticles.

UV-Vis-NIR absorption spectroscopy is, therefore, a useful technique to analyze nanomaterials. Here we shall briefly discuss the essential parts of an experimental set of a UV-Vis-NIR spectrometer.

7.6.2 UV-Vis-NIR Spectrometer

The experimental set up in principle is very simple. A high intensity lamp (or change of lamps in different spectral regions) giving radiation from UV to NIR region is required. A monochromator selects different wavelengths, which fall on the sample. Depending upon its properties, the sample reflects or absorbs certain wavelengths and transmits the rest. The transmitted (or reflected) intensity at different

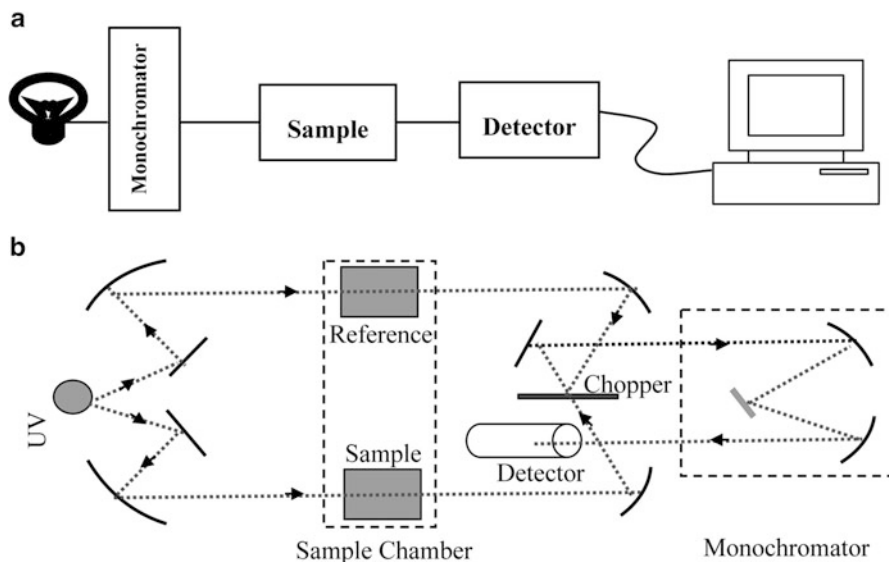


Fig. 7.33 (a) Schematic of UV-Vis spectrometer. (b) Schematic of dispersive spectrometer

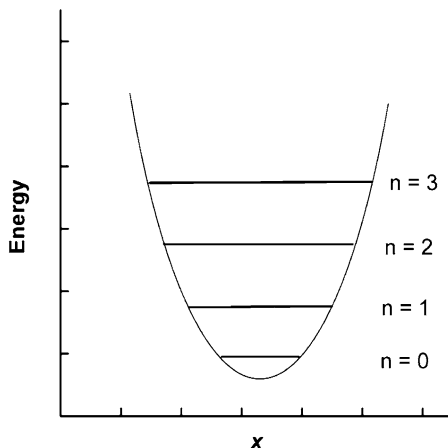
wavelengths are detected by a photodetector and given as an input to a recorder or computer. When the monochromator signals corresponding to selected wavelengths falling on the samples are the signals for X-axis and photodetector signal for Y-axis, absorption plots as shown earlier in Figs. 7.31 and 7.32 are obtained. Figure 7.33a illustrates the schematic of the spectrometer. In practice it is advisable to use a reference, so as to eliminate the effect of path traversed by photons in sample holder, solutions in which colloids are suspended etc. Therefore beam from the radiation source is split into two parts. One part traverses through the reference compartment (cuvett) and through the nanoparticles dispersed in a liquid forming the sample under analysis. The reference is a liquid which is used to disperse the nanoparticles, without actually having nanoparticles. The intensities of light from reference and sample are finally compared to give the signal only due to the sample. A schematic of a typical double beam spectrometer is given in Fig. 7.33b.

Usually H_2 discharge lamp, deuterium lamp or tungsten lamp are used. Common monochromators are gratings or prisms. Photon detectors that are normally used are photomultiplier tube or photodiode. Liquid sample holders are usually made up of quartz glass.

7.6.3 Infra Red Spectrometers

We have seen earlier that inorganic nanomaterials are often surface passivated with organic molecules or sometimes they form composites with organic materials.

Fig. 7.34 Vibration levels in a molecule



Identification of such molecules often throws light on the processes occurring at the surface of nanomaterials. Many functional groups like $-\text{OH}$, $-\text{SH}$, $=\text{C}=\text{O}$, $-\text{CH}_2$, $-\text{NH}_2$ etc. have some characteristic absorption bands in the Infra Red regions.

Molecules have characteristic vibration energy levels as schematically shown in Fig. 7.34. Characteristic absorption bands for the molecules occur as molecules undergo transitions from their one characteristic energy level to another.

The allowed energies of vibration levels are given as

$$E_n = (n + 1/2) h\nu_0 \quad (7.35)$$

where n are the quantum numbers for vibrations and take the values 0, 1, 2, 3... and ν_0 is the frequency of oscillator. The transition occurs by absorption of photon following the selection rule $\Delta n = 0$. For the infra-red absorption to take place, the molecules have to have a dipole moment. The molecules are then said to be IR active. Thus an H_2O molecule is IR active and would show characteristic absorption band. However symmetric CO_2 molecule does not have any permanent dipole on it. Hence CO_2 is not IR active.

It is useful to perform infra red spectroscopy of nanomaterials to find out if there are any active molecules present. Infra Red band is roughly divided into three regions viz. Near Infra Red ($912,500 \text{ cm}^{-1}$ to $4,000 \text{ cm}^{-1}$), Infra Red ($4,000 \text{ cm}^{-1}$ to 400 cm^{-1}) and Far Infra Red (400 cm^{-1} to 10 cm^{-1}).

There are two types of infrared spectrometers in common use. One is simple dispersive spectrometer and is often just referred to as IR spectrometer and the other is Fourier Transform Infra Red Spectrometer or FTIR spectrometer. FTIR spectrometer is mainly used to simultaneously record the entire wavelength range thus reducing data acquisition time as well as improving the signal to noise ratio. FTIR does not have any monochromator to select wavelengths. The dispersive IR

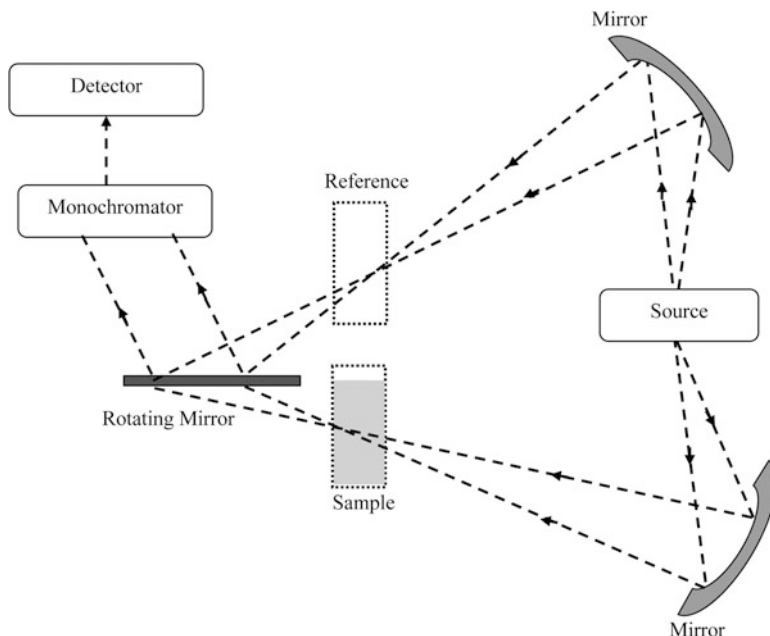


Fig. 7.35 Schematic of Dispersive IR Spectrometer

spectrometer uses either grating or prism as a monochromator to select the wavelength and sequentially scan the spectrum. We shall discuss both the spectrometers in the following sections.

7.6.4 Dispersive Infra Red Spectrometer

Figure 7.35 illustrates schematically a dispersive type IR spectrometer. It consists of an IR source, sample compartment, chopper, monochromator, focusing and collimating mirror and an IR detector. There are slits located at various points in the path of the beams so as to control the size of the beam.

An IR source is a tungsten lamp, Nernst glower or nichrome coil and high pressure mercury arc lamp depending upon whether NIR, Mid IR or Far IR region is to be investigated. Spectrometers are usually able to switch between different lamps.

IR detectors may be photoconductive cells, thermopiles, thermistors, Golay or pyroelectric, depending upon the range of wavelengths to be detected. The beam of IR from a single source is usually split into two and allowed to pass through the slits, then on reference sample and the actual sample. Reference sample signal can be subtracted from sample spectrum signal so as to eliminate effects due to sample

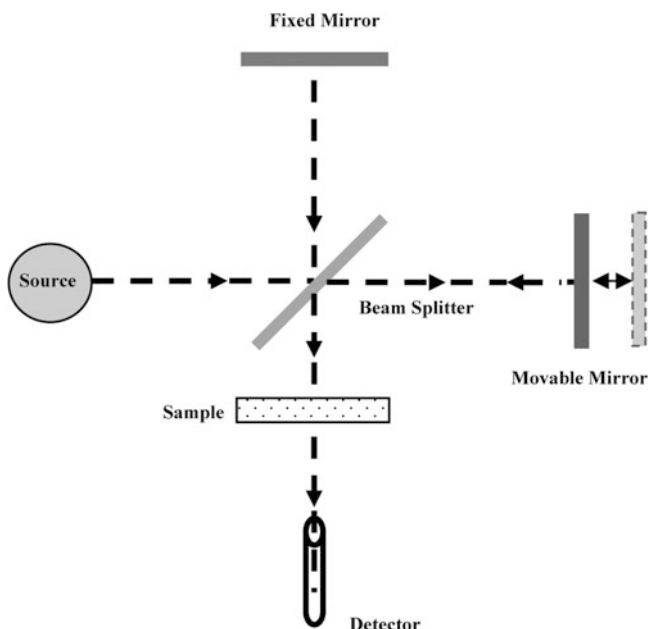


Fig. 7.36 Schematic of Michelson interferometer

preparation. Chopper cuts beam either from sample or from reference alternately so that each beam falls on gating (or prism) and then to the detector. Grating selects different wavelengths sequentially. Thus at the detector, signal from the reference and sample arrives alternately for different wavelengths scanned by the grating. IR detector output from sample and the reference are compared and plotted as absorbed intensity versus wavelength or usually cm^{-1} .

7.6.5 *Fourier Transform Infra Red Spectrometer*

Fourier Transform Infra Red (FTIR) spectrometer makes use of the Michelson interferometer for recording the spectra. As shown in Fig. 7.36 a parallel beam of infra red rays falls on the beam splitter BS.

IR source and detector are similar to those used in dispersive type IR spectrometer discussed earlier. Part of the beam falls on a movable (0–1 cm) mirror M_1 and a fixed mirror M_2 .

The rays are reflected back from both the mirrors along the same path and interfere at BS. A part of this combined beam falls on the sample and the detector. Constructive and destructive interference occurring at BS depends upon the path length of the rays. A white beam i.e. the beam containing a broad continuous spectrum of wavelengths produces constructive and destructive interference pattern

of every wavelength with all others. Intensity as a function of the position of movable mirror position is given by

$$I(x) = \int_{-\infty}^{\infty} I(v) \cos(2\pi xv) dv \quad (7.36)$$

Inverse transform gives

$$I(v) = \int_{-\infty}^{\infty} I(x) \cos(2\pi xv) dx \quad (7.37)$$

The recombined beam passing through the sample produces absorption spectrum in which certain characteristic frequencies are absorbed by molecules present in the sample.

With modern computers, it is quite an easy job to carry out a Fourier Transform. Detector collects, for example, signal from sample every millisecond and stores each spectrum in different locations. Spectra are then Fourier transformed and resultant spectra are obtained as an output. Thus, spectra can be generated very fast and Fourier Transform also is very fast. Thus better and fast data acquisition is possible using an FTIR spectrometer. Hence most of the modern commercial infrared spectrometers are FTIR spectrometers. An FTIR spectrum of TiO_2 is shown in Fig. 7.37.

7.6.5.1 Sample Preparation

Sample preparation is a difficult task in IR range as there is no transparent material for cuvettes. Alkali halides (such as NaCl, KBr) are usually used which are transparent even at longer wavelengths. Powder samples are mixed with alkali halides

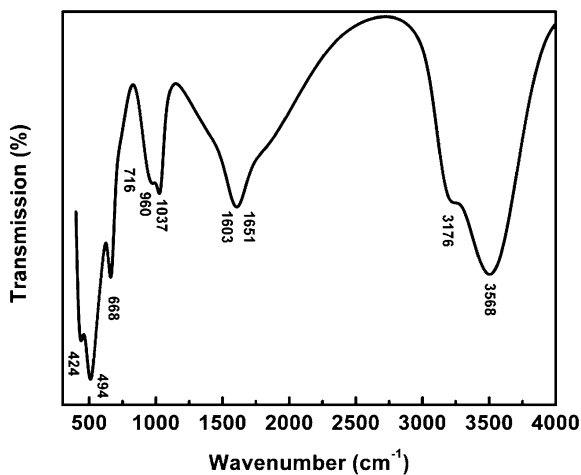


Fig. 7.37 FTIR of TiO_2 , suggesting vibrations due to TiO_6 octahedron and adsorbed hydroxyl group ($3,568 \text{ cm}^{-1}$)

and pressed in the sample holder. For liquid samples, there are single crystals of KBr or NaCl and liquid is sandwiched between the two. But in this case one cannot use aqueous solutions because alkali halides are soluble in water. For such samples silver chloride is used. One can use Teflon also but it shows absorption bands for C-C and C-F. For frequencies less than 600 cm^{-1} one can use polyethylene cell also.

7.6.6 Raman Spectroscopy

Raman spectroscopy is another powerful technique for the analysis of molecules or particles. This also is sensitive to the vibration spectrum of molecules. However it is complementary to the IR analysis. Unlike IR active molecules, Raman active molecules do not depend upon the presence of a dipole moment but on the polarizability of the molecule. The technique is based on the Raman effect discovered by Sir C.V. Raman in 1928. Whenever scattering of the light occurs, the scattered light consists of two types viz. *Rayleigh scattering* and *Raman Scattering*. Rayleigh scattering is strong and has the same frequency (elastic scattering) as the incident beam (ν_0), and the other is called *Raman scattering*. Raman scattering is inelastic scattering and has frequencies $\nu_0 \pm \nu_m$ where ν_m is a vibrational frequency of a molecule. Raman scattering is very weak ($\sim 10^{-5}$ of the incident beam). The decreased frequency ($\nu_0 - \nu_m$) and increased frequency ($\nu_0 + \nu_m$) lines are called the *Stokes* and *anti-Stokes* lines, respectively. The scattering is described as an excitation of the molecule to a virtual state which is lower in energy than a real electronic transition, with nearly coincident de-excitation and a change in vibrational energy. The scattering event occurs in 10^{-14} s or less. In Raman spectroscopy, the vibrational frequency ν_m is measured as a shift from the incident beam frequency ν_0 (Boxes 7.9 and 7.10).

Box 7.9: Fourier Transform Spectroscopies

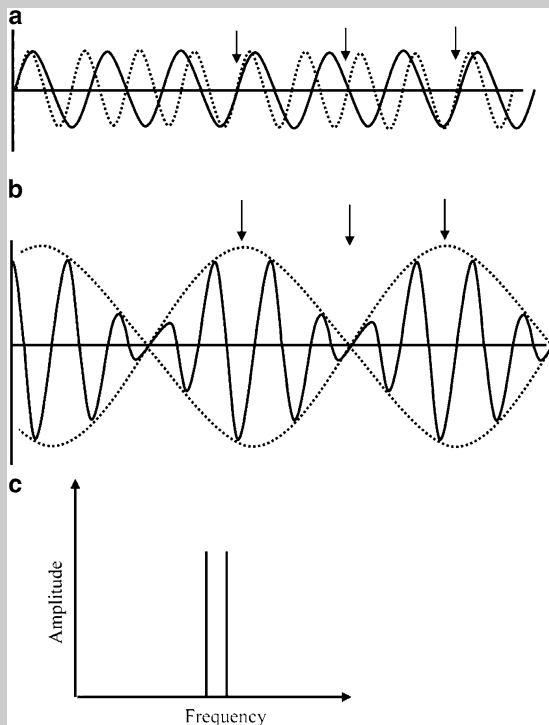
Technique of Fourier Transform is used in many spectroscopies like Raman Spectroscopy, Infra Red Spectroscopy, and Nuclear Magnetic Resonance Spectroscopy. It collects simultaneously the data at various frequencies and is much superior to conventional way of sequentially scanning different frequencies (or wavelengths). Fourier Transform is a mathematical method and is capable of resolving complex spectrum into various frequencies as follows.

Consider that two sine waves of different wavelengths have mixed as shown in Fig. 7.38a. They have different frequencies and their sum would appear as in Fig. 7.38b. A Fourier Transform of the resultant waves gives two separate frequencies as showed in Fig. 7.38c.

(continued)

Box 7.9 (continued)

Fig. 7.38 Fourier transform of two waves having different frequencies. Fourier Transform bears its name due to Jean Baptiste Fourier, a French mathematician who developed this procedure in early 1800

**Box 7.10: C.V. Raman**

C.V. Raman was born on 7th November 1888 in a city Tiruchirapalli in Tamil Nadu, India. His full name is Chandrasekhara Venkata Raman. His father was a Professor in Mathematics and Physics. Born in a highly educated family Raman was known as a very bright student. He joined Presidency College in Madras (now known as Chennai) in 1902 and received B.A. degree in 1904 bagging Gold Medal in Physics. In 1907 he received M.A. degree with distinction. Immediately he joined Indian Finance Department in 1907 as the conditions prevailing in India were such that even a bright student like Raman had no proper opening in science. In spite of this, young Raman did not lose his passion for science and continued his experiments after his office hours. He used to regularly visit Indian Institute of Cultivation of Science and pursue the experiments in light scattering. Later in 1917, he was

(continued)

Box 7.10 (continued)

offered a prestigious position of Sir Tarakanath Palit Professorship in Calcutta University. He remained with Calcutta University for 15 years during which period he became known all over the world as a great scientist.



His recognition started with prestigious Fellowship of Royal Society of London in 1924. He was made 'Knight of British Empire' in 1929. In 1930 he received the Nobel Prize for his discovery of the light scattering effect known as 'Raman Effect' after him. He is the only Indian scientist to receive this honour. He became the director of Indian Institute of Science, Bangalore in 1934. In 1949 he established in Bangalore 'Raman Research Institute', one of the well known institutes of India. Raman received 'Bharat Ratna' award from Government of India and Lenin Peace Prize in 1957 to name a few amongst the various honours. Raman was not only interested in the light scattering but also in acoustics particularly science of musical instruments. Raman was a great lover of diamonds and other stones and had a huge collection of them which he turned into a museum in Raman Research Institute. Raman died on 21st November 1970.

A Raman spectrometer comprises four components which are: (1) excitation source (laser), (2) sample illumination and collection system, (3) wavelength selector and (4) detector and computer processing system. FT-Raman spectrometer is preferred over the normal Raman spectrometer due to the advantage of measuring information of all frequencies at the same time. The instrumentation of FT-Raman is similar to normal Raman spectrometer with an additional inclusion of a Michelson interferometer, which enables the simultaneous acquisition of signals of all frequencies along with the improved resolution. Figure 7.39 shows a schematic of a typical FT-Raman spectrometer. The laser is incident on the sample by means of a lens and a parabolic mirror. The scattered light from the sample is collected and passed to a beam splitter and to the moving and fixed mirrors in the interferometer head. It is then passed through a series of filters and focused onto a liquid-nitrogen-cooled detector.

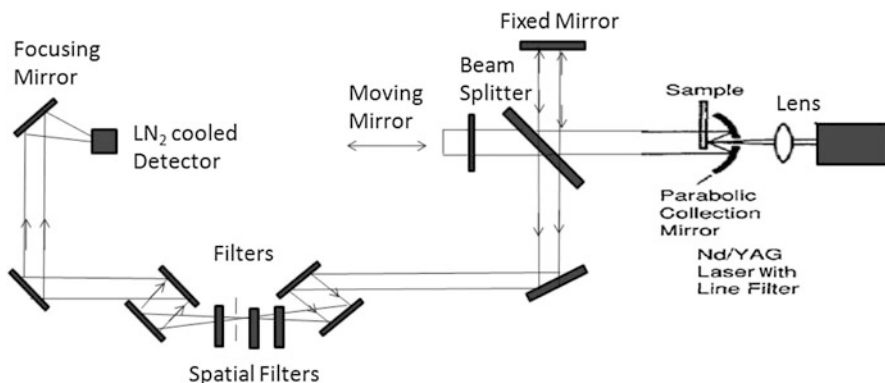
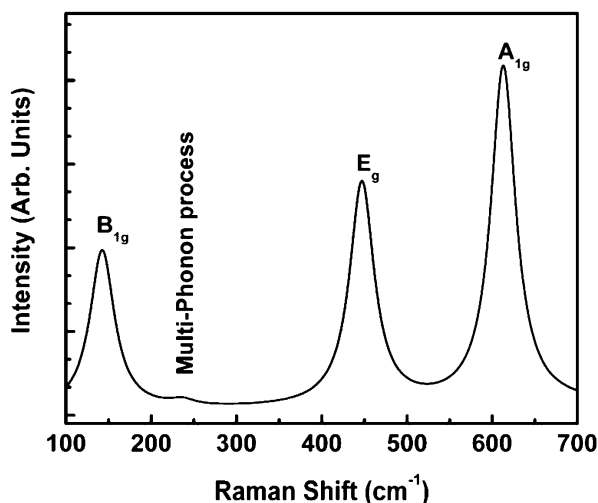


Fig. 7.39 Schematic diagram of FT-Raman spectrometer

Fig. 7.40 Raman spectrum of TiO_2 . Different modes of vibrations can be seen which are characteristic of rutile phase of TiO_2



Raman spectra are shown as 'Raman shift'. Raman spectra are considered to be indispensable for carbon nanotubes and other carbonaceous materials as amorphous, crystalline etc. characteristic forms can be easily identified. Figure 7.40 depicts the Raman spectrum for rutile TiO_2 .

7.6.7 Luminescence

Some materials when excited with an external source of stimulus like electrons or light emit light in the visible range, UV or IR. This phenomenon is known as *luminescence*. The word luminescence was coined by E. Wiedermann in 1888 from a Latin word *lumen* which means *light*. The word luminescence includes

fluorescence and *phosphorescence*. They differ in the duration of time over which light is emitted. Usually, the terminology fluorescence is used if the emission of light takes place within $\sim 10^{-4}$ s of stimulus. If the emission persists for a longer duration of few tens of milliseconds to ~ 10 s after the stimulus is removed it is termed as phosphorescence. Phosphorescence is sometimes also called as *afterglow*.

Many nanomaterials exhibit enhanced (increased intensity) luminescence as compared to their bulk counterparts. Some materials like silicon which are not luminescent in their bulk form become luminescent in nano form, like porous silicon (discussed in Chap. 11). Therefore luminescence investigations of nanomaterials are often carried out.

There are various types of luminescence like photoluminescence, cathodoluminescence, electroluminescence and so on depending upon the external stimulus used to excite the material. Luminescence investigations are useful to analyse the electronic structure of the material. It not only gives information about the transitions between conduction and valence band but also the localized states due to impurities or doping. We shall discuss their details in Chap. 8. Here we shall consider only the experimental technique to obtain photoluminescence. Further it should be mentioned that we discuss here only necessary parts of the spectrometer and actual instrument may be quite elaborate.

7.6.7.1 Photoluminescence Spectrometer

A schematic diagram of photoluminescence set up is shown in Fig. 7.41 which can be suitably modified to observe other types of luminescence by using the relevant source of excitation in place of photon source. Basically there has to be

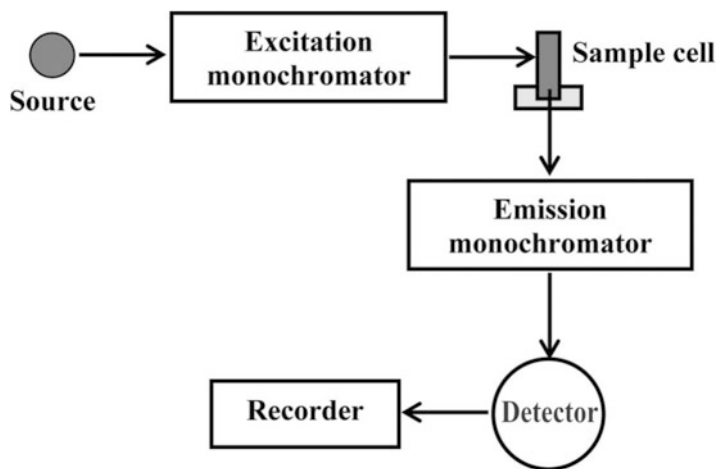


Fig. 7.41 Schematic layout of photoluminescence set up

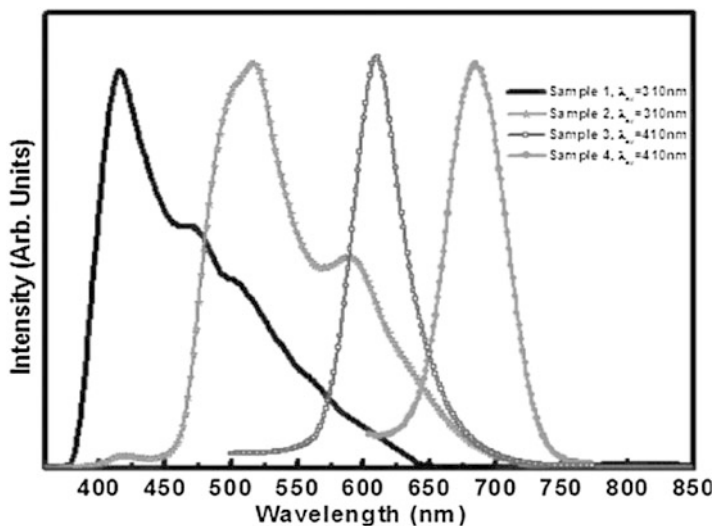


Fig. 7.42 Photoluminescence spectra from CdSe nanoparticles of two different sizes

a source of photons ranging from UV (~ 200 nm) to IR (~ 800 nm), a filter to throw away large band of wavelengths, wavelength selectors or monochromators, sample holder, a detector and a recording system like an X-Y recorder or a computer. Two types of arrangements viz. transmission type and reflection type layouts are normally used. The transmission type of geometry is, however, unsuitable for solid samples.

Figure 7.42 illustrates an example of a photoluminescence spectrum obtained from CdSe nanoparticles of different qualities and sizes.

7.6.8 X-Ray and Ultra Violet Photoelectron Spectroscopies (XPS or ESCA and UPS)

X-ray Photoelectron Spectroscopy (XPS), popularly known as Electron Spectroscopy for Chemical Analysis (ESCA), was developed by K. Siegbahn. A large number of books and reviews are available on ESCA. The technique is based on the photoelectric effect, as explained by Einstein in 1905. Accordingly (Fig. 7.43), photon of fixed energy $h\nu$ incident on an atom ejects an electron of binding energy E_B with kinetic energy E_K according to the equation

$$h\nu = E_K + E_B \quad (7.38)$$

Knowing $h\nu$ and by measuring with an energy analyzer, the kinetic energy E_K of electron, binding energy of electron in an atom can be determined. The equation is

Fig. 7.43 Schematic of photoemission

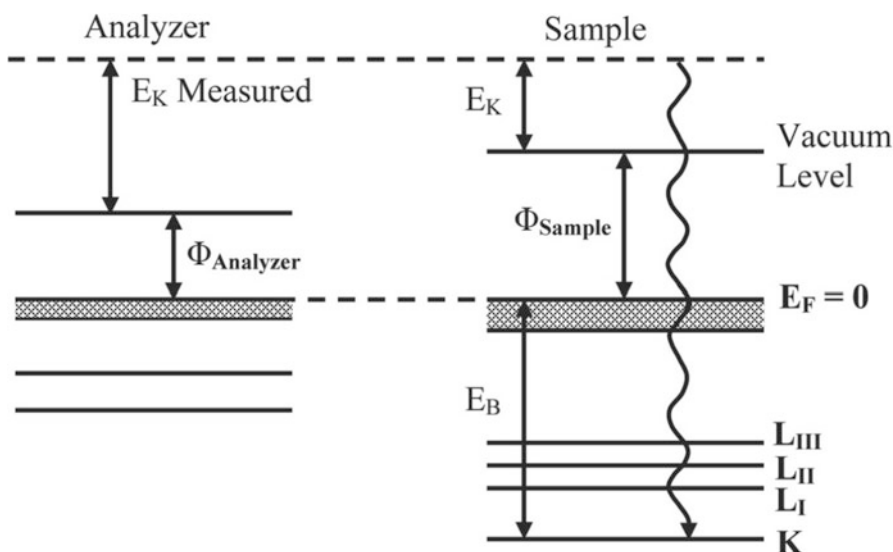
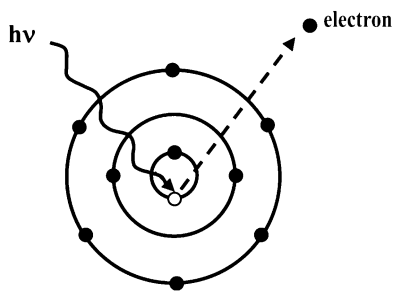


Fig. 7.44 Energy level diagram for photoelectron spectroscopy of solids

valid for atoms in gases, liquids or solids. In a solid (see Fig. 7.44), additional energy (ϕ) the work function of the solid is required for the electron to get emitted as

$$h\nu = E_K + E_B + \phi \quad (7.39)$$

Work function of solids changes from material to material. It also depends on their cleanliness and purity. Fortunately, it is not necessary to know ϕ of the solid. Kinetic energies of emitted electrons are measured with reference to Fermi level and ϕ can be replaced by ϕ_{SP} , the work function of the spectrometer. Therefore, one gets the measured E_K , which differs from the kinetic energy of photoelectron coming out from the sample. When Fermi level of the sample and the analyzer are aligned (by keeping both at earth potential) there is a constant difference between measured kinetic energy and the kinetic energies of electrons emitted from different

samples. As the work function of the spectrometer is known, it is not necessary to know the work function of the samples and one can easily find out the kinetic or binding energies of the samples.

One obtains intensity of electrons versus electron energy which is the Electron Distribution Curve (EDC). Core level electrons in atoms of the solid are very sensitive to their neighbouring atoms. Therefore analysis of photoelectrons was originally known as 'Electron Spectroscopy for Chemical Analysis'. However, there are many electron spectroscopies available now, which yield chemical information. It is therefore safer to call the technique as X-ray Photoelectron Spectroscopy (XPS), which specifically indicates that photoelectrons are produced using X-rays.

When an electron is ejected from a solid sample, a hole is created. Binding energy measured, therefore, gives the energy of the photoelectron in presence of the hole. When an electron leaves an atom, remaining electrons of the atom (and even the surrounding atoms) interact with the hole. The interaction energy depends on the atomic number as well as the energy level in which the hole resides. Corresponding relaxation energies can be large. Thus binding energies measured in an experiment are not the initial state energies of photoelectrons as depicted in Eq. (7.39). However, measured binding energies are still quite useful as they are characteristic binding energies of a given element. As mentioned earlier, the photoelectrons have energies which are quite sensitive to their local environment. Thus measurement of binding energies results into useful chemical information.

7.6.8.1 Ingredients of X-Ray Photoelectron Spectra

Photoelectron spectra are usually rich in their contents. Following ingredients can be made use of in understanding the properties of materials.

- (i) Chemical shift
- (ii) Valence band
- (iii) Auger peaks
- (iv) Spin-orbit splitting
- (v) Multiplet splitting
- (vi) Satellites
- (vii) Plasmon loss

Origin of these features is briefly described below. Details can be found in some review articles on XPS.

Chemical shift – Core electrons of atoms in solids are very sensitive to their surrounding. Whenever there is a charge transfer between outer electrons of different atoms, core electrons also respond to these changes by changing their energies. Thus, there are changes in the binding energies of electrons in a solid. These changes can be studied by analyzing the kinetic energies of photoelectrons.

Valence band – Photoelectron spectroscopy using X-rays or UV rays from helium can be used to infer about the density of states in the vicinity of Fermi level. Photoemission cross-sections are very sensitive to photon energy. Taking advantage

of the situation that in typical commercial instrument Al K_{α} (1,486.6 eV), Mg K_{α} (1,253.6 eV), He I (21.2 eV) and He II (40.8 eV) sources of widely different energies are present, useful cross-section dependent analysis can be made. Valence band spectra show considerable changes when recorded with different photon energies.

Auger peaks – Along with photoelectron peaks, some Auger peaks characteristics of elements present in the solid also can be obtained in a spectrum. Origin of Auger peaks will be discussed in the next sub-section.

Spin-orbit splitting – Some of the peaks in photoelectron spectra appear as doublets. These are due to splitting of p shells [$2p_{3/2} - 2p_{1/2}$, $3p_{3/2} - 3p_{1/2}$ etc.] as well as d and f sub-shells. Spin-orbit splitting for a given atom decreases with increase in the principal quantum number. It increases for a given principal quantum number with increase in atomic number. For large atomic number compounds, spin-orbit splitting can be used for finding out the oxidation state. Spin-orbit splitting differences are quite large for oxides of high Z atoms compared to pure forms.

Multiplet splitting – Large magnetic moments on some of the atoms/ions arise due to unpaired electrons in their $3d$, $4d$ or $4f$ shells. Photoelectron spectroscopy can be used to detect the presence of such unpaired electrons. For example, in Fe^{3+} there are two paired electrons in $3s$ level and five unpaired electrons with parallel spins in $3d$ level. When photoelectron is ejected from $3s$ level two situations can arise viz. (a) an electron left in $3s$ is parallel to electrons in $3d$ level or (b) antiparallel to $3d$ electrons. The energy difference due to situation in (a) and (b) can be as large as few electron volts and depends upon the number of electrons in d level. In fact, both the cases have certain probability to exist. Therefore $3s$ photoelectron spectrum exhibits two peaks. Splitting of $3s$ levels can be correlated to electrons in d level and becomes a measure of magnetic moment.

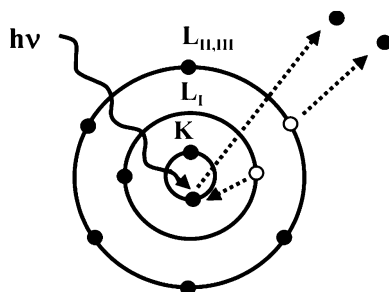
Plasmons loss – Photoelectrons with large kinetic energy can excite plasmons (plasmons are collective but quantized oscillations of electrons in a solid) with energy

$$\omega_p^2 = \frac{4\pi Ne^2}{m} \text{ (C.G.S.)} \quad (7.40)$$

where ω_p is plasmon frequency, n – number of electrons per cm^3 and m^* is the effective mass of electron. In such a case, photoelectron peaks are accompanied by peaks on their higher binding energy side. Bulk plasmon peaks along with second and higher harmonics can be observed with decreasing intensity. Surface plasmon peaks at energies $\omega_p/\sqrt{2}$ also can be observed.

Satellites – Intense or weak peaks can some times be observed along with main photoelectron peak. Origin of these peaks can be understood in terms of electron shake up, shake off or charge transfer and many body interactions collectively referred to as satellites. In the process of photoelectron emission, a hole is created. Remaining electrons respond to this by monopole excitation to a discrete level (shake up) or continuum (shake off). The energy for excitation is derived from outgoing photoelectron. Such photoelectron appears as a peak on the higher binding energy with respect to the ones that come without losing energy.

Fig. 7.45 Schematic of an Auger process



7.6.9 Auger Electron Spectroscopy

Auger electrons were discovered by P. Auger in 1925 but their potential for surface analysis was realized after few decades. As illustrated in Fig. 7.45, when one of the electrons from a core level is removed, an electron from outer level combines with the core hole. The energy difference between the two levels is either emitted as X-ray (photons) or utilized in emitting an electron from one of the outer levels. An electron removed by the later process is known as Auger electron. It is a three-level process and as can be seen from the illustration, three electrons are necessary for the Auger process to take place. Therefore, except H, He and Li, all the elements can produce Auger electrons. Energy of an Auger electron in Fig. 7.45 can be written as

$$E_{K,L_I,L_{II-III}} = E_K - E_{L_I} - E_{L_{II-III}} \quad (7.41)$$

where E_K , E_{L_I} and $E_{L_{II-III}}$ are binding energies of electrons in K, L_I and L_{II-III} levels respectively. Emissions of photons and Auger electrons are competing processes.

Production of core hole for emitting Auger electron is independent of the process in which core hole is created. Photons, electrons or even ions can be used to produce core hole. Auger electrons are therefore produced along with photoelectrons. Auger electrons are characteristics of atom from which they are released and are useful in the elemental analysis. One can use electrons, photons or even ions to produce Auger electrons. More details can be found elsewhere in some review article on Auger spectroscopy.

7.6.9.1 Surface Sensitivity of Photoelectron and Auger Electron Spectra

Photoelectron spectroscopy can be performed using X-rays or Ultra Violet (UV) rays. X-rays and UV rays can penetrate from few micrometres or fraction of a micrometre depth depending upon the energy of the radiation. However, electrons have much shorter mean free path in solids. Therefore, even if photoelectrons can be generated deep within a solid, few electrons can escape out of solid, depending upon their kinetic energy. An empirical curve of electron escape depth as a function

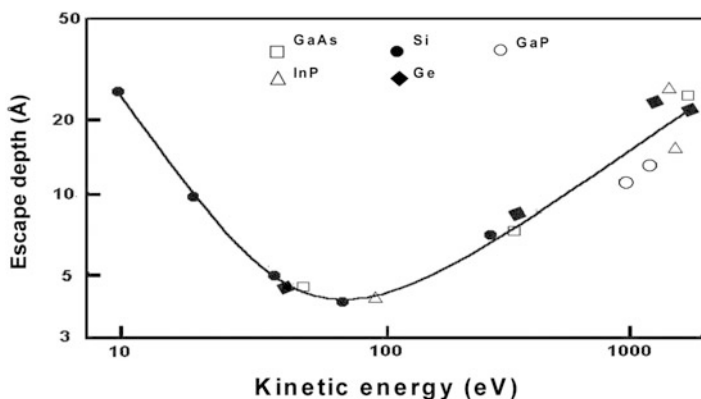


Fig. 7.46 Dependence of escape depth of electrons on kinetic energy

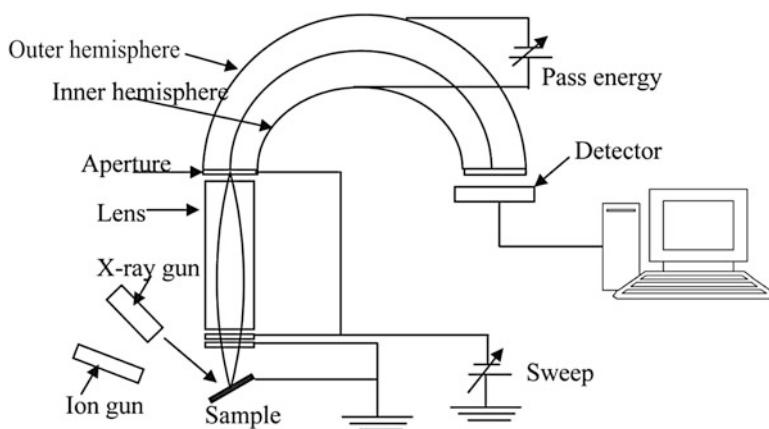


Fig. 7.47 Schematic of ESCA using CHA

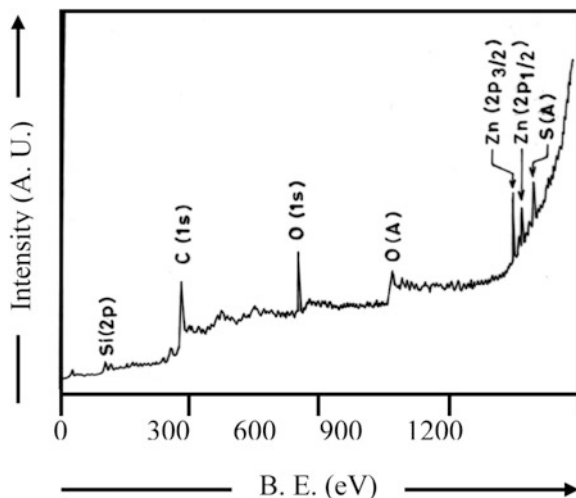
of their kinetic energy is shown in Fig. 7.46. It can be seen that if the electron kinetic energy is around 100–200 eV, escape depth of the electrons is very small, less than a nanometer.

One can always find for a given photon source, some photoelectrons from core or shallow energy level which will come from few nano or sub nanometer depth. Thus photoelectron spectroscopy is a surface sensitive technique. Same is true for Auger electrons.

7.6.9.2 Experimental Set Up

In Fig. 7.47, basic requirements of photoelectron and Auger spectroscopy are illustrated. Except for the source of excitation, instrumentation required by both the techniques is same. This is very useful. AlK_{α} with photon energy 1,486.6 eV and

Fig. 7.48 XPS spectrum of ZnS sample obtained using MgK_{α} source



MgK_{α} with photon energy 1,253.6 eV are available from a twin anode as source of X-rays. Alternatively, a helium discharge lamp also can be used to emit HeI and HeII with photon energies 21.2 and 40.8 eV respectively. HeI and HeII are obtained by operating the UV discharge lamp with helium gas flowing in it.

Auger electrons can be a part of photoelectron spectrum. However, it is common to use an electron gun (2–5 keV) as the source of incident electrons to create core holes. Electrons have the advantage that they can be generated easily and focused to a small spot. They also can be rastered on the sample surface.

Photoelectrons and Auger electrons are analyzed in the same analyzer using Concentric Hemispherical Analyzer (CHA) or double pass Cylindrical Mirror Analyzer (CMA). Analyzer is controlled using spectrometer control unit (SCU). Electrons passing through them are selected according to their energies, are detected and amplified using channeltron or channel plate. Amplified signal after suitable noise filtering is an input for an X-Y recorder or a computer. An X-Y plot or a spectrum of intensity versus electron kinetic/binding energy can be obtained. One can make use of readily available data books to identify the elements (Fig. 7.48).

Core levels due to elements present in the sample are marked on Fig. 7.49. Besides the prominent peaks, a number of other peaks appear on high binding energy sides. They also are useful in understanding electronic structure of the sample. A typical Auger spectrum can be seen in Fig. 7.49.

7.7 Magnetic Measurements

7.7.1 Vibrating Sample Magnetometer (VSM)

Using this technique, magnetization of a sample on application of magnetic field can be measured. Other sensitive technique is Superconducting Quantum Interference

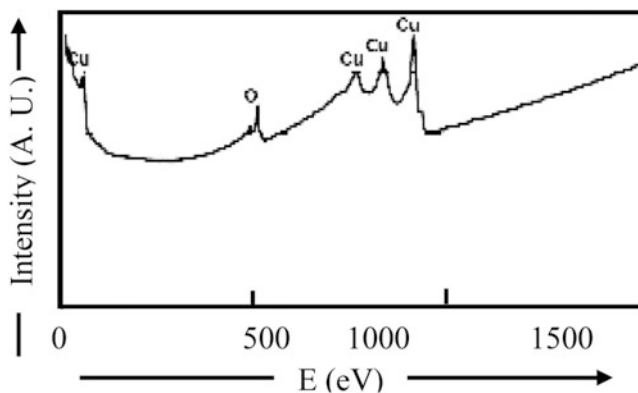
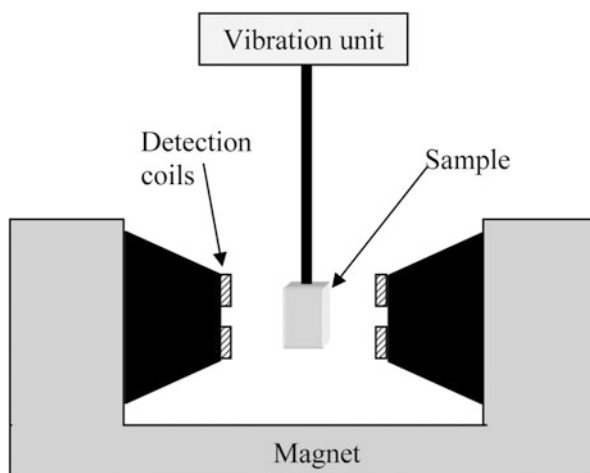


Fig. 7.49 Auger spectrum of Cu sample

Fig. 7.50 Schematic diagram of VSM



Device (SQUID), not discussed here. The instrument is based on the principle that an oscillatory magnetic field can be created by vibrating a magnetic sample. Magnetization in the sample is induced by applying a uniform magnetic field to the sample. The induced changes in the magnetic field are detected by a search coil. Applied magnetic field is usually quite large, but being constant is not detected by search coil. Generally the constant magnetic field is not really constant but slowly varied so that induced magnetization in the sample at different fields can be investigated. One can assume this slowly varying magnetic field as constant compared to vibrating field (Hz). Figure 7.50 gives the schematic diagram of a typical VSM. In most of the standard set ups a large electromagnet with 0 to 2.5 T field is used, although much higher magnetic field also can be provided. The pole pieces of the electromagnet are such that a large uniform magnetic field prevails for a sample, may be a thin film or powder sample held in a glass capillary attached to a special drive.

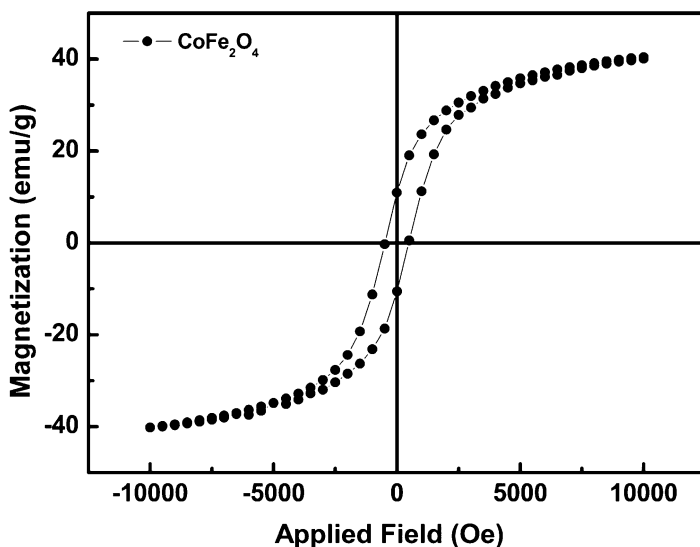


Fig. 7.51 M-H plot for CoFe₂O₄ nanoparticles

Paramagnetic, diamagnetic, ferromagnetic and other samples can be distinguished by plotting magnetization versus applied magnetic field. Provision can be made to heat or cool the samples during magnetic measurements. This enables one to study magnetic phase transitions.

In Fig. 7.51, the hysteresis curve obtained for CoFe₂O₄ nanoparticles obtained using a VSM is shown.

7.8 Mechanical Measurements

Mechanical properties like elastic properties, hardness, ductility or friction of different nanostructures need to be investigated. It should be, however, noted that measurements on single nanoparticles, rods or tubes would inherently be difficult, though not impossible. However measurements on nanocrystalline solids, thin films etc. are possible using some conventional methods. Techniques like nanoindentation are available. Before discussing about it let us first revise some of the concepts in mechanical properties.

7.8.1 Some Common Terminologies Related to Mechanical Properties

Mechanical properties of materials can be understood from Young's modulus, toughness and hardness.

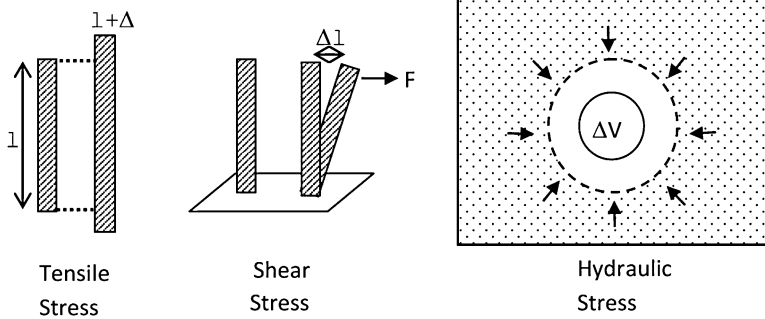


Fig. 7.52 Different types of stresses

$$\text{Young's modulus: } Y = \text{stress/strain} \quad (7.42)$$

It is a measure of elongation (deformation) in the limit of small stress.

Plastic deformation: When the material is elongated further, it may break. Before breaking, it can exhibit non-linear behaviour which can be described as plastic deformation.

Toughness: It is the amount of work necessary for mechanical failure of the material. It is thus a measure of energy that can be absorbed by the material. Both linear and non-linear contributions need to be taken into account to determine the toughness. Area under stress-strain curve is a measure of toughness.

Stiffness: Greater the modulus of elasticity, stiffer is the substance.

Ductility: It is the plastic deformation that can be sustained at fracture.

Hardness: Material's resistance to deformation or to produce indentation or abrasion.

Elastic Moduli: In some solids atoms are arranged such that the objects are stiff or rigid. For example, glass, table and ceramic cup are rigid. On the other hand, some objects are flexible like thread, plastic bag, garden hose in which atoms are arranged in long flexible chains. However all objects are elastic to some extent i.e. their dimensions can be changed (strain) by application of force (stress). There are various types of stresses viz. tensile stress, shear stress and hydraulic stress depending upon the direction of stress and strain as illustrated in Fig. 7.52.

Young's Modulus (Y):

$$Y = \frac{\text{Stress}}{\text{Strain}} \quad (7.43)$$

Tensile and compression strength can be different. For example cement has large compressive strength but poor tensile strength (Fig. 7.53).

Shear Modulus (G):

$$G = \frac{\text{Stress}}{\text{Strain}} \quad (7.44)$$

Fig. 7.53 Relation between stress and strain

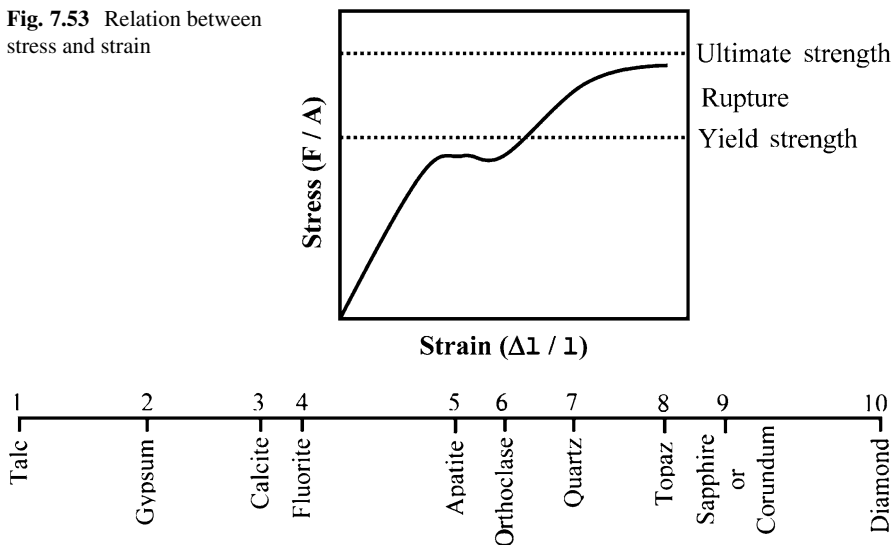


Fig. 7.54 Hardness of some standard minerals

Bulk Modulus (B):

$$B = \frac{P}{\Delta V / V} \quad (7.45)$$

Hardness: Hardness is one of the important mechanical properties of solid materials. It is a measure of a material's ability to resist local, plastic deformation.

There are different scales to measure hardness viz. Brinell hardness test, Rockwell hardness test, Knoop's hardness test, Vicker's hardness test and Mohs hardness test. Out of these, Mohs hardness test is one of the oldest and gives a scale from 1 to 10. One is for the softest material viz. talc and ten is for the hardest materials viz. diamond. The middle numbers, 2–9 are related to different minerals as standards of hardness as shown in Fig. 7.54.

The scale is based on the ability of harder material to scratch the softer material.

Other hardness tests are quantitative. They are based on the specific indentors used to scratch the material to be tested for the hardness. A controlled load determines the depth and size of the indentation in the material under test.

Vicker's Hardness (VHN) test is often used in which a small diamond indentor is used. It has a pyramid shape (See Fig. 7.55). The resulting impression of the diamond is observed under a microscope and used to determine the VHN.

Nanomechanical properties are measured using a 'nanoindenter' integrated with an atomic force microscope (AFM). In the conventional techniques, a hard tip is pressed in the material under certain known load and the load is removed after some time. The impression left behind in the material is investigated

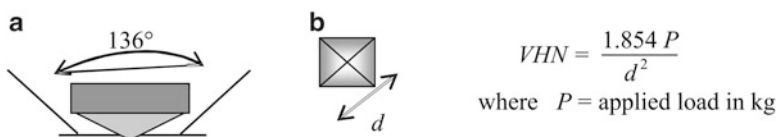


Fig. 7.55 (a) Pyramid shaped diamond tip and (b) top view of the typical impression of the tip in the material

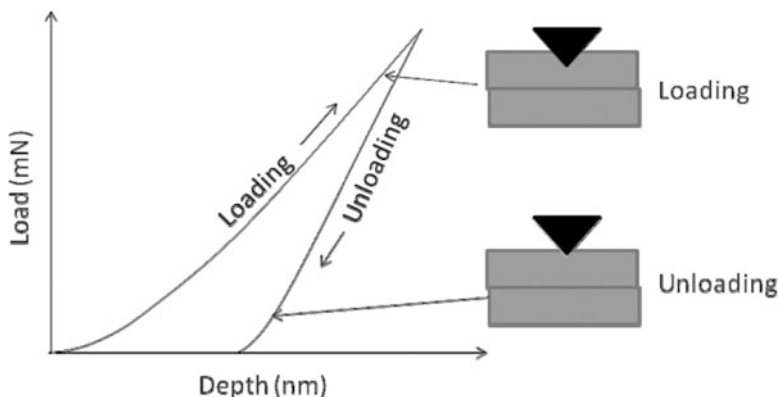


Fig. 7.56 Nanoindentation method

(area of the impression) using a suitable microscope. In nanoindentation, force and displacement of the tip in the material are continuously monitored. The depth of tip penetration as a function of load gives rise to a 'loading' curve (see Fig. 7.56). The tip movement is also monitored while 'unloading' leading to complete history of the displacement. Typical loads are upto about 500 mN and displacements upto 50 μm .

Using nanoindentors, it is possible to investigate mechanical properties like hardness, modulus of elasticity, creep and plastic flow.

Further Reading

- T.A. Carlson, *Photoelectron and Auger Electron Spectroscopy* (Plenum Press, New York, 1978)
 B.D. Cullity, S.R. Stock, *Elements of X-Ray Diffraction*, 3rd edn. (Prentice Hall, Upper Saddle River, 2001)
 H.H. Willard, L.L. Merritt Jr., J.A. Dean, F.A. Settle, *Instrumental Methods of Analysis*, 7th edn. (CBS Publishers, New Delhi, India, 1986)
 J.B. Wachtman, Z.H. Kalman, *Characterization of Materials* (Elsevier/Butterworth-Heinemann, Boston, 1993)

Chapter 8

Types of Nanomaterials and Their Properties

8.1 Introduction

It is an interesting question as to, starting with a few atoms, how the bulk materials reach their structure and related properties. Do they undergo any structural changes or even their smallest unit cell is similar to that in bulk material? This question has been addressed by many. There are reasons to believe that the small clusters or nanoparticles are not just the fragments of bulk materials. There can be entirely different structures as well as bonds and bond strengths in clusters which can even differ from nanomaterials.

It has been well established now that all the materials, may be metals, semiconductors or insulator clusters or nanomaterials, have size dependent physico-chemical properties. In most of the cases a cluster size is below 1 nm and that of nanoparticles is in the 1–100 nm range. Interestingly at such a small size even the shape of the material and interactions between clusters or nanomaterials decide the properties of the material. This opens up a huge possibility of tailor making the materials, which have different properties just due to their size, shape and/or assembly. In this chapter we will discuss to some extent clusters, semiconductor metal nanomaterials and magnetic nanomaterials. Nanomaterials would mean 0-D, 1-D or 2-D materials with the relevant dimensions smaller than ~ 100 nm to qualify them as nanomaterials. The general analysis methods for all the nanomaterials are as mentioned in Chap. 7. In this chapter some additional relevant analysis is discussed for semiconductor and metal nanoparticles. The chapter closes with some additional properties like mechanical, structural, electrical and thermal properties of the nanomaterials in general.

8.2 Clusters

Clusters are aggregates of small number of atoms and can be considered as intermediates of atoms and nanomaterials. Depending upon the number of atoms they can be roughly divided into following:

1. Microclusters – number of atoms between 2 and 10–13 atoms
2. Small clusters – number of atoms between 10 and 13 to about 100 atoms
3. Large clusters – number of atoms between about 100 and 1,000 atoms.

With number of atoms larger than 10^3 and below 10^6 atoms it is a nano-particle. Although the classification of clusters and nanomaterials cannot be very precise, usually clusters are smaller than ~ 1 nm.

One may also ask the question, what is the difference between a large atom, small molecule with couple of atoms and a cluster? For example a 92-atom sodium atom cluster and uranium atom with atomic number 92. A sodium atom cluster with one valence electron and 92 free electrons of sodium atom would make it appear like a uranium atom. However, in uranium atom the positive charge is concentrated in the nucleus in a very small volume and 92 electrons have some specific electron configuration. On the other hand, in a sodium atom cluster, each atom's positive charge is spatially separated from the other atoms and each one would have its ionic shell configuration; only 92 electrons would be free to form an electron gas of the cluster. As far as molecules are concerned they have stable configuration of certain types and number of atoms. They have fixed composition, structure, bond lengths, bond angles and, hence, fixed properties. The clusters on the other hand are the aggregates of atoms and may or may not be stable. The structure, bonds etc. vary with number of atoms and so do the properties. In clusters large fraction of atoms is on the surface. For example in a cluster of 55 atoms like that of argon or sodium 32 atoms are on the surface and only 23 are inside. Often we are only able to understand their presence in some experiments but may not be able to collect them. However, their theoretical as well as experimental investigations can lead to the understanding of the basic processes that take place while the clusters are formed and also lead to applications in catalysis, epitaxial materials and so on.

8.2.1 Types of Clusters

Besides the classification of clusters depending upon the number of atoms they contain, it is possible to realize various types of clusters as follows:

1. Homogeneous clusters – containing only single type of elements, e.g. Ar, Na, K, Si, Au, Pb etc.
2. Heterogeneous clusters – $3\text{Na} + 3\text{K}$, $2\text{Na} + 3\text{K}$ etc. $(\text{NaCl})_n$ or $\text{Na}_n\text{Cl}_{n+y}$, $(\text{Na}_{n+1}\text{Cl}_n)^+$, $(\text{Na}_n\text{Cl}_{n+1})^-$ and many other kinds and compositions.

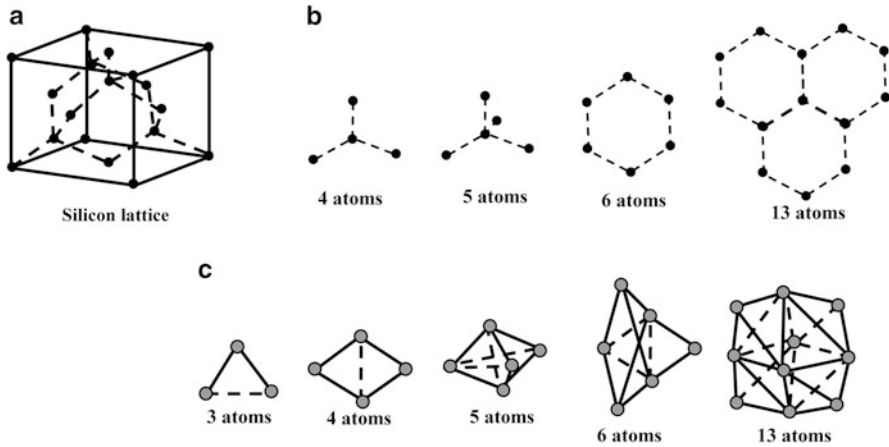


Fig. 8.1 (a) Silicon atoms forming a unit cell, (b) fragments of silicon unit cell, and (c) stable clusters of silicon

3. Oppositely charged clusters may be attracted and stay together.
4. Covalently bonded clusters like carbon, silicon etc. Fullerenes or C60 atom clusters are examples of such clusters and will be discussed in Chap. 9.

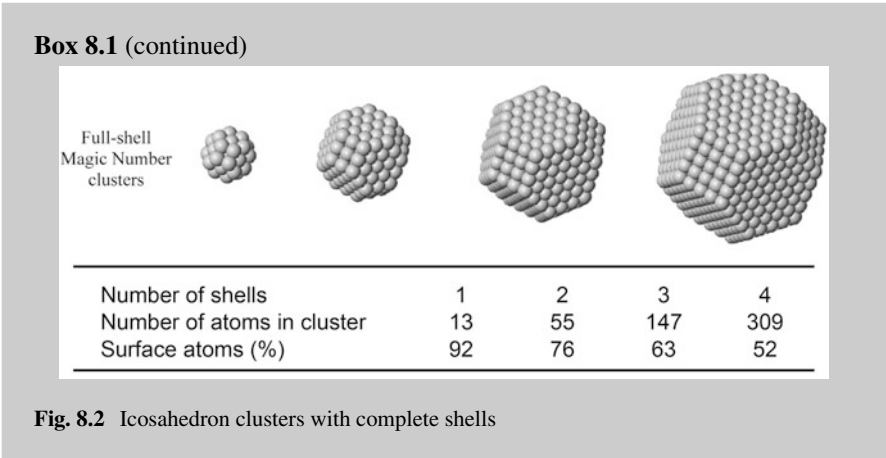
In a silicon cluster, the stable configurations are not the fragments of bulk lattice and instead of having 4, 5, 6 atoms in clusters, stable clusters having 3, 4 ... atoms with different structure rather than bulk fragmentation are energetically favoured (Fig. 8.1) (Box 8.1).

Box 8.1: Clusters and Platonic Solids

A cluster may be considered as being formed shell by shell around a central atom. As the number of shells increase, total number of atoms goes on increasing and percentage of surface atoms compared to interior goes on decreasing. The magic number of atoms in a cluster can be obtained by first considering some basic types which were formulated by ancient mathematicians like Pythagoras and Plato around 450 B.C. The basic shapes considered were cube, octahedron, icosahedrons, tetrahedron and dodecahedron related to earth, air, water, fire and heavenly constellation respectively. One of the common metal cluster is 13-atom icosahedrons. It has one central atom and 12 surrounding atoms, 20 faces, 12 vertices and 30 edges. Number of atoms in the n th shell of an icosahedron can be obtained from $10n^2 + 2$.

Figure 8.2 shows few icosahedrons clusters.

(continued)



5. Metal atom clusters i.e. a phase in which they would ultimately develop if the cluster and nanoparticle size is exceeded. In fact even before a material reaches a bulk size with some critical number of atoms characteristic of the material it undergoes various transitions assuming local stable sizes. This has been very well exemplified for potassium clusters (see Fig. 8.3). Here the ionization energy is plotted as a function of the number of atoms in a cluster. More the ionization energy more stable is the cluster. One can see that 8, 18, 20 ... atoms clusters are locally stable compared to the other neighbouring clusters. After having about 92

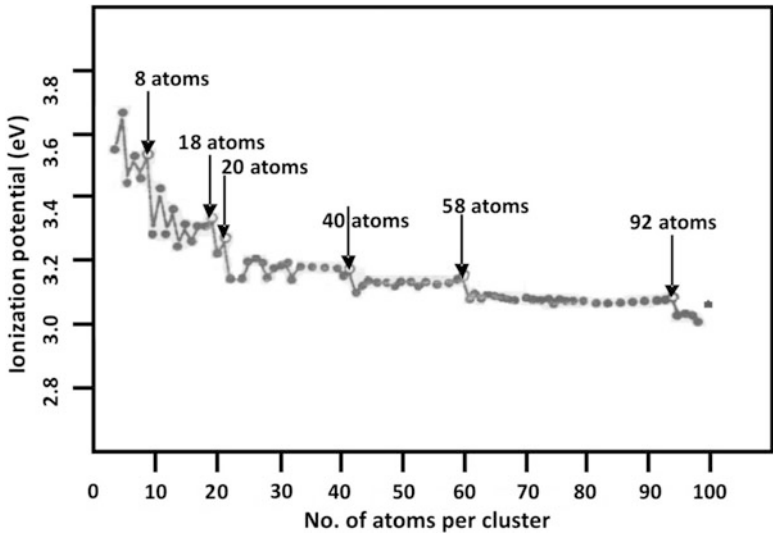


Fig. 8.3 Variation of ionization energies of different potassium clusters having different number of atoms in them

atoms, the cluster reaches the ionization potential (or work function value) of a bulk potassium metal. The clusters also can have different structures compared to the bulk material. One can get more insight about the metal clusters by using a jellium model.

6. Clusters of rare gases – close shell atom clusters of Ar, Kr etc. are held together by weak van der Waals forces and formed at low temperatures.
7. Clusters with hydrogen bonds, for example a $(\text{H}_2\text{O})_n$ cluster. It is a strongly held molecular cluster, stronger than cluster of rare gases but much weaker than metallic, covalent or ionic clusters.
8. Magnetic clusters – clusters of atoms with resultant magnetic moment. The magnetic moments on clusters like Co, Fe and Ni may show the moment close to magnetic moment/atom.

Usually shell theory applies to atoms, nuclei as well as clusters. In all the three cases fermions are confined in $\sim 10^{-12}$ cm for nuclei, 10^{-8} cm for atoms and $\sim 10^{-7}$ to 10^{-6} cm in case of clusters. Using appropriate angular momentum conditions for the spherical harmonic, square well or spherical potential shells are developed which are a direct result of application of quantum mechanics. In all these cases interacting fermions obeying Pauli exclusion principle result ultimately to shells with magic numbers or in case of clusters some stable clusters are referred to as ‘magic clusters’.

In order to obtain very small clusters, starting with 2–3 atoms upto few tens of atoms, special, very sophisticated machines are built. Usually an intense beam of laser is made incident on a bulk piece of material whose clusters are to be investigated. The evaporation normally contains clusters of different sizes which need to be mass selected using a mass spectrometer. The mass selected cluster would then have abundance of clusters of uniform size (same number of atoms per cluster). Such a cluster is then fragmented by another beam of laser which is then analyzed with another mass spectrometer. This analysis would then lead to the understanding of stability of the clusters. For example, if a four-atom cluster is less stable than a three-atom cluster, the fragmentation should show the abundance of three-atom clusters than four-atom clusters.

In some cases where ionization of the clusters is investigated, an oven is used as a source of clusters, where the source material from the hot oven is expanded through a nozzle, ionized in a chamber using photons or electron beam and then analyzed using a mass spectrometer. Of course in both the types of apparatus discussed here, adequate vacuum system is a must in order to avoid any contaminations.

8.3 Semiconductor Nanoparticles

Every material has a characteristic size below which size dependent properties are realized. In semiconductors this size is nothing but the size of the exciton. Understanding the concept of exciton and estimating its size for different semiconductor materials is the first step towards understanding the semiconductor nanoparticles.

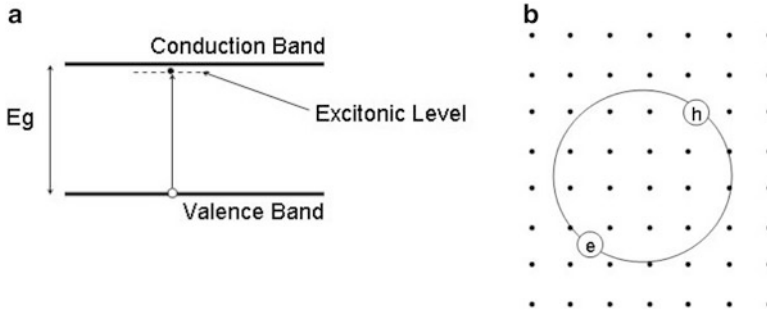


Fig. 8.4 (a) Energy level diagram for a typical semiconductor along with the excitation energy level. (b) Mott-Wannier exciton in a typical lattice

8.3.1 Excitons

In semiconductor or insulator, valence and conduction bands are separated by some finite energy gap characteristic of the material. When an electron from the valence band gets sufficient energy to overcome the energy gap, may be by thermal excitation or absorption of photon, and go to conduction band, a hole is left behind. The electron-hole pair so formed is a quasi-particle called exciton (see Fig. 8.4). An exciton, whose centre of mass motion is quantized, can move in the crystal.

Different kinds of excitons are identified in a variety of materials. Note that the spins of electron and hole in an exciton pair can be either parallel or antiparallel to each other. If the spins of hole and electron in an exciton are parallel to each other the 'dark exciton' results and 'bright exciton' otherwise. The properties of dark and bright excitons are different. As its name suggests dark exciton is optically inactive but the bright exciton is optically active. The optical exciton is useful in optical sensing and emission. Dark exciton, one may think, would not be useful. However, dark excitons and bright excitons can be interchanged by spin flip-flop and this has potential application in spin storage and qubits.

Depending upon the intensity of the light used to excite the excitons, it is also possible under intense illumination that a large number of excitons (exciton liquid) are formed which start interacting with each other. Two interacting excitons can attract each other forming a 'biexciton'. One can write

$$\delta E_{\text{exc}} = E_{\text{biexc}} - 2E_{\text{exc}} \quad (8.1)$$

Here E_{exc} is the energy of a single exciton, E_{biexc} is the energy of the two interacting excitons and δE_{exc} is the difference in the energies given by equation 8.1. Obviously the bound pair will be formed if the energy $\delta E_{\text{exc}} < 0$. Although there is a lot of interesting research being carried out to study the lifetime, dynamics and optical properties of biexcitons, we will not go into more depth here as it is beyond the scope of this book.

We shall consider here single (or widely isolated), optically active excitons. When the electron-hole pair is tightly bound with distance between electron and hole comparable to the lattice constant of the material, it is called Frenkel exciton. At the other extreme, one may have an exciton with electron-hole separation much larger compared to the lattice constant. Such a weakly bound electron-hole pair is called as Mott-Wannier exciton. The energy of such an exciton is slightly less than the energy gap (E_g) between valence and conduction band. In fact the energy gap E_g can be also defined as the energy necessary to create a free electron and a free hole. The Hamiltonian for Mott-Wannier exciton is given as

$$H = \frac{P_e^2}{2m_e} + \frac{P_h^2}{2m_h} - \frac{e^2}{\epsilon |r_e - r_h|} \quad (8.2)$$

where m_e and m_h are effective electron and hole mass respectively, and ϵ is the dielectric constant of semiconductor material. It should be remembered that the mass of an electron and hole in a solid material and even a nanomaterial is not same as that of a ‘free electron’ mass. It may be smaller or larger than the free electron mass depending upon the curvature of a band. For more details one should refer to a text book on Solid State Physics. First two terms on the right hand side of Eq. (8.2) are kinetic energies of electron and hole respectively and the third term is the Coulomb energy of electron-hole attraction or potential energy.

The Bohr radius of such an exciton is given as

$$r_B = \frac{\hbar^2 \epsilon}{e^2} \left[\frac{1}{m_e} + \frac{1}{m_h} \right] \quad (8.3)$$

In Table 8.1 radii for few semiconductor spherical nanoparticles are illustrated.

It may be noticed that these values of r_B are of the order of a few nanometres. We observe in a nanoparticle and other nanostructures a situation in which exciton is confined in the nanostructure. The sizes of excitons are often comparable to the sizes of nanoparticles or nanostructures.

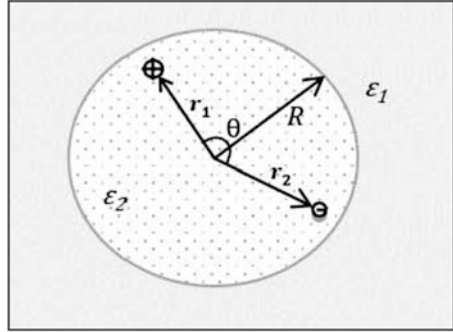
8.3.2 Effective Mass Approximation

There are various theories like Effective Mass Approximation (EMA) and Tight Binding (TB) developed to explain the size dependent energy gap variation (being most important property of semiconductor materials on which other properties

Table 8.1 Bohr radius (r_B) for exciton in some semiconductors

Semiconductor	E_g (eV)	m_e	m_h	ϵ	r_B (nm)
GaAs	1.52	0.067	0.082	13.1	18.8
InP	1.42	0.077	0.64	12.4	9.6
InSb	0.24	0.0145	0.39	17.6	66.7

Fig. 8.5 Spherical semiconductor nano-particle embedded in a medium with different dielectric medium



depend) in nanostructures. Depending upon the material or group of materials one theory may be more applicable than the other. Nonetheless, EMA works out to be reasonably well in most of the cases. It nicely explains the dependence of energy gap opening with the reduction of size and will be discussed here. Historically, Efros and Efros first formulated this theory which was extensively developed by L.E. Brus. Briefly it is outlined below.

Consider a dielectric sphere with dielectric constant ϵ_2 embedded in a medium of dielectric constant ϵ_1 . Let the radius of the sphere be R (see Fig. 8.5).

If $\mathbf{r}_1$ and $\mathbf{r}_2$ are the positions of the charges inside the sphere, the potential energy can be written as

$$V(\mathbf{S}_1, \mathbf{S}_2) = \pm \frac{e^2}{\epsilon_2 |\mathbf{r}_1 - \mathbf{r}_2|} + \mathbf{P}(\mathbf{r}_1) \pm \mathbf{P}(\mathbf{r}_2) \pm \mathbf{P}_M(\mathbf{r}_1, \mathbf{r}_2) \quad (8.4)$$

where polarization $\mathbf{P}$ is given as

$$\mathbf{P}(\mathbf{r}) = \sum_{n=0}^{\infty} \alpha_n \left(\frac{r}{R} \right)^{2n} \frac{e^2}{2R} \quad (8.5)$$

and

$$\alpha_n = \frac{(\epsilon_2 - 1)(n + 1)}{\epsilon_2(\epsilon_2 n + n + 1)} \text{ with } \epsilon = \frac{\epsilon_2}{\epsilon_1}$$

$$\mathbf{P}_M(\mathbf{r}_1, \mathbf{r}_2) = \sum_{n=0}^{\infty} \alpha_n \frac{\epsilon_2 + r_1^n r_2^n}{R^{2n+1}} P_n(\cos \theta) \quad (8.6)$$

Here, θ is the angle between $\mathbf{r}_1$ and $\mathbf{r}_2$ and P_n is a Legendre polynomial.

When R is larger than r_1 and r_2

$$V \rightarrow \frac{e^2}{\epsilon_2 |\mathbf{r}_1 - \mathbf{r}_2|} \quad (8.7)$$

This is the potential energy (third) term given in Eq. (8.2) while writing the Hamiltonian term for the Mott-Wannier exciton.

We can now write the Schrödinger equation as

$$\left[-\frac{\hbar^2}{2m} \nabla_e^2 - \frac{\hbar^2}{2m} \nabla_h^2 + V_0(\mathbf{r}_e, \mathbf{r}_h) \right] \Phi(\mathbf{r}_e, \mathbf{r}_h) = E \Phi(\mathbf{r}_e, \mathbf{r}_h) \quad (8.8)$$

Consider $V_0 = \infty$ outside the sphere.

For small values of R (nanoparticles) we consider wave function $\psi_n(\mathbf{r})$ of the type

$$\psi_n(\mathbf{r}) = \frac{Cn}{r} \sin\left(\frac{n\pi r}{R}\right) \quad (8.9)$$

with

$$E_n = \frac{n^2 \pi^2 \hbar^2}{2mR^2} \quad (8.10)$$

Here m is the effective mass of the charge (electron or hole).

Consider

$$\Phi_0 = \psi(\mathbf{r}_e) \psi(\mathbf{r}) \quad (8.11)$$

as the wave function for Eq. (8.8).

In such a case the energy of the first excited state would be

$$\Delta E = \frac{\hbar^2 \pi^2}{2R^2} \left[\frac{1}{m_e} + \frac{1}{m_h} \right] - \frac{1.8e^2}{\epsilon_2 R} + \text{polarization energy} \quad (8.12)$$

The first term on the right hand side of the equation is localization energy of electron and hole whereas the second term is due to Coulomb interaction energy between electrons and hole. The third term is due to polarization of the cluster but is usually very small and size independent, hence neglected. Details of the formulations are not given here but can be found in the original paper by L.E. Brus. The first two terms have signs opposite for each value of R . But as the particle size decreases, the first term starts dominating and ΔE starts increasing rapidly. This would effectively mean that the energy gap in a semiconductor which is characteristic of a material (e.g. CdS $E_g = 1.42$ eV, ZnS $E_g = 3.6$ eV, GaAs $E_g = 1.42$ eV) would increase with decreasing particle size or in other words we obtain size dependent energy gap which is

$$\text{Total energy} = E_{g(\text{bulk})} + \Delta E (\text{exciton}) \quad (8.13)$$

Obviously, ΔE becomes important when cluster size is comparable to the excitation size. This can be seen from the example given in Table 8.2, where ΔE for various values of R for GaAs are shown.

Table 8.2 Kinetic energy, Coulomb energy, energy shift (ΔE) and total energy gap calculated using Eq. (8.11) for spherical particles of GaAs of various radii

Energies (eV)	Particle radius (nm)			
	20	10	5	2
Kinetic energy	0.015	0.059	0.24	1.48
Coulomb energy	−0.012	−0.024	−0.489	−0.119
Energy shift (ΔE)	0.003	0.039	0.19	1.361
Total energy gap (E_g)	1.42	1.46	1.61	2.78

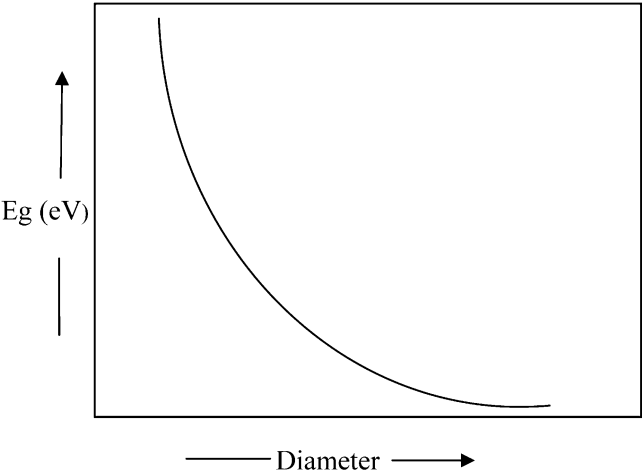


Fig. 8.6 Schematic representation of variation of energy gap with particle size

8.3.3 *Optical Properties of Semiconductor Nanoparticles*

The simplest experiment to determine the size dependence in semiconductor nanoparticles is to study absorption spectrum of the material as a function of wavelength of incident photons. When photons are incident on semiconductor material they will be absorbed only when the minimum energy of photons is enough to excite an electron from the valence band to conduction band, i.e. when the photon energy equals the energy gap of the semiconductor. If lower energy photons are incident, there cannot be any absorption. Therefore there is a sudden rise in absorption when the photon energy is same as the energy gap. This is the onset of absorption. If the energy gap increases there should be a shift in the onset of absorption towards the shorter wavelength. As shown in Fig. 8.6, one expects a blue shift with absorption in smaller and smaller particles, which is indicative of increasing energy gap. For every small particle (see Fig. 8.7) one often uses the terminology borrowed from chemistry, viz. Highest Occupied Molecular Orbital (HOMO) and Lowest Unoccupied Molecular Orbital (LUMO) instead of top of valence band and bottom of conduction band used in case of extended solid.

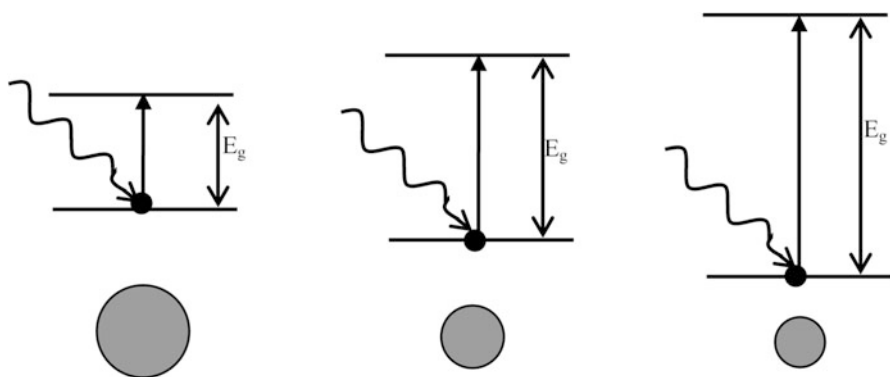


Fig. 8.7 With decreasing particle size, energy gap increases. This would give rise to blue shift in the absorption spectra

This change in absorption would have an interesting effect on the originally coloured materials. It is known now from literature that Cd_3P_2 is a dark brown semiconductor with energy gap of approximately 0.5 eV. When its particles are made, it progressively passes through a series of colours like brown, red, yellow and white with particle size changing from ~ 30 to ~ 15 Å. For ~ 15 Å particles the band gap increases to 4 eV.

The same is true for CdS (Figure in the preface of this book). The bulk semiconductor with energy gap of 2.42 eV is orange in colour. As the particles become smaller and energy gap increases it becomes yellowish and ultimately white. It is quite easy to show by chemical analysis techniques that this white material is CdS and nothing else. In fact, observation of different colours due to CdSe in glass matrix led scientists to think that CdSe nanoparticles of different sizes might have been formed.

Many nanomaterials exhibit enhanced luminescence as compared to their bulk counterparts. Some materials like silicon which are not luminescent in their bulk form become luminescent in the nano form. Therefore luminescence investigations of nanomaterials are often carried out and are quite interesting.

We shall discuss here the principles of (a) photoluminescence, (b) electroluminescence, (c) cathodoluminescence and (d) thermoluminescence.

8.3.3.1 Photoluminescence

When the external stimulus is electromagnetic radiation or photons, the observed luminescence is known as photoluminescence. An electron from the valence band can be excited to a level in the conduction band if photon of sufficient energy to make a transition is available. This process leaves a hole in the valence band as shown in Fig. 8.8. The excited electron can lose energy by emission of phonon in a relatively shorter time ($\sim 10^{-12}$ – 10^{-13} s) before it can relax and make a radiative

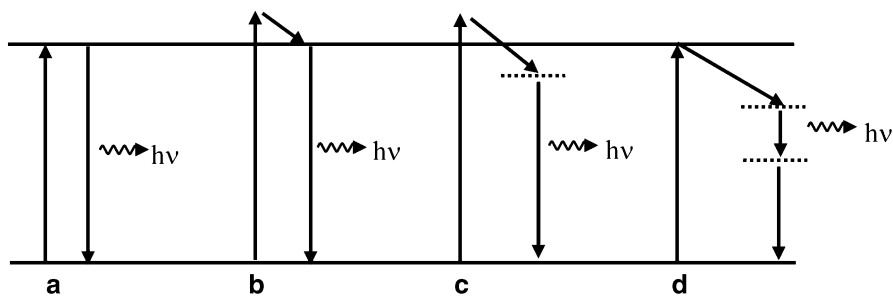
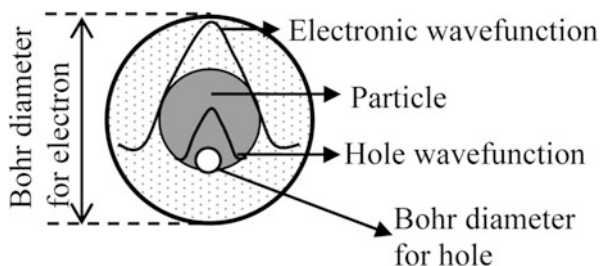


Fig. 8.8 Various processes of luminescence

Fig. 8.9 Localization of electron and hole wavefunction inside a nanoparticle



transition (Fig. 8.8b). The life time of radiative process is much longer $\sim 10^{-7}$ to few milliseconds (fluorescence) or even few seconds (phosphorescence) than the life time of nonradiative transition.

Various possibilities of luminescence mechanisms are illustrated in Fig. 8.8. When emission of light due to the transition of an electron from the valence band maximum to the conduction band minimum, as shown in Fig. 8.8a, takes place it is called the band edge emission. The energy of emitted light or photoluminescence is the energy difference between the conduction band minimum and the valence band maximum. It is also possible as in Fig. 8.8c that the electron is first excited to the conduction band from where it moves in the material and gets trapped in some localized level. Subsequently it makes a radiative transition to valence band emitting the photon of a longer wavelength compared to that emitted by band edge luminescence. If there are localized levels introduced by impurities in the energy gap then non-radiative transition from conduction band minimum to impurity level occurs, followed by radiative transition to the lower levels (Fig. 8.8d) and then another non-radiative transition to valence band maximum.

The wavelength of luminescence, luminescence efficiency and radiative lifetime depend upon the material. In case of nanomaterials additionally the size of the material plays an important role. This can be readily understood by considering the schematics in Fig. 8.9 shown for a spherical nanoparticle.

In a particle with dimension comparable to Bohr diameter of electron, electron and hole wavefunctions can overlap. Therefore the probability of efficient recombination increases. Due to altered electronic structure in nanomaterials, the transition

rules also relax in certain cases and some transitions which are forbidden in bulk material become possible. The fluorescent efficiencies in general are found to be enhanced in nanoparticles. Interestingly, doping of nanoparticles also is possible. Dopants can introduce some energy levels in the band gap of the host material as discussed above. Such levels can be responsible for altering the luminescence properties of the host material. As an example, consider the doping of ZnS nanoparticles with manganese. ZnS is a wide direct band gap material. The band gap luminescence of bulk ZnS is at ~ 344 nm. If Mn is introduced in small quantity (~ 0.1 at. %) it produces discrete energy levels in the band gap. A photoexcited electron can make a transition from conduction band minimum to upper level and then combine with a trapped hole in lower level to emit luminescence at ~ 590 nm. In nanoparticles also such luminescence is observable but with more probability or higher luminescence efficiency, due to increased overlap, arising because of small size in nanoparticles, of wavefunctions. The process can be accompanied by band edge luminescence and/or defect emission.

8.3.3.2 Electroluminescence

Luminescence observed by the application of an electric field to a material is known as electroluminescence. It can be observed by applying either low or high ($>10^4$ V/cm) field; accordingly it is classified as ‘injection luminescence’ and ‘high field electroluminescence’ respectively.

Injection luminescence: Light emitting diodes are based on the principle of minority carrier injection in a diode. When a p-n diode is formed, a depletion region is set up with V_d as the diffusion potential between the p-n regions (see Fig. 8.10).

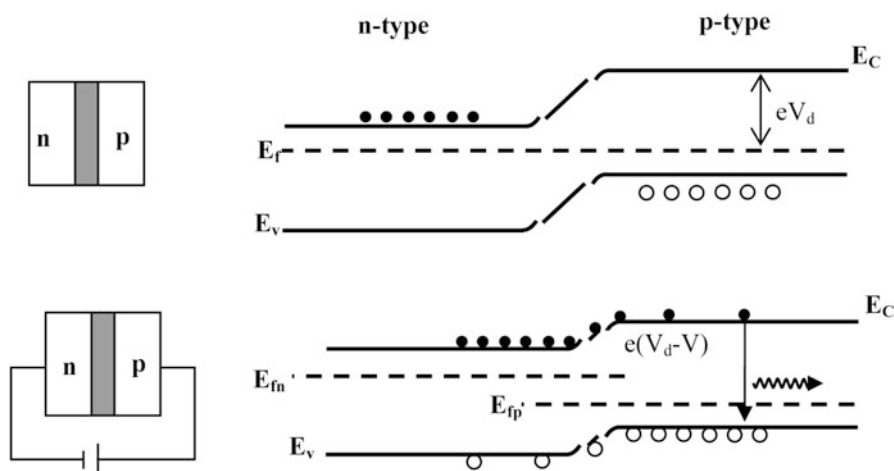


Fig. 8.10 p-n junction at thermal equilibrium (*upper*) and with forward bias

In a p-type semiconductor, holes are the majority carriers and electrons are minority carriers. On the other hand, in an n-type semiconductor, electrons are majority carriers and holes constitute minority carriers.

By forward biasing the diode, the diffusion barrier can be reduced and potential becomes $V_d - V$, where V is the applied voltage. Under this condition, the electrons from n-region find it easier to travel to p-region and holes can easily migrate from p-region to n-region. Thus minority carriers are injected on each side. The injected carriers can combine with easily available opposite charge carriers in each region to produce luminescence.

It can be observed that the injection luminescence is a 'band edge' type luminescence. Therefore for a fixed semiconductor, the wavelength of luminescence is fixed. However in nanomaterials we know that energy gap can be tuned with particle size. Hence the luminescence also can be tuned to desired wavelength in nanomaterials.

High field electroluminescence: This type of electroluminescence is used in 'display panels'. Emission of electrons by application of very high electric field ($\sim 10^5 - 10^7$ V/cm) is known as *field emission*. As illustrated in Fig. 8.11 an energy barrier existing for a material has to be overcome by an electron in order to leave it. The barrier, however, reduces on application of the electric field, the width of which depends on the applied field. As can be seen from the figure, more the applied field, smaller is the width of the barrier. In general, an electron needs energy sufficient to overcome the barrier in general. However quantum mechanically, if the barrier width is reduced sufficiently, the electrons can tunnel through the barrier. This type of field emission requires very high field as mentioned above.

In a metal semiconductor contact, the barrier height is reduced due to transfer of electrons either from metal to semiconductor or semiconductor to metal depending upon the type (n or p) and value of work functions of both metal and semiconductor. A metal-semiconductor contact is known as Schottky barrier. Electrons can be

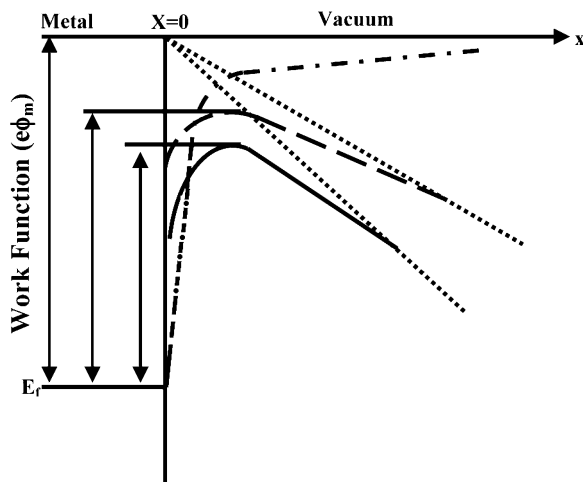


Fig. 8.11 Field emission from a metal

injected from metal to p-type semiconductor by high electric field. Note that here electrons are emitted from metal into the semiconductor by a different phenomenon as compared to that in injection luminescence. In injection luminescence, charge carriers overcome the potential barrier and penetrate the other material whereas in high field luminescence it is through tunneling as the electrons pass from metal to semiconductor. After passing from metal to semiconductor, if the electron comes in the vicinity of a hole it combines radiatively.

High field luminescence in nanomaterials: Nanomaterials are useful in display panels using high field. The fine particles resolve the images better as compared to microparticles. However due to high quantum efficiency of luminescence, electroluminescence also is an efficient process.

8.3.3.3 Cathodoluminescence

Electrons of very high energy striking a semiconductor material produce luminescence known as 'cathodoluminescence'. The incident electrons here are from some filament or field emission cathode. They strike the luminescent material in vacuum at high energy. The electrons on impact with the material are able to lose their energy by various processes. Phenomenon of cathodoluminescence is used in oscilloscope and old television display screens.

Nanomaterials are also useful to produce high efficiency cathodoluminescence due to the same reasons as for high field luminescent materials. Cathodoluminescence in ZnS doped with Mn and ZnO nanoparticles has been investigated. It has been observed that nanoparticles are better luminescent materials than microparticles of the respective material.

8.3.3.4 Thermoluminescence

In semiconductors with large band gaps it is found that if they are excited at very low temperatures with photons in the UV range, on heating to some temperature which depends upon the dopant ions, light is emitted even in the absence of any other stimulus. The phenomenon is known as 'thermoluminescence' or 'after glow'.

As illustrated in Fig. 8.12, luminescence is due to trapped electrons. The trapped electrons are activated by thermal energy and can combine with hole making radiative transition or causing luminescence. Figure 8.12 also gives some examples in which luminescence due to different co-activation is observed for ZnS doped with Cu. Observation of thermoluminescence is an effective method of studying the trap levels in materials.

Nanomaterials also have defect levels. The surface atoms can act as efficient traps for electrons/holes. Therefore thermoluminescence can be quite strong in nanomaterials. Thermoluminescence has been reported for ZnS nanoparticles doped with copper and other co-dopants (or co-activators) illustrated in Fig. 8.12.

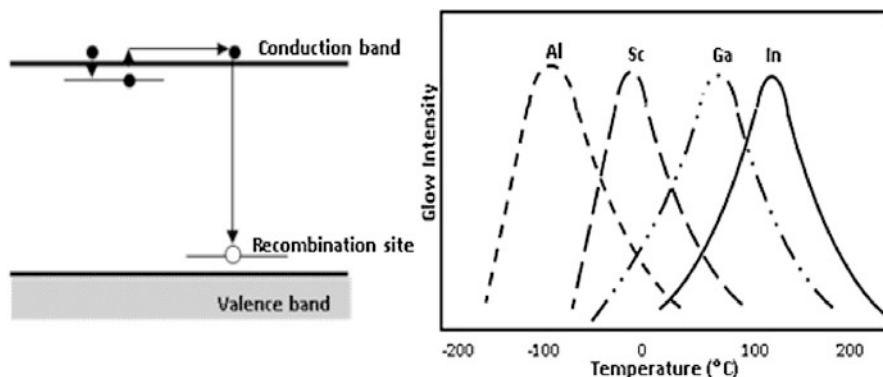


Fig. 8.12 Thermoluminescence energy level diagram

8.4 Plasmonic Materials

Windows in old churches, palaces, houses etc. are designed with beautiful tinted glasses. Such glasses are made by dissolving small amount (<5 %) of metal particles like gold, silver, cobalt, iron and nickel. Such glasses are transparent but have different colours like red, pink, blue, green or other shades depending upon the dissolved metal particles. Indeed the colour of glasses is due to metal nanoparticles. Note that the colours produced by dissolving the metal powders are not same as their bulk metal colour. For example gold metal has yellow 'golden' colour in the bulk form but the colour produced by dissolving small amount of gold in glass may be reddish, pink or even blue. Although the art of making tinted glasses is more than 2,000 years old, scientific interest in metal nanoparticles started with M. Faraday's synthesis of gold nanoparticles in 1857. He reduced chloroauric acid (HAuCl_4) using citric acid [$\text{CH}_2(\text{COO})_2\text{H}_2\text{O}$] and showed that Au metal nanoparticles produced intense magenta-red colour as against yellow appearance of bulk gold metal. Attempts were since then made to explain the observed intense colours due to metal nanoparticles. Apart from their beautiful appearance, it is now realized that metal nanoparticles have applications in sensors, solar cells, cancer therapy, cloaking devices and photonic devices for communication or optical circuits. Here we will try to understand the theories developed to explain the observed behaviour of metal nanostructures. It should be remembered that metal nanoparticles differ from the semiconductor particles. Metals have a large free electron density of valence electrons $\sim 10^{22}\text{--}10^{23}/\text{cm}^3$. On the other hand semiconductors have much lower electron density $\sim 10^{16}\text{--}10^{18}/\text{cm}^3$ depending upon the level of doping. Doping is a process of introducing electron donor or acceptor atoms in small quantities in order to make the n or p type semiconductors. Therefore in semiconductor nanomaterials the observed properties are explained based on the electron confinement in a potential well of appropriate shape depending on the nanomaterial dimensionality (0-D, 1-D or 2-D). In metal nanostructures a

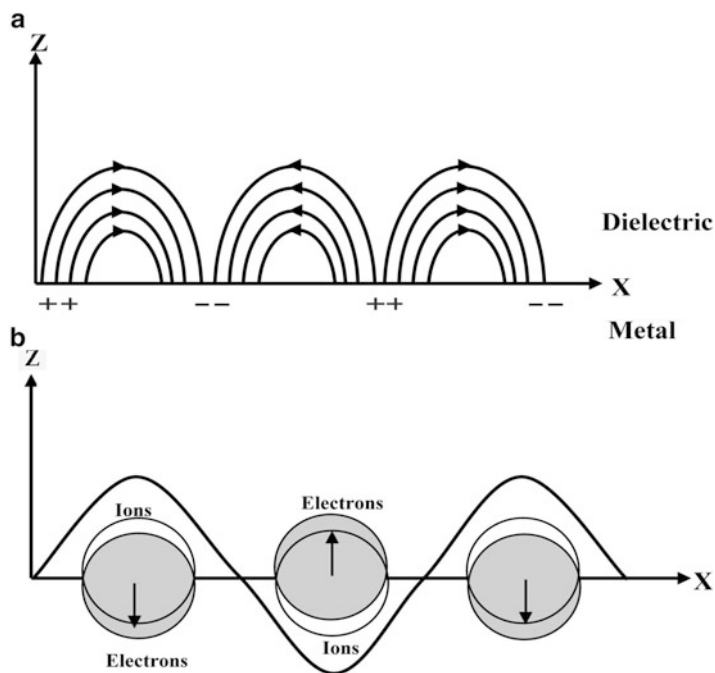


Fig. 8.13 Schematic of SPR: (a) Propagating SPR and (b) Localized SPR

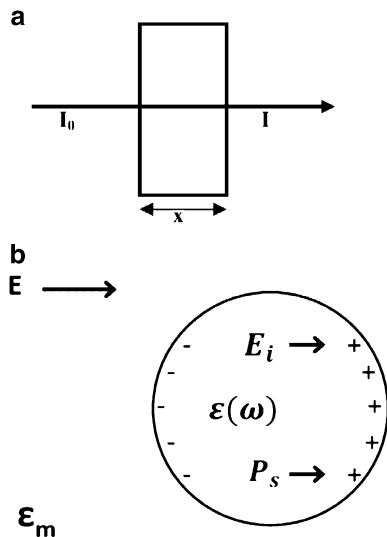
large number of electrons are considered together and the confinement is often known as ‘dielectric confinement’ rather than ‘electron confinement’. We now proceed to discuss how to describe observed behaviour of metal nanoparticles and nanostructures.

The small particles of different shapes like spherical particles, wires, rods, cubes and thin films or surfaces of metals need to be treated differently. The small particles exhibit *localized surface plasmon resonance* (SPR) and surfaces of metals or thin films have propagating surface waves viz. *surface plasmon polariton* as separately discussed below and schematically illustrated in Fig. 8.13.

8.4.1 Localized Surface Plasmon Resonance

In 1908, G. Mie explained by using Maxwell’s equations the scattering (and absorption) of light from very small particles. When electromagnetic radiation is incident on the spherical particles of uniform size, embedded in a medium, reduction in intensity is observed. It is necessary to consider mainly the dielectric constant of medium in which the particles float and the dielectric constant of the particles. Interaction between the particles is neglected in Mie theory (Fig. 8.14). His theory can be briefly described as follows. The details can be found in some standard text books on optics.

Fig. 8.14 (a) Incident radiation I_0 passes through a medium of thickness x and (b) A particle embedded in a dielectric medium (Mie theory)



A beam of electromagnetic radiation of intensity I_0 and wavelength λ passes through a medium having dielectric constant ϵ_m . If the particles are embedded uniformly in the medium and multiple reflections do not take place, the transmitted intensity would be given by D'Lambert equation

$$I = I_0 e^{-\mu x} \quad (8.14)$$

where μ is extinction coefficient.

$$\mu = \frac{N}{V C_{ext}} \quad (8.15)$$

N is the number of particles in medium, V – the volume of colloidal particles and C_{ext} is extinction cross section of a particle.

$$C_{ext} = C_{abs} + C_{scatt} \quad (8.16)$$

In the Mie theory

$$C_{ext} = \frac{2\pi}{k^2} \sum_{n=1}^{\infty} (2n+1) R_e(a_n + b_n) \quad (8.17)$$

where a_n and b_n are scattering coefficients. They are functions of a and λ in terms of Ricatti-Bessel functions. The details of the Mie theory can be found in the book by Born and Wolf (see the complete reference at the end of this chapter).

$$k = \frac{2\pi\sqrt{\epsilon_m}}{\lambda} \quad (8.18)$$

For very small particles ($kR \ll 1$) having radius R , extinction is mainly due to absorption. The extinction coefficient is given by

$$C_{\text{ext}} = \frac{24\pi^2 R^3 \varepsilon_m^{\frac{3}{2}}}{\lambda} \frac{\varepsilon_2}{(\varepsilon_1 + 2\varepsilon_m)^2 + \varepsilon_2^2} \quad (8.19)$$

where

$$\varepsilon(\omega) = \varepsilon_1(\omega) + i\varepsilon_2(\omega) \quad (8.20)$$

$\varepsilon(\omega)$ is the dielectric constant of the particles. As can be seen from above equation, C_{ext} depends on R^3 . Absorption coefficient μ is inversely proportional to $1/V$ i.e. $1/R^3$. Thus absorption is independent of particle size. From the equation (8.19), it is also seen that extinction would be maximum when

$$\varepsilon_1 + 2\varepsilon_m = 0 \text{ or } \varepsilon_1 = -2\varepsilon_m \quad (8.21)$$

if ε_2 is small. This gives rise to strong resonance band. Bandwidth and peak height is mainly determined by $\varepsilon_2(\omega)$ (interestingly C_{ext} vanishes if $\varepsilon_2 = \infty$ as well as $\varepsilon_2 = 0$).

Although Mie theory successfully explains the observation of an absorption band for metal nanoparticles in the visible range of wavelengths, for particle size less than ~ 10 nm, it hardly explains the observed shifts in the absorption peaks with change in the particle size. One, therefore, should consider dielectric constant which depends not only on the wavelength (or frequency) but also on the particle size. For metals the dielectric response of electrons is better described by Drude model. The theory of LSPR is based on the Drude model of dielectric response of free electrons in a metal. Consider that the electric field vector of an electromagnetic wave with frequency ω is written as

$$E = E_0 \exp(-i\omega t) \quad (8.22)$$

on a metal particle. The electrons of mass m are then accelerated. The equation of motion is given as

$$m \frac{d^2 x}{dt^2} + \gamma \frac{dx}{dt} + m\omega_0^2 x = eE e^{-i\omega t} \quad (8.23)$$

An electron in a metal experiences collisions (which give rise to resistivity in materials) with lattice and impurities. The relaxation time of an electron is $\tau = 1/\gamma$. Solving the Eq. (8.23) we get

$$x = \frac{eE/m}{\omega_0^2 - \omega^2 - i\omega_d \omega} \quad (8.24)$$

where $\omega_0^2 = f/m$, in which 'f' is the restoring force constant, and $\omega_d = \gamma/m$ is the damping. The resonance frequency here is ω_0 and it may or may not be equal to the frequency of the incident radiation.

The electron displacement 'x' is in the same direction of applied field when the resonance frequency ω_0 is more than the incident frequency ω . The displacement can be negative if $\omega > \omega_0$.

Consider that the displacement x is not just for a single electron but for N electrons/cm³ in a particle, then they would create polarization $\mathbf{P}$

$$\mathbf{P} = N e \mathbf{x} \quad (8.25)$$

Polarization can also be written in terms of the dielectric function as

$$\mathbf{P} = \epsilon_0 (1 - \epsilon) \mathbf{E} \quad (8.26)$$

Using the equations

$$\epsilon(\omega) = 1 + \frac{Ne^2/m\epsilon_0}{\omega_0^2 - \omega^2 - i\omega_d\omega} \quad (8.27)$$

and

$$\epsilon(\omega) = \epsilon'(\omega) + i\epsilon''(\omega) \quad (8.28)$$

where ϵ' and ϵ'' are the real and imaginary parts of the dielectric functions respectively. Therefore,

$$\epsilon'(\omega) = 1 + \frac{(\omega_0^2 - \omega^2) \omega_p^2}{(\omega_0^2 - \omega^2)^2 + \omega^2 \omega_d^2} \quad (8.29)$$

$$\text{When } \omega_0 = 0 \text{ and } \epsilon(\infty) = 1, \omega_p^2 = \frac{4\pi Ne^2}{m} \text{ is the volume plasma,} \quad (8.30)$$

$$\epsilon'(\omega) = \epsilon(\infty) - \frac{\omega_p^2}{\omega^2 + \omega_d^2} \quad (8.31)$$

and

$$\epsilon''(\omega) = -\frac{\omega_d \omega_p^2}{\omega(\omega^2 + \omega_d^2)} \quad (8.32)$$

$\omega_0 = 0$ is possible when the electrons do not have any restoring force. $\epsilon(\infty)$ is the dielectric constant at very high frequency.

Fig. 8.15 Real and imaginary parts of the dielectric function with respect to the frequency

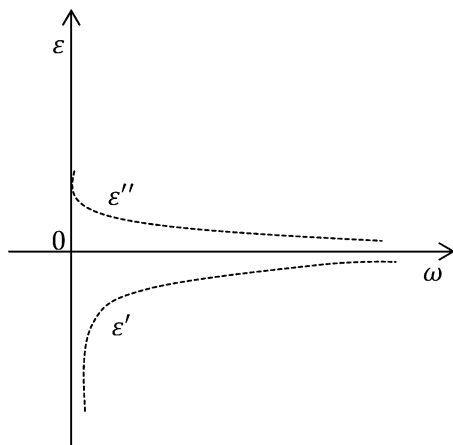
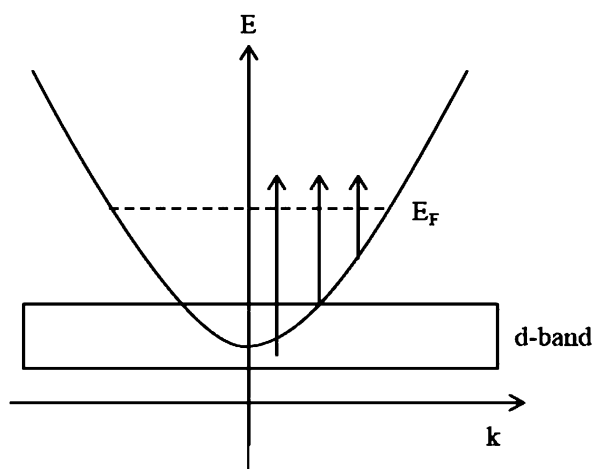


Fig. 8.16 Schematic diagram of various transitions from filled inter and intra bands in a metal



When $\omega_0 = 0$, the electrons are out of phase with the electric field $\mathbf{E}$. This implies that the real part of the dielectric constant would always be negative. This is schematically shown in Fig. 8.15.

Additionally, as shown schematically in Fig. 8.16, the electrons from the conduction band (s) and the filled band (d) to empty states above the conduction band also get photo-excited.

Thus the dielectric constant may get further modified by interband transitions. In silver and gold, the real part of the dielectric constant gets modified due to contribution from the 'd' and 'sp' bands. In the visible and ultra-violet regions, interband transitions contribute substantially whereas the 'red' or the long wavelengths derive contributions from intraband transitions.

Further, one needs to take into account that at small sizes of metal particles, the collisions of electrons with surface of the particle would increase. Therefore, the damping term γ would become

$$\gamma = \gamma_0 + g \frac{v_{F0}}{D} \quad (8.33)$$

where g is a proportionality factor (~ 1), v_F – the velocity of electrons at the Fermi surface, and D is the diameter of the particle. If D is much smaller than the wavelength of incident radiation and the resonance condition given by $\epsilon' + 2\epsilon_m = 0$, is satisfied, the resonance will occur at

$$\Omega(r) = \frac{\omega_p}{\sqrt{\epsilon'_{\text{interband}} + 2\epsilon_m}} \quad (8.34)$$

One can see from the above equation and the discussion in this section, how the dielectric constant of the metal particles (alongwith the dielectric constant of the medium in which they are dispersed) determines the position of the SPR peak. Therefore the phenomenon of SPR is aptly referred to as the ‘dielectric confinement’. This equation works well for silver, but complications arise for gold and copper due to the overlap of interband transitions with the surface plasmon band. Thus in order to explain the observed difference in the surface plasmon resonance positions of gold and silver nanoparticles in spite of almost same lattice constants and electron density (hence ω_p which is 8.98 eV for silver and 9.01 eV for gold), it is necessary to remember that interband transitions need to be taken into account. In Table 8.3 values for some metals in vacuum ($\epsilon_m = 1$) are given.

For the silver nanoparticles, surface plasmon resonance peaks are much sharper compared to those due to gold nanoparticles. This is due to the occurrence of silver plasmon resonance peak where imaginary part of dielectric constant (ϵ'') has nearly zero value. This is not the case for gold, for which absorption due to interband transition for ϵ'' has finite value.

Interestingly only few metals like gold, silver, copper and palladium nanoparticles show surface plasmon resonance in the visible range or close to visible range. There are many metals which surprisingly do not exhibit strong resonance peaks in the visible range. This is because they have interband transitions spread over UV-VIS-IR region which overlap with their resonance energies making them difficult to observe.

Table 8.3 Values for some metals: $\hbar\Omega(\text{ib})$ – interband energy and $\hbar\Omega(r)$ – surface plasmon energy

	m_e (effective mass of electron)	N (e/m^3)	l (nm)	v_F (10^6 m/s)	ω_p (eV)	$\hbar\Omega(\text{ib})$ (eV)	$\hbar\Omega(r)$ (eV)
Ag	1.03	5.85	55	1.4	8.98	3.9	3.5
Au	1.01	5.9	42	1.4	9.01	2.3	2.4
Cu	1.17	8.48	28	1.05	9.05	2.1	2.0

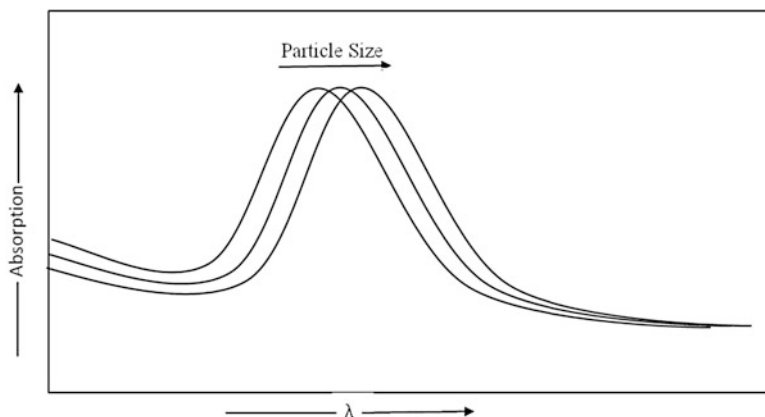


Fig. 8.17 The surface plasmon resonance absorption peak position red-shifts with increase of the particle size

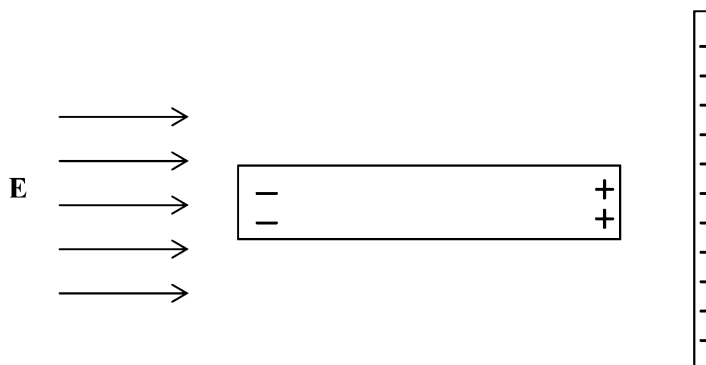


Fig. 8.18 Illustration of the change in the magnitude of the polarization due to the orientation of nanorods

For a given metal nanoparticle surface plasmon resonance energy is affected by the medium (dielectric constant) in which they are embedded. The surface plasmon resonance peak position also depends on the particle size (Fig. 8.17).

Usually small particles are dominated by absorption and large particles by scattering. Gold and silver nanoparticles also show effects on surface plasmon resonance due to shape. This is easily attributed to the differences in their polarization/depolarization effects. The orientation of nanoparticles with respect to the incident radiation also would get modified (see Fig. 8.18). For example a nanorod standing parallel or transverse to the incident beam would have different value of C_{ext} .

A nanorod also readily shows two surface plasmon resonance peaks as shown in Fig. 8.19.

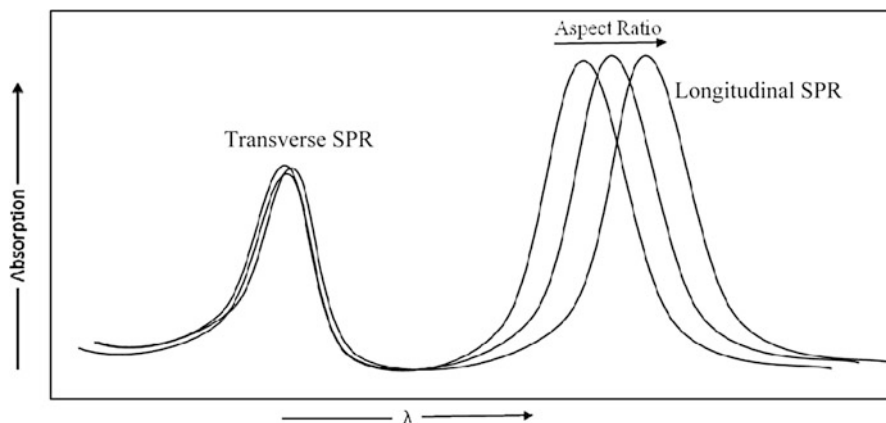


Fig. 8.19 Transverse and longitudinal surface plasmon peaks of nanorods. With increasing rod length the longitudinal plasmon resonance peak keeps on shifting to the longer wavelength but transverse resonance peak remains at the same position if the diameter of the rod remains unchanged

The two peaks appear due to the cross section (diameter) and length of the particle. If the diameter of the nanorod is kept constant and length changed, it is also possible to keep the resonance peak due to cross section at same energy and vary the longitudinal peak (due to length). The length/diameter or aspect ratio variation is often made in order to meet the needs of some applications.

Numerous other shapes like cubes, stars, tetra pods, flowers and so on are observed by tuning the synthetic conditions. They produce tunable surface plasmon resonance peaks.

Temperature is another parameter which can affect the surface plasmon resonance. With increase in temperature, volume expansion takes place shifting the electron density. This causes a red shift in the resonance peak position. The peak broadening also can occur as a result of increased scattering of electrons. The solvent temperature decreases its refractive index which shifts the resonance peak to blue side.

8.4.2 Surface Plasmon Polariton

Surface plasmon polaritons have become very important in nanotechnology because the scientists are quite hopeful of developing a branch called *nanophotonics* using them. Using some metallic nanostructures or thin films, unlike in localized surface plasmon resonance, it is possible to let the waves of surface plasmon polaritons to propagate along the interface between metal surface and a dielectric (vacuum, air, glass or any other dielectric medium) medium over large distances. This is because of the coupling of the electromagnetic radiation and the surface plasmons.

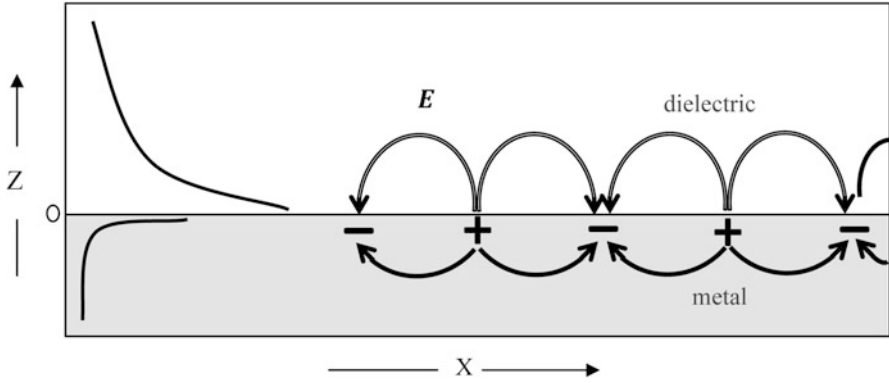


Fig. 8.20 Surface plasmon polariton. Metal shown in grey and white portion of the figure is dielectric

Coupling of a photon and surface plasmon is known as *surface plasmon polariton*. It should be remembered that polariton is a general term used to refer to any coupling between a photon and an elementary excitation (exciton, phonon and plasmon). Thus surface plasmon polaritons are confined and can travel over large distances of even few micrometres which makes them very useful to transfer the light from one end to the other end of the nanostructure. This becomes useful to design nano optical devices and circuits.

As shown in Fig. 8.20, there are two components of the surface plasmon polariton viz. longitudinal (along X-axis) and the transverse (along Z-axis).

The electric field inside the metal can be written as

$$E_m(x, z, t) = E_{m,0} e^{i(k_x x - k_z z - \omega t)} \quad (8.35)$$

and in the dielectric as

$$E_d(x, z, t) = E_{m,0} e^{i(k_x x - k_z z - \omega t)} \quad (8.36)$$

Note that m and d denote metal and dielectric. The electric field has both x and z components, although the amplitudes decay exponentially inside the metal as well as dielectric medium as shown on the left side of Fig. 8.20. Using the Maxwell's equations the dispersion relation can be shown as

$$k_x = \frac{\omega}{c} \left(\frac{\epsilon_m \epsilon_d}{\epsilon_m + \epsilon_d} \right)^{1/2} \quad (8.37)$$

with

$$k_x = \frac{\omega}{c} \sqrt{\epsilon}$$

ϵ_m and ϵ_d are dielectric constants of the metal and dielectric respectively.

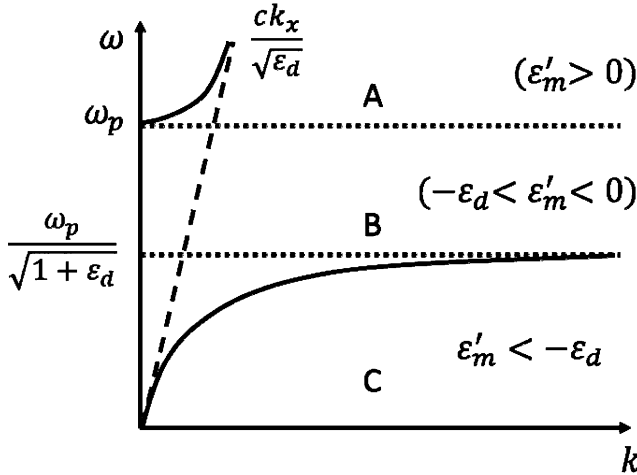


Fig. 8.21 Dispersion curve for the surface plasmon polariton

The dispersion behaviour of surface plasmon polariton can be understood from Fig. 8.21.

The dotted line represents the dispersion of light in the dielectric medium, air with $\epsilon_d = 1$ here. Obviously the dispersion is divided into three parts, A, B and C, depending upon the dielectric constant values of the metal. Here we consider the real part of the dielectric constant viz. ϵ' .

When $\omega > \omega_p$ (region A),

$$\epsilon_m \text{ and } \epsilon_m + \epsilon_d \text{ are } > 0 \quad (8.38)$$

The modes are radiative. In this region, k_x and k_z are both real.

In the region B, we have

$$\frac{\omega_p}{\sqrt{1 + \epsilon_d}} < \omega < \omega_p \quad (8.39)$$

Here $\epsilon_m < 0$ but $\epsilon_m + \epsilon_d > 0$ and k_x is imaginary but k_z is real and no propagating modes exist.

In the region C,

$$0 < \omega < \frac{\omega_p}{\sqrt{1 + \epsilon_d}} \quad (8.40)$$

The quantities ϵ_m and $\epsilon_m + \epsilon_d$ are < 0 and k_x is real but k_z is imaginary. When the values of k_x are small they approach the dotted line for dispersion of light in the dielectric.

In region C, we have bound modes and through the coupling of the electromagnetic waves and surface plasmons, surface plasmon polaritons can be propagated along the surface without radiative decay. This is what is expected in this branch of nanophotonics.

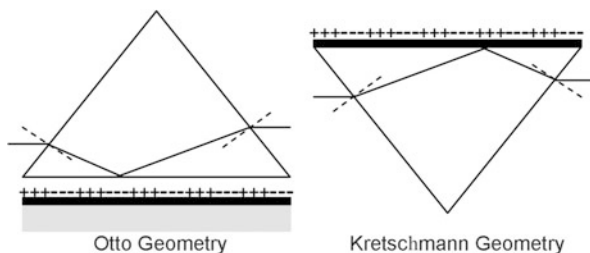


Fig. 8.22 The light rays passing through the prisms couple with the surface plasmons shown as positive–negative charge waves on the metal thin film (*dark black*). By keeping the analyte between the prism and the metal thin film, either in Otto or Kretschmann geometries, analytes can be studied by finding out the changes in the reflected rays

We shall not go into more details about surface plasmon polaritons here but just mention that surface plasmon modes and electromagnetic modes can be coupled also at the glass prism/dielectric and metal surface as illustrated in Fig. 8.22. Two different configurations can be set to investigate the surface plasmon polariton or to make its application as a sensor.

The light rays are internally reflected at glass prism and a dielectric between metal thin film, known as Otto configuration as in the left of Fig. 8.22 or by depositing the thin metal film directly on the prism. The silver or gold films are often used.

8.5 Nanomagnetism

Magnetism in bulk materials is due to magnetic moments on constituting atoms, ions or molecules. Presence of a magnetic field around a moving charge is well known. In an atom, spin and orbital motion of electrons and change of orbital motion in the presence of external applied magnetic field $\mathbf{H}$, gives rise to magnetic moment. Nuclear magnetic moment is usually very small and will be neglected in our discussion.

Total spin angular momentum due to spin of electrons in incomplete shells is given by

$$\mathbf{S} = \sum_i^N \mathbf{S}_i \quad (8.41)$$

Where $\mathbf{S}_i$ is the spin on the i th shell and the sum is over N shells. Corresponding dipole moment is given by

$$\boldsymbol{\mu}_s = -\frac{e}{m} \mathbf{S} \quad (8.42)$$

Total orbital angular momentum due to that of electrons in incomplete shells is given by

$$\mathbf{L} = \sum_i^N \mathbf{L}_i \quad (8.43)$$

Where $\mathbf{L}_i$ is the orbital momentum of i th shell and the sum is over all the N shells. The magnetic orbital angular momentum is given by

$$\mu_L = \frac{e}{2m} \mathbf{L} \quad (8.44)$$

Total angular momentum $\mathbf{J}$ is

$$\mathbf{J} = \mathbf{L} + \mathbf{S} \quad (8.45)$$

Interaction between $\mathbf{L}$ and $\mathbf{S}$ is known as spin-orbit interaction. The total magnetic moment of an atom is given by

$$\mu = \mu_L + \mu_S \quad (8.46)$$

It can be shown (see Fig. 8.23) that in free space magnetic moment on an atom is given by

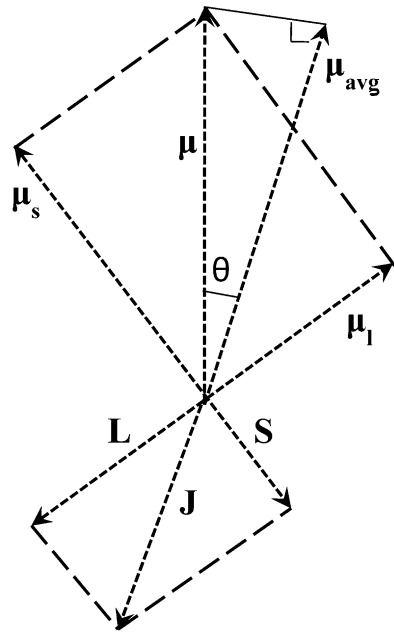


Fig. 8.23 Spin-orbit interaction

$$\mu_{\text{avg}} = g \left(-\frac{e}{2m} \right) \mathbf{J} \quad (8.47)$$

with

$$g = 1 + \frac{J(J+1) + S(S+1) - L(L+1)}{2J(J+1)} \quad (8.48)$$

In the presence of an applied magnetic field, angular momentum can have $2J+1$ orientations. One can determine S , L and J values if one knows the angular momentum of incomplete shells and number of electrons using Hund's rules as follows.

Hund's Rules

1. The maximum allowed value of S will be decided by Pauli's exclusion principle.
2. The maximum value of L is consistent with the value of S from L .
3. The value of J will be $|L - S|$ when the shell is less than half filled and $|L + S|$ when the shell is more than half filled.

The resultant magnetic moment of an atom could be defined as:

$$|\mu| = \mu_B g (J(J+1))^{1/2} \quad (8.49)$$

where μ_B is the Bohr magneton ($9.3 \times 10^{-24} \text{ Jm}^2/\text{Wb}$) and g is Lande factor having a value between 1 and 2.

8.5.1 Types of Magnetic Materials

Depending upon the response of a material to the external magnetic field, i.e. how much magnetization (M) is induced in the material by an external magnetic field most of the materials can be classified as diamagnetic, paramagnetic, ferromagnetic, antiferromagnetic and ferrimagnetic. The ratio of M to H is called the susceptibility:

$$\chi = \frac{M}{H} \quad (8.50)$$

$$\mu_0 = \frac{B}{H} \quad (8.51)$$

where B is magnetic induction and μ_0 is the permeability of the free space.

Permeability of the medium μ_m is given by

$$\mu = 1 + 4\pi\chi \quad (8.52)$$

We shall now discuss these materials in brief.

8.5.1.1 Diamagnetic Materials

Diamagnetism can be considered as the atomic manifestation of Lenz's law. When a magnetic field is applied to a conducting loop, a current is induced which opposes the change of flux. Electrons in atom are equivalent to conducting loop. In the applied magnetic field, electrons precess about the magnetic field direction. The precession is known as Larmor precession. It induces a magnetic moment

$$\mu = -\left(\frac{\mu_0 e^2}{4m}\right) r^2 H \quad (8.53)$$

where r is the radius of the electron orbit. There is no diamagnetic component of magnetic moment in the absence of an external magnetic field. However, as soon as an external magnetic field is present, diamagnetic component would always be present.

The magnetic susceptibility for N atoms having atomic number Z is given by

$$\chi = -\left(\frac{N\mu_0 e^2 Z}{6m}\right) a^2 \quad (8.54)$$

where a is the average radius of electrons in an atom.

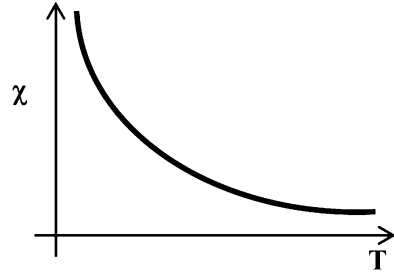
Thus diamagnetic materials have temperature independent negative susceptibility.

8.5.1.2 Paramagnetic Materials

Materials having some net magnetic moment present on atoms, molecules or ions even in the zero magnetic field are paramagnetic. In the zero magnetic field these magnetic moments would be randomly distributed and the material can show zero magnetization. In the presence of the magnetic field the magnetic moments try to align themselves in the direction of the magnetic field. Thermal energy kT tries to disturb this alignment. Diamagnetic component also arises in the opposite direction of the magnetic field but paramagnetic response can be larger resulting into a net positive response. Magnetic susceptibility is given by

$$\chi = \frac{Np^2\mu_B^2 Z}{3kT} = \frac{C}{T} \quad (8.55)$$

Fig. 8.24 Temperature variation of magnetic susceptibility in case of paramagnetic materials



where C is known as Curie constant and

$$p = g[J(J + 1)]^{1/2} \quad (8.56)$$

p is known as effective Bohr magneton.

Equation (8.55) is known as Curie law. It shows that the magnetic susceptibility in case of paramagnetic substances is inversely proportional to the temperature. Figure 8.24 illustrates schematically the temperature dependence of paramagnetic susceptibility.

8.5.1.3 Ferromagnetic Materials

Some magnetic materials have regions or domains in which large number of atoms or ions or molecules having permanent magnetic moments, like in paramagnetic materials, are aligned in a particular direction. Different domains are separated by what is known as domain walls. The domains themselves may be randomly oriented in the absence of a magnetic field. Such materials are said to have spontaneous magnetization and are known as ferromagnetic materials. This is due to the magnetic interaction, known as exchange field, amongst the magnetic moments of atoms. It was proposed by Weiss that the effect of mean magnetic field due to all other moments on a given moment of an atom is given by

$$\mathbf{H} = \lambda \mathbf{M} \quad (8.57)$$

$$\mathbf{H}_{\text{total}} = \mathbf{H} + \lambda \mathbf{M} \quad (8.58)$$

Modified Curie law, known as Curie–Weiss law for the susceptibility of ferromagnetic materials, is given by the equation

$$\chi = \frac{C}{(T - T_c)} \quad (8.59)$$

Fig. 8.25 Hysteresis loop of a ferromagnetic material

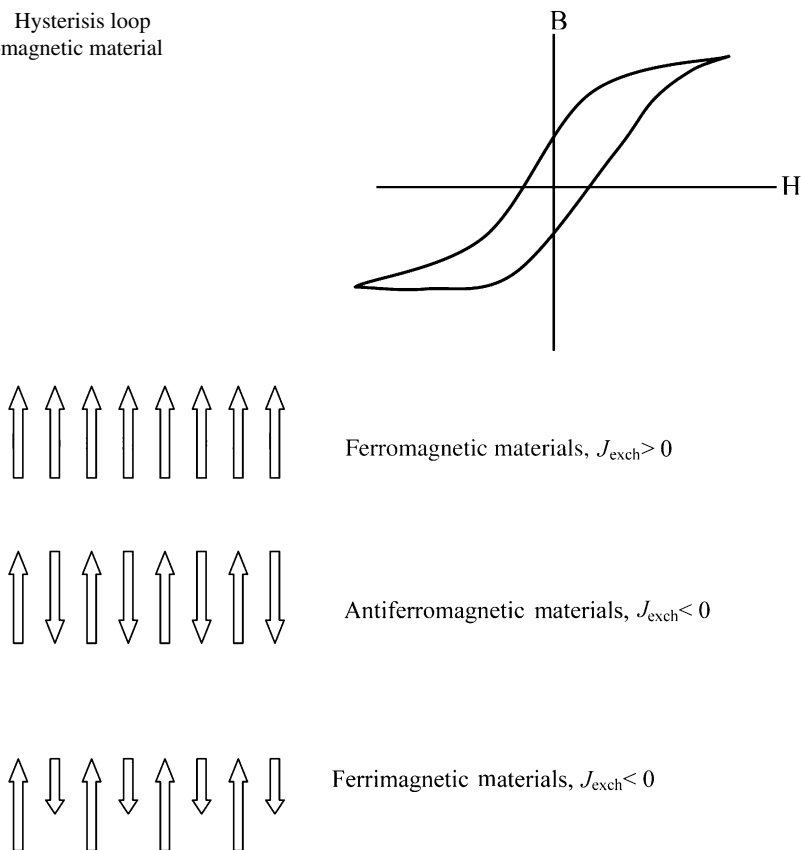


Fig. 8.26 Schematic of spin orientations on neighbouring atoms in case of ferromagnetic, antiferromagnetic and ferrimagnetic materials

where T_c is the Curie temperature. In ferromagnetic materials below the characteristic temperature T_c of the material, spontaneous magnetization exists in zero magnetic field. At temperature higher than T_c , the material behaves like a paramagnetic material. Ferromagnetic materials exhibit a hysteresis behaviour (see Fig. 8.25) of magnetization when an external magnetic field is applied.

Heisenberg explained the spontaneous magnetization by considering that the exchange energy U_{exch} which is due to the spin moments say $\mathbf{S}_i$ and $\mathbf{S}_j$ on i th and j th atoms is given by

$$U_{\text{exch}} = -2J_{\text{exch}} \mathbf{S}_i \cdot \mathbf{S}_j \quad (8.60)$$

where J_{exch} is known as the exchange integral.

When J_{exch} is negative, antiparallel configuration of spins is favoured, whereas for positive J_{exch} , parallel configuration is favoured, as illustrated in Fig. 8.26.

Fig. 8.27 A relation between J_{ex} and interatomic distance for transition elements; Here r_a is interatomic distance and r_{3d} is the radius of 3d shell

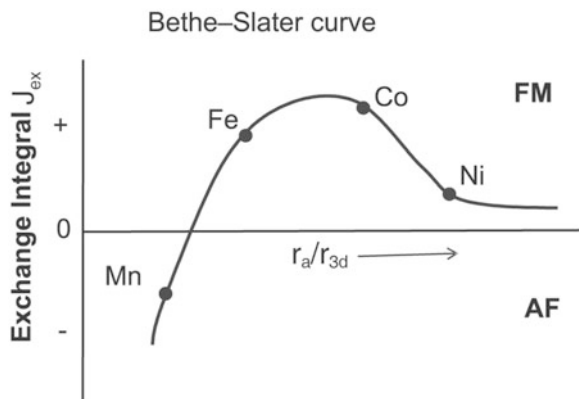
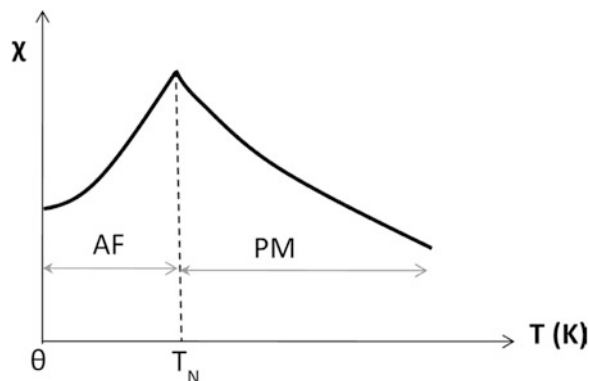


Fig. 8.28 Typical susceptibility curve in case of antiferromagnetic material



Slater showed that there exists a correlation between the interatomic distance and radius of incomplete filled d shell in case of transition elements, which is responsible for the ferromagnetism or antiferromagnetism, as shown in Fig. 8.27. It can be seen that this makes Fe, Co and Ni ($J_{\text{exch}} > 0$) ferromagnetic but Mn and Cr ($J_{\text{exch}} < 0$) antiferromagnetic materials.

8.5.1.4 Antiferromagnetic Materials

For antiferromagnetic materials the J_{exch} is negative and spins on neighbouring atoms are antiparallel to each other. Susceptibility in these materials is small but positive, as applied magnetic field tends to orient the spins in the direction of magnetic field overcoming the diamagnetism. Figure 8.28 illustrates schematically the susceptibility of antiferromagnetic materials in general.

Antiferromagnetic arrangement exists below a critical temperature known as Néel temperature T_N . Above T_N , susceptibility is paramagnetic and has the form

$$\chi_\theta = \frac{2\chi}{T + T_N} \quad (8.61)$$

8.5.1.5 Ferrimagnetic Materials

As illustrated in Fig. 8.26, there is also ferrimagnetism which is similar to antiferromagnetism but alternate spins are of smaller magnitudes than the rest. In general the ferromagnetic and antiferromagnetic materials are quite complex in their behaviour.

8.5.1.6 Nanomagnetic Materials

With this background of magnetic materials we now proceed to understand properties of some special types of nanomagnetic materials in which materials are reduced at least in one of the dimensions. Magnetic nanoparticles, assemblies of nanoparticles, magnetic nanowires, magnetic thin films or multilayer films and some metal oxide films, doped semiconductor particles or thin films have become the focus of attention due to their interesting magnetoresistive or magneto-optical properties they exhibit. The field of nanomagnetic materials is quite vast and still expanding. We shall discuss some of these materials.

8.5.1.7 Magnetic Nanoparticles

Ferromagnetic materials like Fe, Co, Ni, Fe_3O_4 , $\gamma\text{-Fe}_2\text{O}_3$ and many others have very interesting behaviour below a critical size, characteristic of each material. Bulk ferromagnetic materials have spontaneously magnetized domains. Formation of domains occurs in order to minimize the total magnetostatic energy of the system. However below the critical size, domain formation is not energetically favoured and material prefers to be single domain. In such a situation all the spins of atoms are oriented in one direction. Typically, the particles with a size below 100 nm are likely to be single domain. Coercivity H_c in such particles is given by

$$H_c = \frac{2K}{M_s} \left[1 - \frac{T}{T_B} \right] \quad (8.62)$$

where M_s is the saturation magnetization due to applied magnetic field, K – the anisotropy constant and T_B is the blocking temperature. Anisotropy (preferred

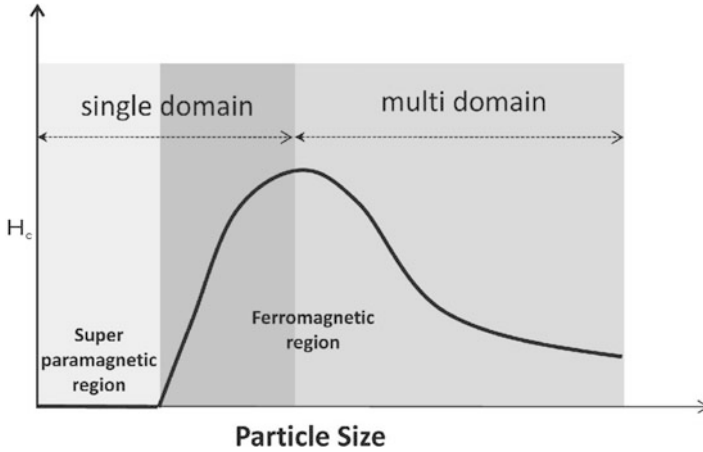
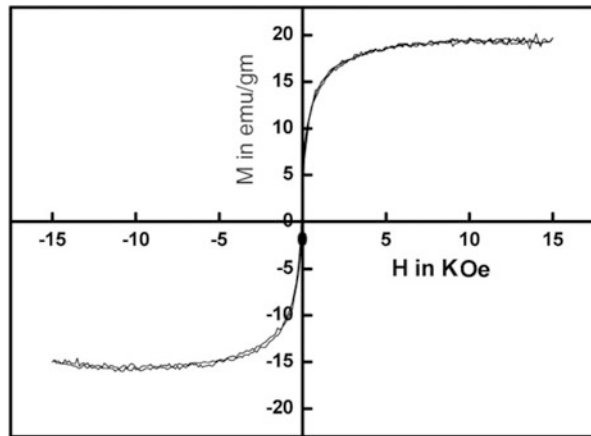


Fig. 8.29 A relation between particle coercivity and size

Fig. 8.30 Magnetization behaviour in case of a superparamagnetic particle



orientation of magnetic moments) can arise due to crystal structure, shape, stress and surface. Blocking temperature is the temperature, above which thermal energy is able to set the orientation of magnetic moments free. Below the blocking temperature, they are as if frozen. Figure 8.29 shows a common behaviour in materials of H_c as a size of particles.

Single domain particles of extremely small size which do not show coercivity or hysteresis (see Fig. 8.30) are known as superparamagnetic particles.

In superparamagnetic particles, spins are oriented in one direction and switch coherently in the opposite direction. The switching time is of the order of few nanoseconds and is given by

$$t = \tau_0 e^{-KV/kT} \quad (8.63)$$

where KV is the barrier for total spin orientation. Typical measurements of magnetization require 10–100 seconds. Thus magnetization (M) versus applied field (H) would exhibit the curve as shown in Fig. 8.30 with no coercive field. If magnetization M versus (H/T) are plotted then all the points for a superparamagnetic material lie on a single curve.

Small particles are characterized by large surface to volume ratio. Therefore surfaces and interfaces play an important role in their magnetic properties of nanostructures. At surfaces there is not only the symmetry breaking of the bulk crystal structure but there is a change in the coordination number as well as change in the lattice constant. Such effects can give rise to observation of ferromagnetic behaviour in materials which are not ferromagnetic in the bulk form. The coupling of magnetic nanoparticles spread at short distances also is an interesting area of current research.

8.5.1.8 Magnetic Oxide Materials

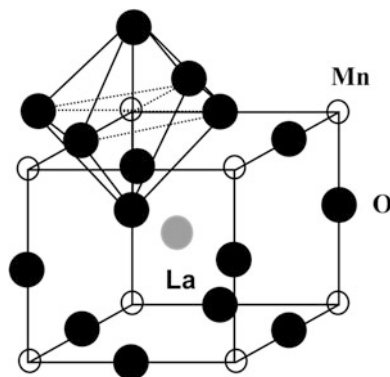
Magnetic multilayers are artificially obtained materials in which thin layers of magnetic layers separated by nonmagnetic layers are coupled and show Giant Magneto Resistance (GMR). In some metal oxides it has been observed that some planes are magnetically coupled to each other. This gives rise to huge change in the resistance when magnetic field is applied. The change in resistance is even larger than that in magnetic multilayers and is referred to as Colossal Magnetoresistance or CMR.

Colossal Magnetoresistance has been predominantly discovered in manganese-based perovskite oxides. In these crystals strong mutual coupling of spin and charges on atoms is observed. Hence not only high temperature superconductivity, but also new magnetoelectronic properties are increasingly discovered in materials with perovskite structures.

The 3d transition metal oxides, particularly the manganites are being exploited from the prospects of having improved device performance as compared to the GMR materials. These oxides display a diverse nature of properties such as paramagnetic to ferromagnetic transition accompanied by insulator to metal transition and realization of high magnetoresistance on application of comparatively low magnetic field. Historically magnetoresistance in manganites is known since 1950s. But recent interest is developed in doped manganites for exploring further high magnetoresistance value at room temperature at low magnetic field value to interplay of spin, charge and lattice distortion.

For example, undoped LaMnO_3 is an antiferromagnetic charge transfer insulator wherein Mn is in +3 valence state. As shown in Fig. 8.31 it has a perovskite structure. If any divalent atom (such as Ca, Sr, Ba, Pb) is doped at La site for instance, $\text{La}_{1-x}\text{Ca}_x\text{MnO}_3$, it converts Mn^{+3} to Mn^{+4} state in equal proportion of the doping concentration. This is equivalent to the hole doping in the system. The hole doping introduces a number of dramatic changes in electric and magnetic

Fig. 8.31 Perovskite structure of LaMnO_3



properties from the parent LaMnO_3 compound, such as insulator-metal transition, paramagnetic-ferromagnetic transition, charge ordered state, phase separation etc. In the light of these rich physical properties, hole doped LaMnO_3 has a potential for promising device applications such as magnetic sensors, magnetic valves, read head technology and bolometric application. Thin films of these materials are therefore a topic of great current interest.

8.6 Mechanical Properties of Nanomaterials

Mechanical properties of materials depend upon the composition and bonds between the atoms viz. covalent, ionic, metallic etc. As a result purest materials may be inherently weak or strong or brittle. Presence of impurities affects all these properties. Most of the materials have various impurities like C, O, N, P, S etc. present in them as well as point defects, grain boundaries, dislocations etc., which are responsible for the deviations of the properties expected from high purity and ordered materials.

When the size of materials is reduced to nanoscale, materials tend to be single crystals. However we need to consider different types of nanomaterials as schematically illustrated in Fig. 8.32 to specify the properties.

Indeed it is possible to determine various mechanical properties like elasticity, hardness, ductility etc. of different nanostructures. Techniques like bending measurement, velocity of sound measurement, nanoindentation etc. can be used. It should be, however, noted that measurements on single nanoparticles, rods, tubes etc. would inherently be difficult, though not impossible. However measurements on nanocrystalline solids, thin films etc. are possible using some conventional methods.

It has been shown in case of metallic nanocrystalline materials (as in Table 8.4) that elastic moduli reduce dramatically. For example in case of magnesium nanocrystalline material (grains ~ 12 nm size), Young's modulus was observed to be $3,900 \text{ N/mm}^2$ as against $4,100 \text{ N/mm}^2$ for polycrystalline (size $> 1 \mu\text{m}$)

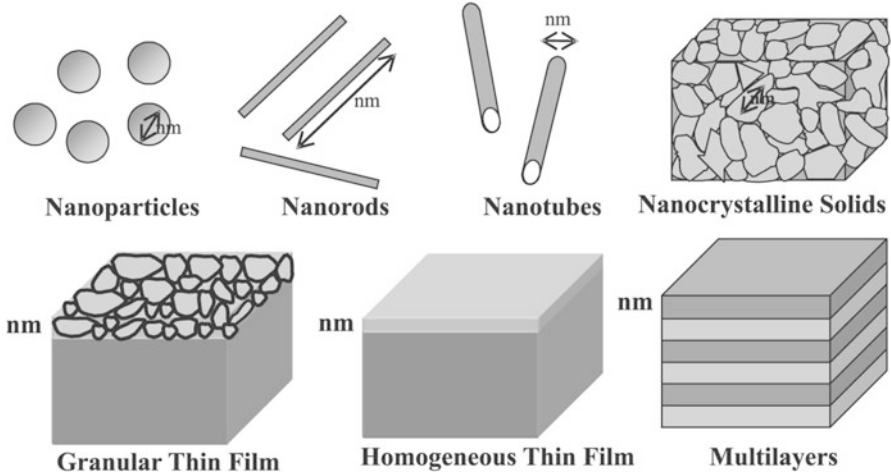


Fig. 8.32 Some common types of nanocrystalline materials

Table 8.4 Elastic properties of some common materials

Material	Density kg/m ³	Young's modulus 10 ⁹ N/m ²
Steel	7,860	200
Al	2,710	70
Glass	2,190	65
Concrete	2,320	25–35
Bone	1,900	9
Polystyrene	1,050	3
Diamond	3,510	1,035
CNT	2,600	1,280

magnesium. In case of palladium, nanocrystallites of ~ 8 nm size had Young's modulus $8,800 \text{ N/mm}^2$ as against $12,300 \text{ N/mm}^2$ for polycrystalline palladium.

Ceramic materials are often compacted and sintered using powder material. This also increases the hardness of materials. It has been shown that in case of TiO_2 nanoparticles (~ 12 nm size) produced in powder form much less temperature was required to densify and achieve the hardness comparable to usual polycrystalline material.

In fact density of nanocrystalline pellet is often low due to some pores left when powders are compressed to form pellets. Nanocrystalline pellets densities approach those of bulk polycrystalline materials as the sintering at high temperature progresses.

Plastic deformation in nanocrystalline materials strongly differs from that of polycrystalline bulk counterpart as shown in Fig. 8.33. It was shown in case of nickel that stress removal results in more effective recovery of the material as compared to corresponding polycrystalline material.

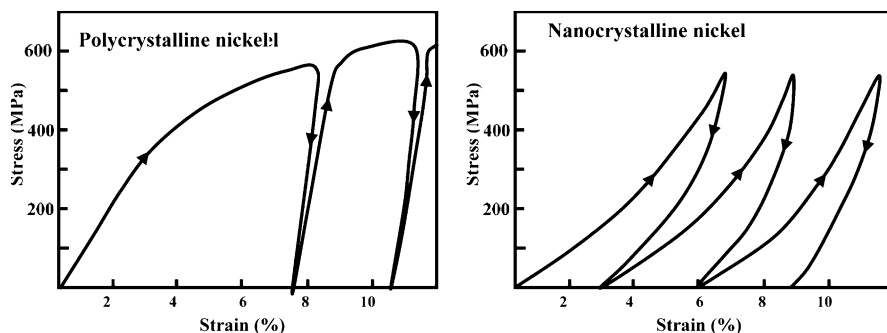


Fig. 8.33 Comparison of stress and strain relation in case of nanocrystalline and polycrystalline material

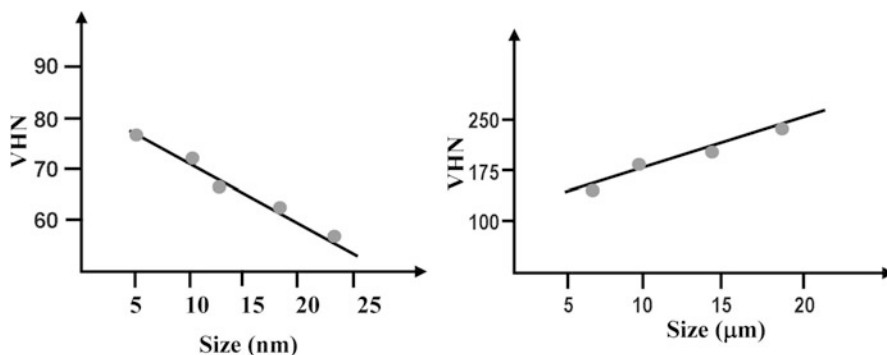


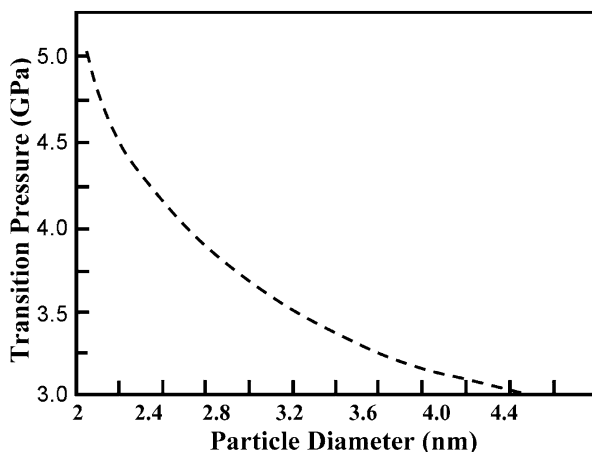
Fig. 8.34 Hardness variation for copper in nanometer and micrometre grain size range

Hardness of materials is also related to the grain size. As illustrated in Fig. 8.34 for copper, even in micrometre grain size range there is a linear dependence of hardness on particle size. It decreases with increase of grain size. However in nanometer size range, the hardness increases with increase of particle size linearly. Similar results are found in case of palladium nanoparticles and microparticles.

8.7 Structural Properties

Even though some nanomaterials with slightly large number of atoms (>50 – 60 atoms) may acquire bulk crystalline structure, it is found that the lattice parameters may not be the same as in the bulk material. For example it has been shown by rigorous analysis of X-ray diffraction patterns of ZnS that as small as 1.4 nm particles had liquid disorder. However larger crystals of ZnS indeed show same

Fig. 8.35 Structural deformation of CdSe nanoparticles



sphalerite (cubic) structure as in the bulk. It has been observed that there is a lattice contraction of $\sim 1\%$ for 1.4 nm ZnS nanoparticles. Other small particles also show up to $\sim 2\text{--}3\%$ lattice constant deviations compared to bulk crystalline materials.

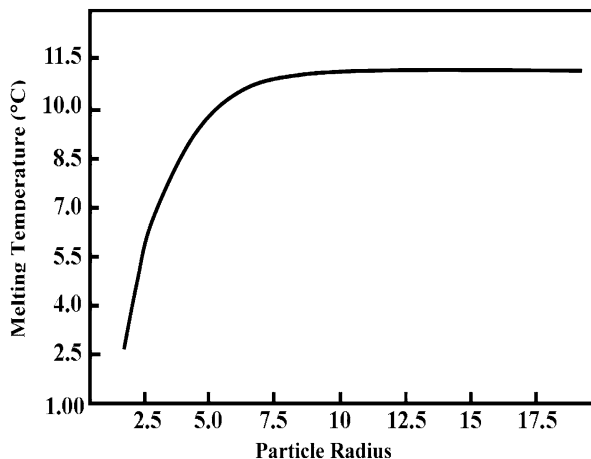
With increase in temperature the disordered structure of small particles of ZnS were found to transform to wurtzite (hexagonal) structure. Further the chemical capping, often used in the synthesis of nanoparticles, gets removed and particles tend to agglomerate or coalesce forming larger particles.

Effect of pressure on structural properties (using X-ray diffraction) has also been well investigated for some nanoparticles. It has been found that indeed the structural transformations do take place in case of nanoparticles with applied pressure. However the pressures required for this are larger for nanoparticles than for corresponding bulk material and depend upon the particle size, as illustrated in Fig. 8.35 for CdSe nanoparticles. Thus CdSe nanoparticles of 1–2.1 nm required 4.9 GPa to 3.6 GPa pressure to transform them from wurtzite to rock salt structure. Bulk CdSe needs just 2.0 GPa for the same transformation.

8.8 Melting of Nanoparticles

A variety of nanoparticles like Au, Ag, CdS etc. have been investigated for their thermal stability and melting. Melting begins at the surface. As the particle size decreases surface to bulk atoms ratio increases dramatically. In small particles or clusters the central atom may be considered as surrounded by 1st, 2nd, 3rd, ... compact shells of atoms. Number of atoms in shells is given as $10n^2 + 2$. Thus 1st shell would have 12 atoms, 2nd shell would have 42 atoms and so on. It can be easily seen that number of surface atoms is quite large in nanoparticles and surface to bulk atoms ratio goes on increasing with decreasing particle size (or shells). Large surface

Fig. 8.36 Variation of melting point with size of nanoparticles



is related to large surface energy. This energy can be lowered by melting. As shown in Fig. 8.36 melting temperature of gold nanoparticles of 3–4 nm size is reduced by $\sim 500^\circ\text{C}$ compared to bulk melting point.

Melting of nanoparticles is usually determined either by X-ray diffraction or electron diffraction. Heating increases the lattice parameter and at melting long range order is lost.

Further Reading

- M. Born, E. Wolf, *Principles of Optics* (Cambridge University Press, Cambridge, 2003)
W.D. Callister, *Materials Science and Engineering: An Introduction*, 4th edn. (Wiley, New York, 1997)
M. Fox, *Optical Properties of Solids* (Oxford University Press, Oxford, 2012)
H. Gleiter, Nanocrystalline materials. *Prog. Mater. Sci.* **33**, 227 (1989)
S. Shionoya, W.M. Yen (eds.), *Phosphor Handbook* (CRC Press, Boca Raton, 1999)

Chapter 9

Nanolithography

9.1 Introduction

We have seen in Chaps. 3, 4, and 5 a variety of methods to make nanoparticles (including spherical and other particle shapes) and thin films (or multilayers). These methods are popularly known as *bottom up* approach. In bottom up approach atoms and molecules are assembled so as to have nanomaterials of required size and shape by controlled deposition or reaction parameters. In the *top down* approach reverse is the case. Atoms and molecules are removed from a bulk material or sometimes thin films so as to obtain desired nanostructure.

Conventional *lithography* is a top down approach. The word lithography has its origin in the Greek word ‘litho’ which means stone. Lithography, therefore, literally means carving a stone or writing on a stone. It is used now to mean a process in which a sample is patterned by removing some part of it or sometimes even organizing some material on a suitable substrate. Lithography is very intensively used in electronics industry so as to obtain integrated circuits (IC) or very large scale integration (VLSI) on small piece of semiconductor substrate often called a ‘chip’.

Immediately after the discovery of solid state transistor by Bardeen, Brattin and Schockley in 1947, the quest for making smaller and smaller electronic components began. First transistor was fabricated in germanium but soon it was found that silicon was a better material and was used commercially. Texas Instruments, U.S.A. marketed first transistor. Scientists found that speed of a device, system or an instrument finally depends on how fast a transistor can switch on and off. Speed of a computer also depends upon transistors. All electronic devices including cell phones, ATMs, etc. need fast transistors. It was further realized that smaller the device (or transistor), faster is the switching. Advantages of making smaller devices are manifold. They require smaller amount of material, space and consume less power for their performance making the resulting product cheaper. Initially the solid state transistors were assembled together along with different components like capacitors and resistors and wired to fabricate desired circuits.

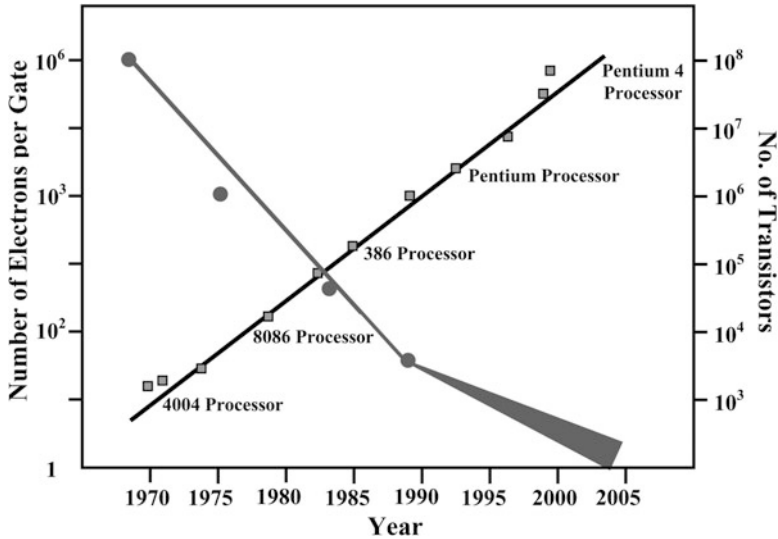


Fig. 9.1 Moore's law

It was proposed by a British engineer G.W.A. Dummer that an entire circuit should be directly made on a silicon substrate instead of wiring together the different components. Robert Noyce at Fairchild Semiconductor and Jack Kilby at Texas Instruments, U.S.A., invented in 1958 what we know today as integrated circuit (IC) which had complete circuit fabricated on a single silicon wafer. This made a great revolution in electronics and Jack Kilby was awarded the Nobel Prize in the year 2000 for this development.

Around 1959, James Moor predicted that there would be a reduction in the transistor size with time. Ever since its depiction the trend in miniaturization of electronic devices has faithfully followed what is known as Moor's law (see Fig. 9.1). It implies that every 18 months the reduction in the size doubles or the number of transistors on a chip doubles or the processing power of computers doubles. However following this law the devices have reached now a lowest size of ~ 100 nm and deviation from the law has begun. It is not only increasingly difficult to achieve smaller and smaller sizes less than 100 nm but also difficult to retain linear nature of the graph that Moor had predicted. Below a size of ~ 100 nm we know that besides the 'surface effect', materials also have size-dependent properties. Therefore nanodevices using active or passive nanocomponents cannot be expected to behave like those of large (micrometre) size devices and components. Interestingly this very size-dependent nature can be used to obtain some novel devices, which were not imagined earlier. For example *single electron transistor* (SET) is a completely new device due to unique properties of quantum dots. *Magnetic Spin Valve* and *Magnetic Tunnel Junction* (MTJ) using nanomaterials are some other high speed devices which are the products of nanotechnology.

Thus a new era in electronics has begun with new devices which have much larger memory for computers, consume low power, are compact and faster in their operations. All this is possible but new devices will be economically viable provided one can pattern small devices perhaps using *nanolithography*.

Over the last 3–4 decades different lithography techniques like optical lithography, X-ray lithography and electron beam lithography have been developed. They depend upon using photons or particle radiations for carving the materials. The lithography technique involves transfer of some pre-designed geometrical pattern (called *master* or *mask*) on a semiconductor (like silicon) or directly patterning (often known as *writing*) using suitable radiation. Mask is usually prepared by creating radiation opaque and transparent regions on glass or some other material. Pre-designed patterns can be transferred on a substrate much faster as compared to direct writing. Direct writing being a slower process is overall expensive.

Common principle in most of the lithography techniques is to expose a material sensitive to either electromagnetic radiation or to particles in some regions. Such a radiation-sensitive material is known as *resist*. The selection of area to be exposed to radiation is made using a *mask*, which is transparent in some regions and opaque in the other regions. This causes selective exposure of the resist, making it weaker or stronger compared to unexposed material depending upon the type of the resist being used. By removing the exposed or unexposed material in suitable chemicals or plasma, desired pattern is obtained. This may be done in a number of steps depending upon the pattern and materials involved (Box 9.1).

Box 9.1: History of Lithography

Roots of present lithography can be found in the art of printing as well as photography. Lithography or writing/carving on stone has a very long history. It is an art invented by Alois Senefelder from Austria in 1798. He covered the surface of a highly porous stone with a mixture of a gum and water. Only the stone part absorbed the greasy solution. He then dipped the stone in ink made of oil, wax, soap and lampblack. The ink could stick only to the greasy part of the stone. When ink coated stone was pressed on a piece of paper, an impression was made.

It was soon realized that complicated figures, designs or patterns can be easily transferred number of times by this process of printing. By 1848 it was possible to print ~10,000 copies in an hour.

Jules Chéret (1836–1933) in Paris made some artistic posters using lithography. He was awarded the Legion of Honor for creating a new branch in art.

When ICs were to be fabricated on an extremely large scale so as to satisfy the needs of large electronics market, it was necessary to adopt a process of lithography, which can make multiple copies of a pattern in a short time.

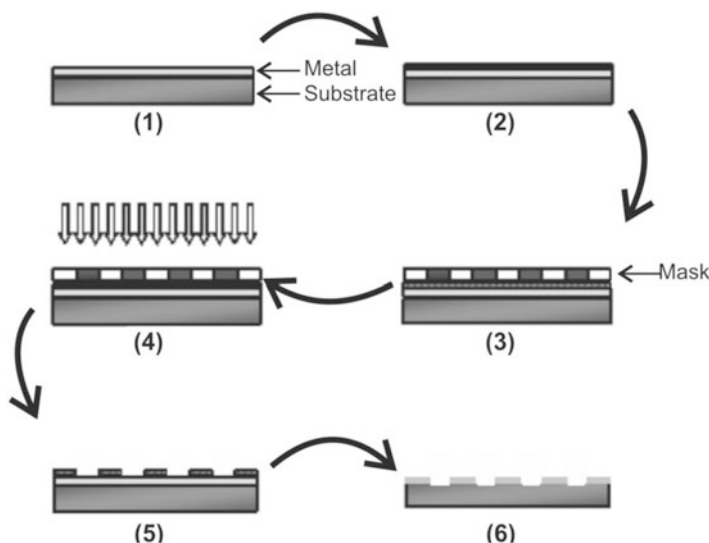


Fig. 9.2 Photolithography process steps: (1) surface is coated with metal, (2) coating of photoresist on the substrate, (3) mask placed over upper layer, (4) exposed UV radiation, (5) resist development and stripping, and (6) etching to get final pattern

Figure 9.2 depicts schematically various steps involved in photolithography, to transfer a pattern on some semiconductor surface. A thin film coating of a metal (like chromium) is deposited on a suitable substrate (for example glass or silicon). A *positive or a negative photoresist*, usually some polymer, is coated on metal thin film. Positive photoresist material has the property that when exposed to the appropriate radiation it degrades or some chemical bonds are broken. Negative resist on the other hand is a material, which hardens (crosslinks) on exposure to a radiation. A mask is placed between the resist-coated substrate and the source of light. By using a suitable chemical (developer) the weakened portion is removed (or image is developed). Remaining unexposed part also can be removed by appropriate chemical treatment. The remaining material can be dissolved in one step and the hardened material in another step. Depending upon the radiation used like visible light, X-rays, electrons or ions, the lithography name is tagged with it.

After the development of Scanning Tunneling Microscopes (STM) around 1982 and other Scanning Probe Microscopes (SPM) thereafter, it was realized that they can be used to carry out lithography in nanometer range. Using SPM probe or fine tip of SPM it is possible to directly write on the material.

In recent years some replication techniques also have emerged which are quite inexpensive and allow patterning necessary for some exotic purposes like ‘lab-on-chip’ or quick diagnostics.

Here we shall outline essentials of lithography (patterning) using photons, electrons, scanning probes and replication methods.

9.2 Lithography Using Photons (UV–VIS, Lasers and X-Rays)

It is possible to use visible, ultraviolet, extreme ultraviolet (EUV) or X-rays to perform lithography. Wherever possible, lasers are used. Highest resolution of the generated features ultimately depends upon the wavelength of radiation used and interaction of radiation with matter as well as mask and optical elements used. Smaller the wavelength used smaller can be the feature size which is limited by diffraction limit, $\sim \lambda/2$. Depth of focus depends upon the penetration of incident radiation. For the lithography using electromagnetic radiation, optical elements and masks have to be used for various purposes. In the visible range ($\sim 700\text{--}400\text{ nm}$), glass lenses and masks can be used. In the UV range, fused silica or calcium fluoride lenses are used.

There are three methods (see Fig. 9.3) viz. ‘proximity’, ‘contact’ and ‘projection’ which can be used to pattern a substrate.

As the name suggests, in ‘proximity’ method, mask is held close to the photoresist coated metallized substrate, whereas in ‘contact’ method the mask is in contact with photoresist. In both proximity and contact methods, a parallel beam of light falls on the mask, which transmits the radiation through some windows and blocks through opaque parts. Although better resolution is achieved with contact method as compared to the proximity method, in contact method the mask gets damaged faster compared to the proximity method. In case of projection method, a focussed beam is scanned through the mask, which allows good resolution to be achieved along with the reduced damage of the mask. However scanning is a slow process and also requires scanning mechanism, adding to the cost.

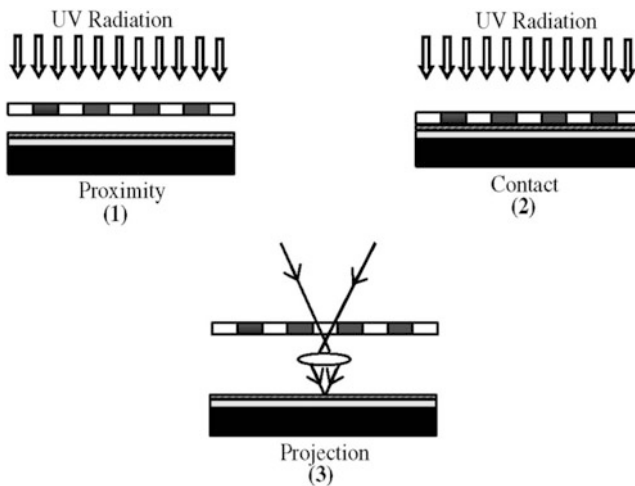


Fig. 9.3 (1) Mask is close to the photoresist, (2) mask is in contact with the resist and (3) focussed beam is scanned through the mask

9.2.1 *Lithography Using UV Light and Laser Beams*

Using monochromatic light in visible to UV light features as small as 1–1.5 μm size can be routinely obtained. Often g-line (436 nm) from the mercury line is used as a source of radiation. Laser beam of KrF (248 nm) or ArF (193 nm) also are employed reaching ~ 150 nm as the smallest feature size. However to obtain feature size below ~ 100 nm using photons is a difficult task, unless one uses *near field optics* (see Chap. 7) based on scanning probe technique.

9.2.2 *Use of X-rays in Lithography*

Smaller features are possible to obtain by employing X-rays also. However it is difficult to make suitable masks for X-ray lithography. X-rays in the 0.1–5 nm range are used with appropriate metal masks in proximate geometry. Absorption of X-rays in materials not only depends upon the thickness of the material but is also complicated by the presence of absorption edges. Depending upon the wavelength of X-rays used, metals of suitable elements are chosen. Metal masks are fabricated in such a way that through thin portions they are transmitted and absorbed in thicker regions. Gold masks are often used. The masks themselves are made using electron beam lithography discussed in the next section.

9.3 *Lithography Using Particle Beams*

We know that all the moving particles have associated wavelength λ known as de Broglie wavelength given by

$$\lambda = h/mv \quad (9.1)$$

where h is Planck's constant, m —the mass and v is the velocity of the particle. All kinds of particles can in principle be used but to achieve high resolution λ . It should be as small as possible. Thus large mass and large velocity of particle makes it possible to get adequate resolution. In fact it is possible using neutral atoms, ions or electrons to bring down the particle-associated wavelength to any desired value, even as small as even 0.1 nm. However ultimate resolution depends upon the interaction of incident particles with resist material. Under certain conditions features as small as 2 nm have been patterned. Due to various reasons like they can be easily generated, accelerated and focussed, electrons are preferred for lithography purpose and often used.

9.3.1 Electron Beam Lithography

Figure 9.4 shows schematically electron beam lithography set up. It is very similar to a scanning electron microscope (SEM) and requires vacuum ($\sim 10^{-2}$ – 10^{-4} Pa). Sometimes SEM is modified in order to use it as a lithography set up. Electron beam lithography is a direct writing method i.e. no mask is required to generate a pattern. Rather, patterns required for other lithography processes like soft lithography (discussed in Sect. 9.5) can be generated using electron beam lithography.

Electrons with high energy (usually larger than ~ 5 keV) are incident on the photoresist. Here also positive or negative photoresists can be used. Common positive resists are polymethylmethacrylate (PMMA) and polybutane-1-sulphone (PBS). Negative resist often used in electron beam lithography is polyglycidylmethacrylate coethylacrylate (COP). Developers used are methylisobutylketone (MIBK) and isopropylalcohol (IPA) in 1:1 ratio. A focussed electron beam in electron beam lithography is used in two modes viz. 'vector scan' or 'raster scan'. In vector scan the electron beam 'writes' on some specified region. After one region is completed the X-Y scanning stage on which the substrate to be patterned is mounted, moves. During its movement electron beam is put off. Then a new region is selected and

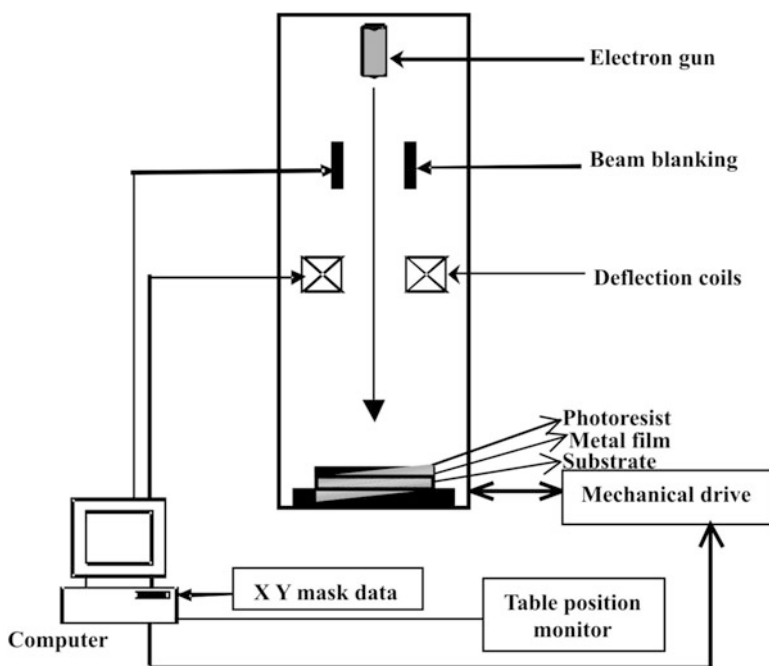


Fig. 9.4 Electron beam lithography set up

‘written’ with the beam. This is continued until whole pattern is generated. In ‘raster scan’ the beam is rastered or moved continuously over a small area, line by line. The X-Y stage of the sample moves at right angles to the beam. The beam is turned off or turned on depending upon the pattern. Although very high resolution (~ 50 nm) is routinely possible using this lithography, due to scanning mode it is rather slow. For example if the optical lithography can generate 40 patterns with $1\text{ }\mu\text{m}$ resolution in 1 h, only five similar patterns would be generated with electron beam. However larger layer depth is an added attraction of electron beam lithography as compared to optical beam lithography.

9.3.2 Ion Beam Lithography

Very small size features ($\sim 5\text{--}10$ nm) having large depth can be written using high-energy ion beams. Major advantages of using ion beams is that resists are more sensitive to ions as compared to electrons and have low scattering in the resist as well as from the substrate. Commonly used ions are He^+ , Ga^+ etc. with energy in the 100–300 keV range.

9.3.3 Neutral Beam Lithography

Neutral atoms like argon or cesium have been allowed to impinge on substrates to be patterned through the mask. Such beams cause less damage to the masks. Self assembled monolayers on gold substrates have been often patterned using neutral beams.

In fact any deposition of neutral atoms (physical vapour deposition, molecular beam epitaxy etc.) through the mask can be considered as lithography of this type.

9.3.4 Nano Sphere Lithography

It is a very simple but useful method of obtaining desired patterns of controlled sizes and shapes with regular spacing. This is achieved usually by using self assembly of polymer or silica colloidal particles on appropriate chemically treated substrate. The formation of silica colloidal particles by chemical method can be found in Chap. 12. Size of the particles can be controlled by managing the reaction time as well as dilution of the chemicals. The self assembly of the particles is achieved by allowing the liquid containing silica particles (or polymeric colloids) to slowly evaporate. Once the film dries, the metal or desired materials are evaporated on the films which can get deposited in the spaces between the spheres. After deposition the spherical

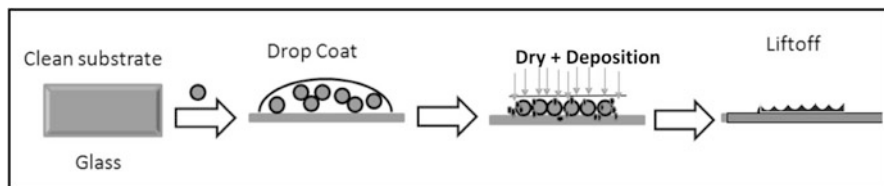


Fig. 9.5 Sequence of patterning steps using nano sphere lithography methods

particles can be simply removed in an ultrasonic treatment or using suitable solvent which only removes self assembled colloids. Figure 9.5 illustrates schematic of nano sphere lithography steps.

Spherical colloids can produce triangular patterns but other shapes of colloidal particles may be employed to produce other shapes.

9.4 Scanning Probe Lithography

With the development of scanning tunnelling microscope (STM) and atomic force microscope (AFM) there began a new chapter in the history of lithography. Microscopes were discussed in Chap. 7. STM and other similar microscopes using sharp tips or probes for imaging can be used for lithography purpose. While using STM, some scientists noticed that repeated scanning on some areas gave different images due to movement of atoms. Systematic observations have evolved a branch known as Scanning Probe Lithography (SPL). One major advantage of SPL is that like optical lithography it also can be carried out in air. There are different ways in which SPL can be carried out viz. mechanical scratching or movement, optical and electrical.

9.4.1 Mechanical Methods

In mechanical lithography, there are different modes like scratching, pick-up and pick-down or dip pen lithography as described briefly below.

There are a large number of experiments in which pits or lines can be produced using either STM tip or AFM tip on the surface of bulk material or surface of a thin film. Often diamond tips can be used.

Formation of pits or lines by scratching is like ploughing, in which scratched material is piled up around the indented region (as shown in Fig. 9.6). Variety of materials like nickel, gold, copper, polymers, Langmuir Blodgett films, and high temperature superconductors are possible to scratch. Pits as small as 30 nm in diameter and 10 nm in depth are possible to make.

Fig. 9.6 Schematic of mechanical scratching by a microscope probe tip

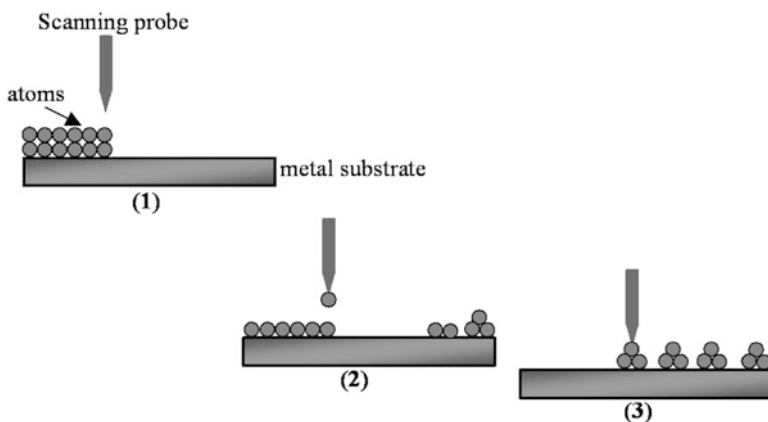
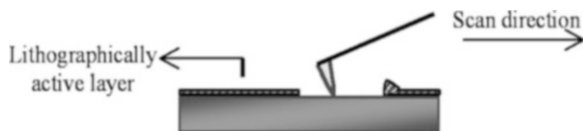


Fig. 9.7 (1) Atoms on a substrate, (2) being picked up one by one and (3) arranged in desired pattern

It was often found in STM/AFM that some loosely adsorbed surface atoms or molecules could be moved with scanning probe. Systematic work by scientists at IBM made it possible for them to pile up xenon atoms on a metal substrate and write a letter pattern “IBM” for the first time. Later they could also organize different metal atoms on metal substrates producing some beautiful patterns. Schematic of the process is shown in Fig. 9.7.

Some scientists moved 30 nm GaAs particles on a GaAs substrate. Letter patterns as high as 50 nm in height were made using AFM tip by some scientists. Now the technique is used to fabricate some circuits.

9.4.2 Dip Pen Lithography

This method is very similar to pick up and pick down method. The method bears a similarity to writing on a piece of paper with ink. That is why the name dip pen lithography is given. An AFM tip is used as a pen and molecules are used as ink (see Fig. 9.8). Appropriate molecules picked up by the tip from the source of molecules can be transported and transferred at desired place on the substrate. Letters with line thickness as small as 15 nm and distance 5 nm have been written. Overwriting and erasing capability of dip pen lithography is quite a unique feature.

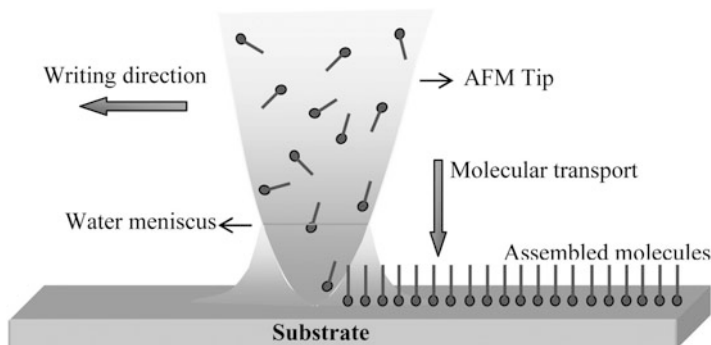
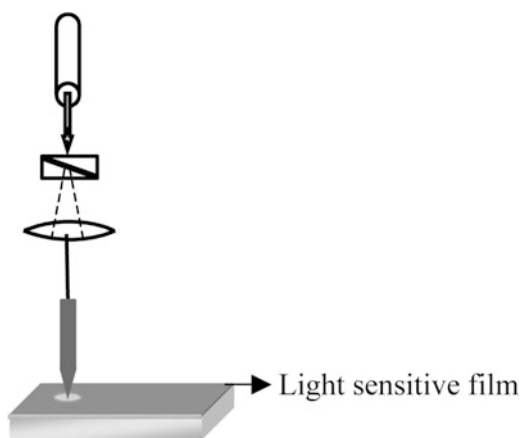


Fig. 9.8 Principle of dip pen lithography

Fig. 9.9 Lithography using SNOM probe



9.4.3 Optical Scanning Probe Lithography

As discussed earlier, very high resolution $\sim 20\text{--}50\text{ nm}$ is possible, overcoming the diffraction limit, even with visible light using Scanning Near-Field Optical Microscope (SNOM). This is attributed to near-field component of electromagnetic radiation. In SNOM (see Fig. 9.9), a fine spot of visible light emerging through an aperture, scans on the surface at a distance of $\sim \lambda/50$, where λ is the wavelength of light used for scanning. By placing the aperture close to the photoresist coated substrate, it is possible to obtain as small as $\sim 50\text{ nm}$ size features routinely.

9.4.4 Thermo-Mechanical Lithography

It is also possible to use an AFM tip along with a laser beam and carry out nanolithography (see Fig. 9.10). While the AFM tip is in contact with coating like

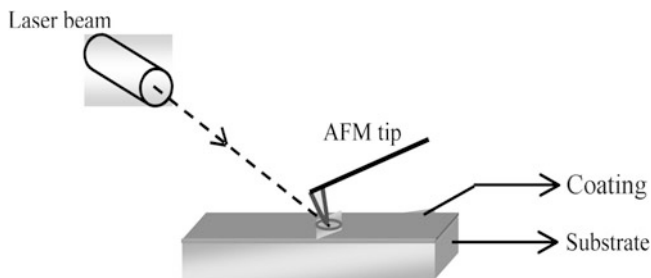


Fig. 9.10 Thermo-mechanical lithography

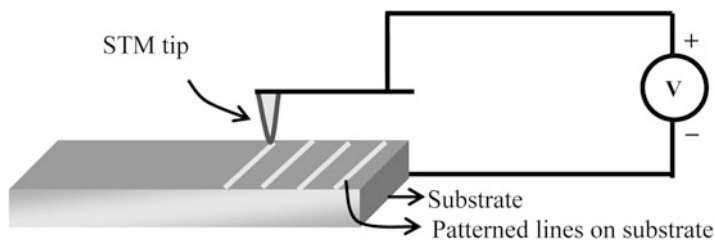


Fig. 9.11 Electrical SPL

PMMA, laser beam strikes the same point of the coating. This heats the film locally enabling the tip to penetrate in the material and make a pit. This thermo-mechanical method is capable of producing resolution as high as ~ 30 nm.

9.4.5 Electrical Scanning Probe Lithography

In this method, as illustrated in Fig. 9.11, a voltage is applied between the STM tip and the sample. Above some critical voltage, if large current flows between the tip and the sample, an irreversible change can occur in sample surface. Variety of bulk solid and thin films surfaces have been patterned using this method. In silicon or modified silicon surfaces, ~ 30 – 60 nm wide and ~ 5 – 10 nm deep lines have been engraved.

9.5 Soft Lithography

There are some inexpensive, non-conventional techniques that have been developed, useful for patterning non-planar and non-routine, inorganic, organic, or biological samples. By conventional techniques, we mean techniques like photon lithography

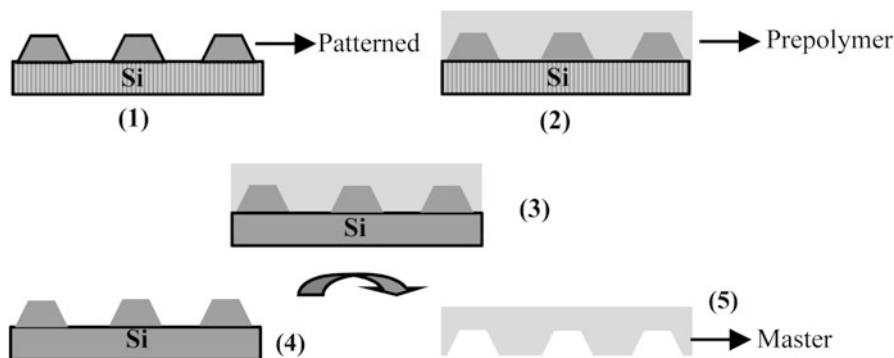


Fig. 9.12 Process of making master: (1) A desired pattern is obtained on silicon substrate using suitable lithography. (2) Prepolymer is poured on silicon-patterned substrate. (3) After proper heat treatment (4) master is removed easily from silicon pattern and is (5) ready for use

or particle lithography. To achieve resolutions better than 100 nm feature size, X-rays, electrons or ions need to be used in conventional lithography. Due to various technical problems related to radiation and matter interactions, expensive and sometimes dedicated instruments using high energy beams are required. These are indeed used in microelectronics industry. However for other purposes, soft lithography is a useful alternative to obtain resolution better than ~ 100 nm at low cost. Moreover, the method is applicable from few nm to few μm size features. The name soft lithography is used to mean the techniques using materials like polymers, organic materials or self assembled films.

In general soft lithography technique involves fabrication of a patterned master, molding of master and making replicas. A master is usually made using X-ray or electron beam lithography. It is supposed to be quite rigid. A mold is usually made using a polymer like polydimethylsiloxane (PDMS), epoxide, polyurethane etc. PDMS is most common amongst the polymers used for molding due to its attractive properties like thermal stability ($\sim 150^\circ\text{C}$), optical transparency, flexibility ($\sim 160\%$ elongation), capability of cross-linking using IR or UV radiation etc. However during molding some distortions can take place and adequate control has to be practiced to achieve reproducible and required results. Figure 9.12 illustrates the method of making a master.

Replication of patterns can be done by different ways, the most common are as described below.

9.5.1 Microcontact Printing (μCP)

A PDMS stamp is dipped in an alkanethiol solution (see Fig. 9.13) and pressed against the metallized (Au, Ag, Cu) substrate. Those parts of substrate which come

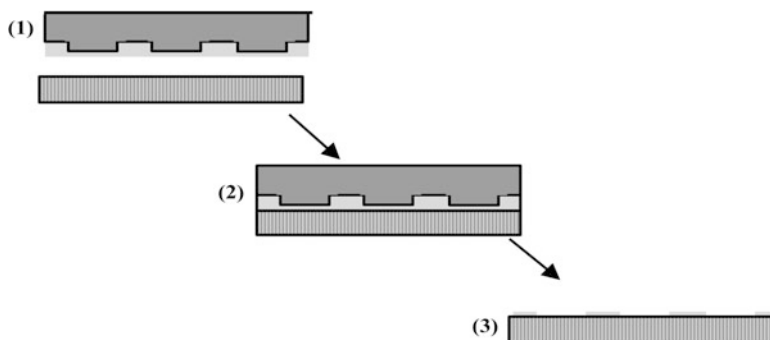


Fig. 9.13 Microcontact printing: (1) Stamp is inked with alkanethiol, (2) stamp is placed on the substrate and (3) stamp is cured and removed

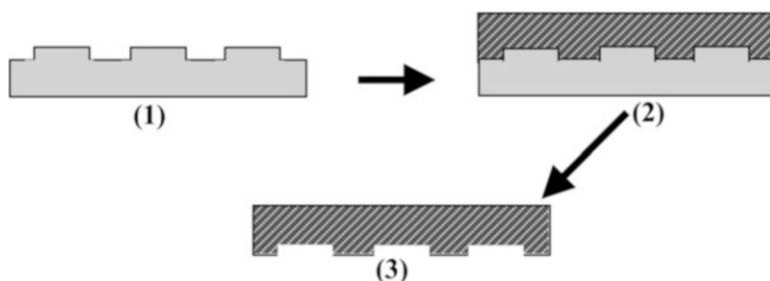


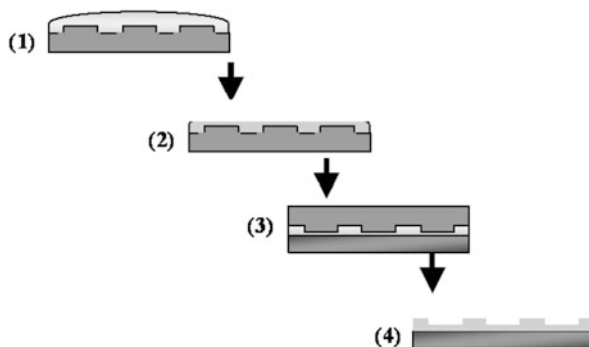
Fig. 9.14 Replica molding: (1) PDMS master, (2) pour a prepolymer on master, and (3) after heat treatment, pattern (its complementary part) is transferred into polymer

in contact with the PDMS receive layers of alkanethiol. The monolayers do not spread on the substrate. Further, these self assembled monolayers can be used as resists for selective etching or deposition. The printing being simultaneous, it is a fast method.

9.5.2 Replica Molding (REM)

In this method, a PDMS master or stamp is used to replicate a number of copies. For example a solution of polyurathene (see Fig. 9.14) is poured in PDMS and cured using UV light or thermal treatment so that polyurathene becomes solid. PDMS can be easily removed so that a pattern opposite to that is produced in polyurathene. By applying small pressure on PDMS, it is possible to further reduce the size of the features smaller than in the original pattern. Nanostructures ~ 30 nm have been achieved using this method.

Fig. 9.15 Microtransfer molding: (1) Prepolymer poured on the stamp, (2) Excess solution is removed using nitrogen blow, (3) Stamp is pressed on the substrate and cured for one hour, and (4) Stamp removed carefully



9.5.3 Microtransfer Molding (μ TM)

As shown in Fig. 9.15, a pre-formed polymer is poured in PDMS stamp. Excess polymer is removed by blowing nitrogen gas on it and the stamp is pressed against a substrate. Using thermal treatment polymer is imprinted on the substrate and mould is removed.

9.5.4 Micromolding in Capillaries (MIMIC)

In this technique (see Fig. 9.16), a PDMS stamp is placed on a substrate to be patterned. A low viscosity polymer is then placed in contact with PDMS. The liquid flows into channels of PDMS by capillary action. After the thermal treatment of curing with UV radiation the polymer gets solidified. PDMS stamp is then removed to obtain patterned substrate.

9.5.5 Solvent-Assisted Micromolding (SAMIM)

A PDMS stamp coated with a solvent is pressed against the substrate coated with a polymer film (see Fig. 9.17). Solvent softens the polymer surface in contact. PDMS can be removed after the solvent has evaporated. PDMS stamp itself is not affected by the solvent. Volatile and substrate dissolving solvents, but not PDMS stamp dissolving, need to be used. Often 'Novalac' coating is given to the substrates. Polymethylmethacrylate (PMMA), cellulose acetate, polyvinyl chloride etc. are used as polymers.

Although soft lithography techniques are fast, economically viable and in principle capable of producing sub-nanometer patterns, the mechanical stability of such stamps is often a problem. Keeping good contact between the substrate and the

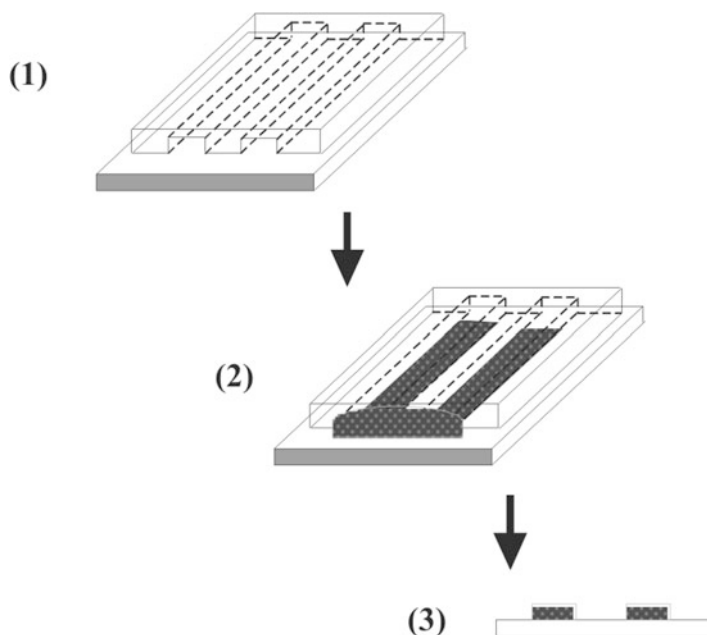


Fig. 9.16 Micromolding in capillaries: (1) Put the stamp over substrate, (2) pour the solution from open channel and cure it, and (3) stamp is removed and the pattern remains on substrate

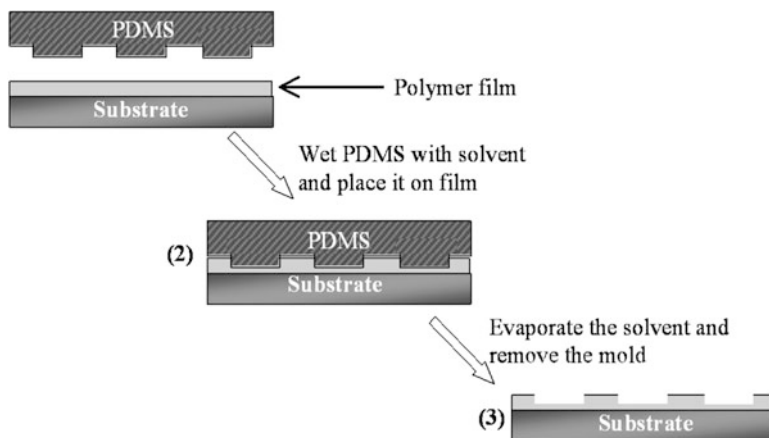


Fig. 9.17 Solvent-assisted micromolding: (1) Prepare a thin film on substrate, (2) put the solvent coated stamp over substrate and (3) mold is removed after evaporating the solvent

mold also is often problematic. Various further treatments of PDMS surface, use of shadow masks and proper choice of chemicals (solvents or organic molecules) etc. are necessary to achieve very high resolution and quality patterns. The field is yet under development and no unique method is so far available.

Further Reading

C.Y. Chang, S.M. Sze, *VLSI Technology* (McGraw Hill, New York, 1999)

S.A. Gangal, S.K. Kulkarni. *Physics Education*, April–June 2002, p 1–9

W.H. Moreau, *Semiconductor Lithography: Principles, Practice and Materials* (Plenum Press, New York, 1985)

Y. Xia, J.A. Rogers, K.E. Paul, M.G. Whitesides, *Chem. Rev.* **99**, 1823–1848 (1999)

X. Younan, G.M. Whitesides, *Angew. Chem. Int. Ed.* **37**, 550–575 (1998)

Chapter 10

Nanoelectronics

10.1 Introduction

In Chap. 9 (also see Fig. 9.1), we saw that the sizes of the electronic devices like transistors are shrinking accompanied by the increase of their density per unit area. With increased device density, in general, the cost of the product not only goes down but the device performance also improves. The performance is limited by the increased heat generation which in turn would restrict the size of the device that should be made. Further, as we go on reducing the size the physics of low dimensional materials does not remain the same as for the bulk according to the discussions in the previous chapters. In fact this gives rise to new nanodevices which would be discussed in this chapter just to understand the basic principle behind their peculiar behaviour which is not seen in the microdevices. Although very interesting, we shall not discuss the nanodevices in which carbon nanotubes or graphene are being used to obtain more efficient devices. Here we are restricting only to certain phenomena rather than special nanomaterials in the devices.

After the discovery of solid state transistor by Bardeen, Brattin and Schockley, revolutionary changes occurred in the field of electronics. Bulky, power hungry, expensive vacuum tubes based equipment slowly got replaced with tiny, less power consuming, light weight devices. This made it possible to produce portable, space saving, compact equipment. Today one can have personal wristwatches, calculators, portable computers (laptop), mobiles, stereos, videos etc. due to advances in electronics and related areas. Satellites, space missions and internets are unimaginable without solid state electronics. It was realized as early as 1958 that one could go on shrinking not only the sizes of individual device but fabricate large circuits on a single ‘chip’ as an ‘IC’ or integrated circuit (Chap. 9). Extrapolation of these ideas and development in device fabrication techniques like lithography, not only made it possible to fabricate Very Large Scale Integration (VLSI) of electronic devices and circuits but also faithfully produced large quantities of them commercially.

This makes it possible for the manufacturers to warrant their products for uniform performance.

In fact as early as in 1960, Moore predicted a trend in electronics device shrinkage which is popularly known as Moore's law which was discussed in Chap. 9. It can be noticed that after 2000 A.D. there has occurred a deviation from the Moore's law. This is quite easy to understand. One can go on reducing the size with the limit of an atom, but there is certainly a limit to the size below which properties of materials are not independent of size. We know it now that this is where the 'nanoscience' and 'nanotechnology' take over microelectronics.

Next revolution is expected in computers with what is known as nonvolatile memory by which we shall not lose any data being stored on a computer if there is a sudden electricity failure or we forget to save the entries. We may also have what is being researched presently as quantum computers which will be much more powerful than the existing computers. Such computers will use the fruits of nanotechnology.

The flat panel television or computer monitors are products of nanotechnology. Even the coatings used on screens of TV or monitors can be of nanoparticles, which have better properties in terms of colour quality and resolution than micro particle coatings.

Here we shall discuss a few peculiar phenomena/nanodevices achieved due to reduced dimensionality (confinement effects) as well as realization through the developments in the lithography techniques to fabricate nanodevices and chips.

10.2 Coulomb Blockade

Materials are often classified as metals, semiconductors and insulators, according to their ability to let current flow through them. Conductivity is defined in terms of the properties of electrons (their number, effective mass, scattering etc.) in the solids and is given by

$$\sigma = \frac{Ne^2\tau}{m^*} \quad (10.1)$$

where σ is electrical conductivity, N – number of electrons per cm^3 , e – electron charge, τ – relaxation time and m^* is effective mass of electron.

Resistivity is the inverse of conductivity. Metals are characterized by very low resistivity ($\sim 10^{-6} \Omega\cdot\text{cm}$). Semiconductors have medium resistivity (few $\Omega\cdot\text{cm}$) and insulators have large resistance ($> 10^3 \Omega\cdot\text{cm}$). The resistivity (or conductivity) in solids can be measured in principle by connecting electrically conducting wires to solid material of known geometry, applying a voltage difference across it and measuring the current flowing through it (Fig. 10.1) (Box 10.1).

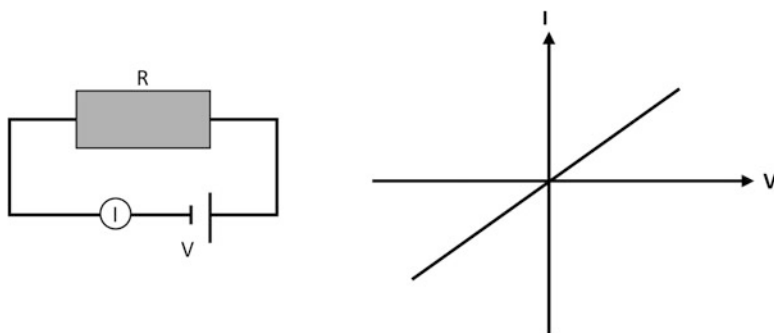


Fig. 10.1 Resistance measurement of a typical metal resistor

Box 10.1: Can Nanoparticles Be Considered as Metals?

Measurement of current through materials having dimensions in nanometer range is a difficult task. Particularly, measurements on single nanoparticles are difficult. Apart from the experimental difficulties one needs to also understand the meaning of metallic nature of particles at reduced dimensions. It is important to know, as the particle size reduces, does the material become semiconductor or insulator and at what size it deviates from being a metal. Is there any other characteristic of materials which can allow us to know when the metal stops being a metal as the particle size reduces. Indeed if one can try to carefully look at the elements in the periodic table, it can be seen that all the atoms have characteristic ionization potential or energy. The elements which form metals are characterized by the low ionization energy. The elements which form semiconductors have ionization energy larger than metals but smaller than that for the insulators. This is depicted in Fig. 10.2. Ionization energy is defined as the energy required to remove an electron from the atom/molecule/cluster/solid to vacuum level.

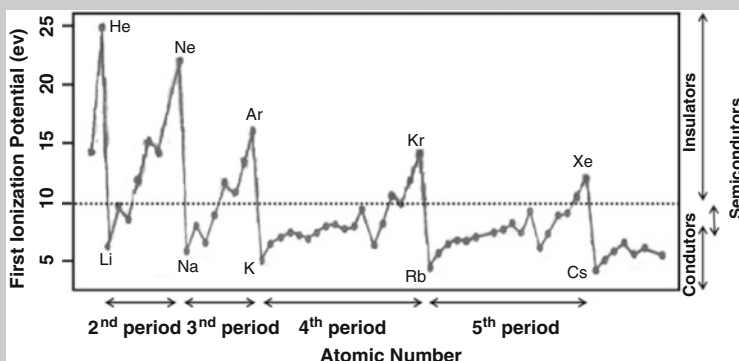


Fig. 10.2 Elements in the periodic table plotted according to their increasing atomic number and their first ionization potential

Fig. 10.3 Metal cluster or semiconductor quantum dot placed between the electrodes in order to measure the electron current flowing through the circuit

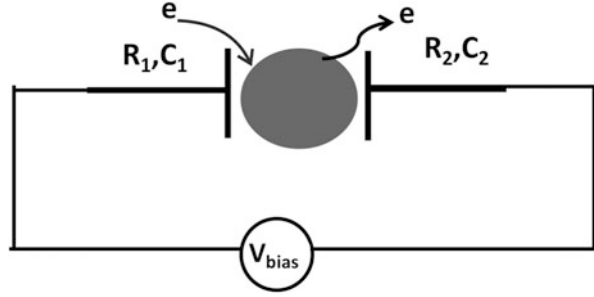
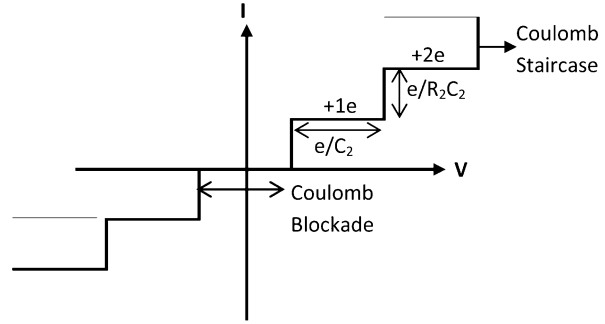


Fig. 10.4 Coulomb blockade and staircase for a quantum dot



If we reduce now the dimensions of metal piece (or introduce a semiconductor nanoparticles or quantum dot) to ~ 100 nm or less and wish to measure its conductivity, then it is useful to put metal electrodes (capacitors) on either side so that direct contact between electrodes and metal particle is avoided (Fig. 10.3). This enables to deduce the correct details of electronic structure. There appears then a region around zero voltage for which there is no current flow (Fig. 10.4). This phenomenon is known as *Coulomb blockade*. This can be understood as follows.

With the arrangement as in Fig. 10.3, there will be step-like current flow as shown in Fig. 10.4.

The electrostatic energy E (charging energy) of a parallel plate capacitor having capacitance ' C ' is given by

$$E = \frac{e^2}{2C} \quad (10.2)$$

For small value of the capacitance and low thermal motion of electrons ($kT \ll e^2/2C$) the charging energy E will be significant. The small metal island connected to electron source and drain by tunnel barriers can be charged in such a way that only a single electron is transferred to it when voltage $\pm e/2C$ is applied. Below this voltage electron cannot be transferred (Fig. 10.5). Therefore the region of *no current* of low bias voltage is known as *Coulomb Blockade region*. Repeated tunnelling of single electrons produces what is known as *Coulomb Staircase*.

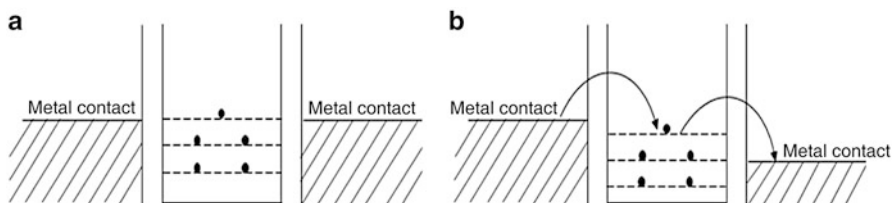


Fig. 10.5 (a) Current cannot flow; (b) Current flows due to the convenient alignment of the Fermi level positions in the contacts and the metal cluster or quantum dot

There are many examples now in which phenomenon of Coulomb Staircase has been demonstrated using quantum dots or metal islands.

The Coulomb blockade can also be very well understood from Fig. 10.5. When the Fermi levels on both the sides are at the same level, no current flows but the moment one of the electrodes as shown in the figure receives higher potential with respect to the quantum dot, the current can flow between the metal electrodes and the quantum dot. Similarly if the cluster is at higher potential compared to the electrode then the electrons from the cluster tunnel towards the electrode.

The single electron transistor is based on the phenomenon of Coulomb blockade.

10.3 Single Electron Transistor (SET)

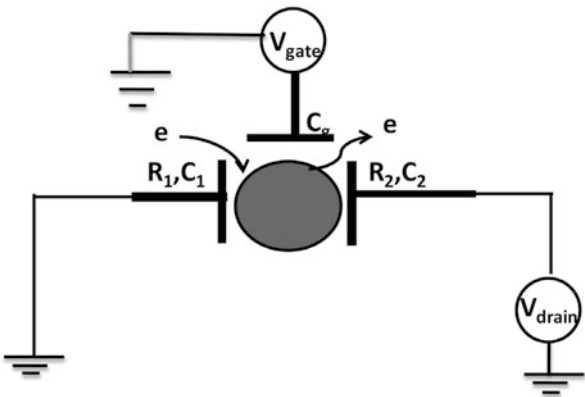
Just as in a usual or classical micro size transistor (see Box 10.2) there are three main components of a single electron transistor viz. source, drain and gate. As illustrated in Fig. 10.6, the quantum dot is placed between the source and the drain. The gate controls the raising or lowering of the electron levels in the quantum dot through the gate capacitor by adding or lowering the number of electrons in the quantum dot. As the phenomenon of tunnelling and, hence, Coulomb blockade are the main phenomenon involved, the electrons are controlled precisely one by one and the name to the transistor is given as *Single Electron Transistor*. Different SETs are fabricated based on mainly the ‘quantum dot’ being used. The requirement is that there has to be nanostructure with discrete energy levels as in ‘particle in a box’ which takes place of the quantum dot so that a single electron control is achieved.

Box 10.2: Diodes and Transistors

We had seen earlier briefly that semiconductor materials can be doped with various elements which change their optical and electrical properties by introducing some electronic states in the energy gap between the valence and the conduction band. Even though in nanostructure of semiconductors the

(continued)

Fig. 10.6 Schematic diagram of a single electron transistor with three terminals

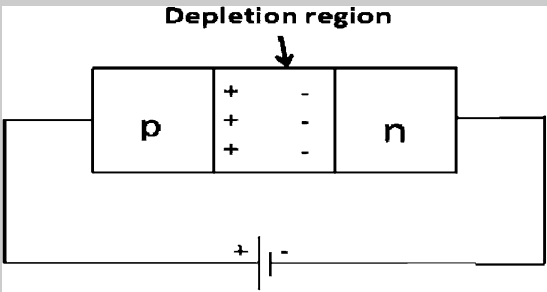


Box 10.2 (continued)

band gaps or overall electronic structure may change with dimensionality and shape, the dopants are still powerful means of altering their properties. The dopants may introduce extra electrons making the semiconductor ‘n’ type or introduce excess holes making it a ‘p’ type semiconductor.

The simplest and the earliest semiconductor device is a two-terminal diode in which, in a semiconductor material, under controlled diffusion of dopants a p-n junction is formed. Thus in the diode one side is ‘p’ type and other is ‘n’ type. This makes diode an asymmetric device. At the junction a depletion region is formed, as illustrated in Fig. 10.7.

Fig. 10.7 A p-n diode



This helps to make the diode forward biased (p connected to the positive terminal and n connected to the negative terminal) or reverse biased. By biasing through an external power source (battery), we are able to control the Fermi level and thereby the nonlinear electron or charge flow in the

(continued)

Box 10.2 (continued)

circuit. Diode is an asymmetric device and is able to allow the current in one direction but to block the current in the reverse direction. This can be seen from Fig. 10.8. Note that there can occur a breakdown known as Zener breakdown on the application of the large negative voltage.

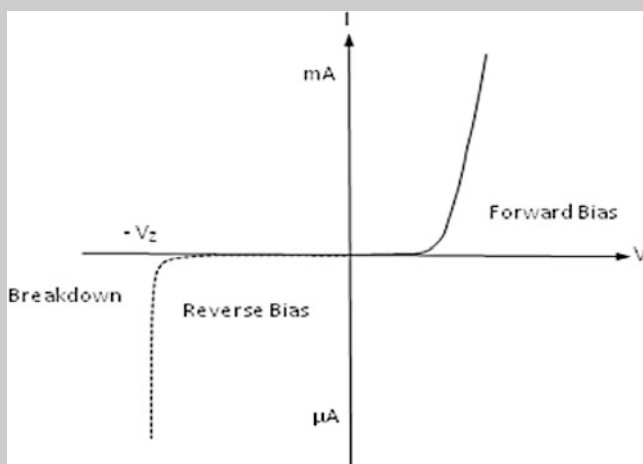


Fig. 10.8 Current flow with applied voltage in a p-n diode

Without going into further details about the diodes it can be just mentioned that the diodes are classified as Schockly diode, Schottky diode, Zener diode, photo diode, light emitting diode, Gunn diode, laser diode etc.

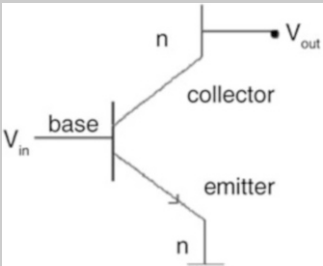
The next and most widely used semiconductor device is the three-terminal semiconductor transistor. It can be used to not only amplify the signal but also as a switching device. Historically, Julius Edgar Lilerfeld had patented in 1925 a field effect transistor (FET) in Canada and in 1926 and 1928 in USA. This was followed by another patent in 1934 by Oskar Heil in Germany. However John Bardeen, William Schockly and Walter Brattain at Bell labs made gold point contacts and observed in 1947 that they were able to get the large output signal compared to the input signal. They were given the Nobel prize in 1956. The device was referred to as a transistor by John R. Pierce because it can be looked upon as 'transfer resistor'. The increase of signal is nothing but 'gain'.

There are mainly two types of transistors viz. bipolar transistors and field effect transistors. The terminals of a bipolar transistor are called emitter, collector and base. In Fig. 10.9 an n-p-n bipolar transistor is schematically illustrated.

(continued)

Box 10.2 (continued)

Fig. 10.9 Schematic illustration of a three-terminal bipolar semiconductor transistor



The field effect transistor (Fig. 10.10) has source gate and drain on the terminals. Current between source and drain is controlled by primarily the voltage at the gate.

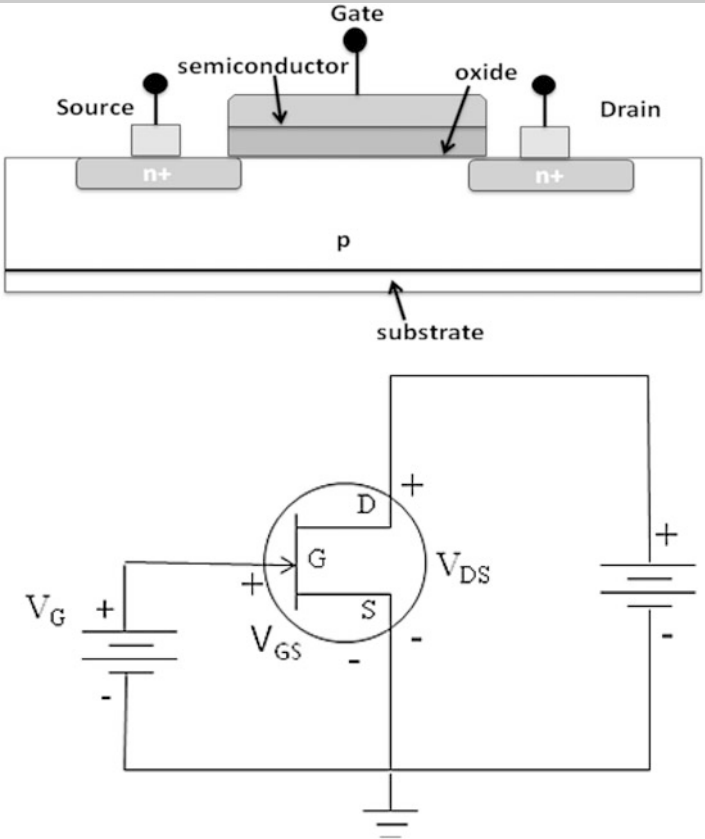


Fig. 10.10 Schematic illustration of a Field Effect Transistor

(continued)

Box 10.2 (continued)

The low operating voltage, hence power saving, has made transistors popular. Transistors also have high efficiency; they are reliable for repetitive operations and extremely long life. They are not affected by mechanical vibrations and render themselves in many applications. There is a huge variety of transistors such as Schottky transistor, bipolar junction transistor (BJT), insulated gate bipolar transistor (IGBT), avalanche transistor, Metal Oxide Field Effect Transistor (MOSFET) and so on. In recent years Carbon Nanotubes also have been used in Field Effect Transistor (CNFET). Transistors using organic semiconductors (OFET) also have been fabricated.

The transistors are lithographically patterned along with other components such as diodes, capacitors and interconnects to obtain densely packed integrated circuits (ICs).

10.4 Spintronics

The electronic devices with typical dimensions of few nanometers in either of three directions display not just the miniaturization but unique properties not known over last 5–6 decades since the beginning of solid state devices. Single Electron Transistor (SET), spin valves, and Magnetic Tunnel Junctions (MTJ) are conceptually new devices based on nanotechnology. Such devices are fast, compact, relatively cheap and finding their way to market. Spin valve type devices are already being used in personal computers to ‘read’ disk which have enabled to increase data storage capacity of hard disks. Interestingly, spin valve and MTJ are based on a concept which itself is growing into an area in itself known as spintronics or spin-based electronics or magnetoelectronics. Some potential spintronics materials are given in Box 10.3. It is well understood that an electron (or hole) has both charge and spin. However electronics has so far used only the charge property of electron (or hole) and spin has been neglected. It has been now realized in recent years that if spin of an electron (or hole) is taken into account, properly fabricated devices would lead to some superior devices. Using an external magnetic field, spin transport can be controlled. Advantage with spin is that it cannot be easily destroyed by scattering from collisions with other charges, impurities or defects. Many spin-based devices like Spin-FET, Spin-LED, Spin-RTD, optical switches with THz frequency, modulators, encoders, decoders, and q-bits for quantum computers are on the hot list of scientists and the technologists. We consider here devices based on Giant Magneto Resistance (GMR), spin valve, Magnetic Tunnel Junction (MTJ) and Spin Field Effect Transistor (SFET).

Box 10.3: Spintronics Materials

Some of the materials which have a potential in spintronics are summarized below.

1. II-VI semiconductors doped with transition metal ions (also known as Diluted Metal Semiconductors or DMS for short)
2. III-V semiconductors doped with transition metal ions (DMS)
3. Metal oxides with large band gap and doped with transition metals. For example TiO_2 , SnO_2 doped with cobalt
4. Eu chalcogenides
5. Heusler alloys like NiMnSb , Mn_2CoGe
6. Ferromagnetic metal oxides like CrO_2 , Fe_3O_4
7. $\text{Mn}_{11}\text{Ge}_8$
8. Lanthanum doped CaB_6

10.4.1 Giant Magneto Resistance

Giant Magneto Resistance (GMR) can be realized in the multilayers. Deposition of one kind of material over the other (sputtering, e-beam evaporation or electrochemical deposition are commonly employed techniques) of a few nanometer thickness, and repeating it several times gives rise to a multilayer. Multilayers are artificially created man-made materials. As can be seen from Fig. 10.11, such multilayers are characterized by the presence of a large number of interfaces. The properties of

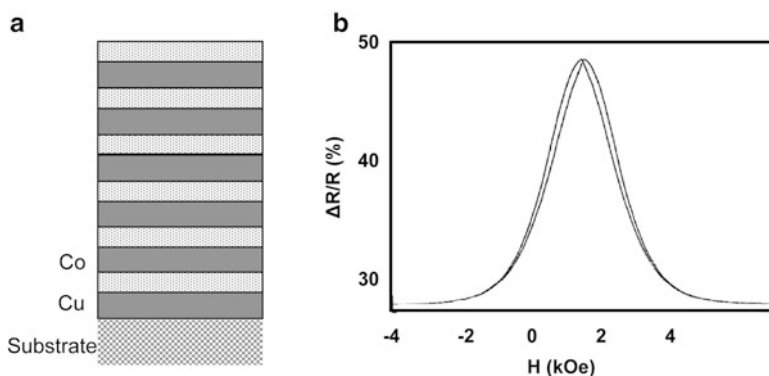


Fig. 10.11 (a) Schematic of a Co/Cu multilayer structure, (b) Magnetoresistance of the Co/Cu multilayer at room temperature

multilayers are, therefore, governed not only by the parent materials but also by their surface and interface properties.

Multilayers can be of metals, semiconductors, insulators, organic materials or combinations of these. Here we are interested in magnetic multilayers, though other types of multilayers also are interesting and have large applications.

Magnetic multilayers in which ferromagnetic layers of materials few nanometers thick are separated by comparable thickness metallic layers attracted the attention around 1988 when French scientist Albert Fert and German scientist Peter Grünberg observed that such ferromagnetic layers can be ferromagnetically or antiferromagnetically coupled. This gives rise to a magnetoresistivity which depends upon the orientation of the magnetic layers. Magnetoresistance (MR) is the relative change in electrical resistance of a material on the application of magnetic field.

It is usually defined as:

$$MR(\%) = \frac{R(0) - R(H)}{R(H)} \times 100 \quad (10.3)$$

where $R(0)$ is the resistance of the material when no external magnetic field is applied and $R(H)$ is the resistance of the material when external field of H is applied.

The change in the resistivity can be quite large and is known as Giant Magneto Resistance (GMR). Except few cases of anisotropic magnetoresistance, magnetoresistance in materials is usually quite low (less than even 0.5 %). However GMR can be as high as even 50–60 %. This is very effective in observing small changes in the magnetic field and useful as a read device of the magnetically stored data.

Obviously, the effect has found huge application in today's computers and can be considered as first direct application of nanotechnology. No wonder that both Fert and Grünberg received in 2007 the Nobel prize in physics for discovery of GMR.

In Fig. 10.11, Co/Cu multilayers are shown in which cobalt ferromagnetic multilayers are separated by copper layers. The corresponding magnetoresistance behaviour is shown in Fig. 10.11b. It can be seen that with the application of the magnetic field the resistance of the sample changes dramatically. This behaviour can be understood as follows.

Figure 10.12 schematically shows the mechanism due to which giant magnetoresistance occurs in magnetic multilayers. When magnetic layers are polarized in a particular direction, the carrier electrons with parallel spin pass through the material easily and resistance in such a case would be low for the flow of electrons. Thus alternate magnetic layers by the application of magnetic field would exhibit low resistance.

On the other hand if the alignment of magnetic layers is antiparallel to each other like in antiferromagnetic material, electrons, with spin parallel or antiparallel, are bound to see the layers of opposite spin and perceive the resistance to flow. Thus initially antiferromagnetically coupled magnetic layers would offer larger resistance as in (b) as compared to those that are ferromagnetically aligned as in (a). Depending

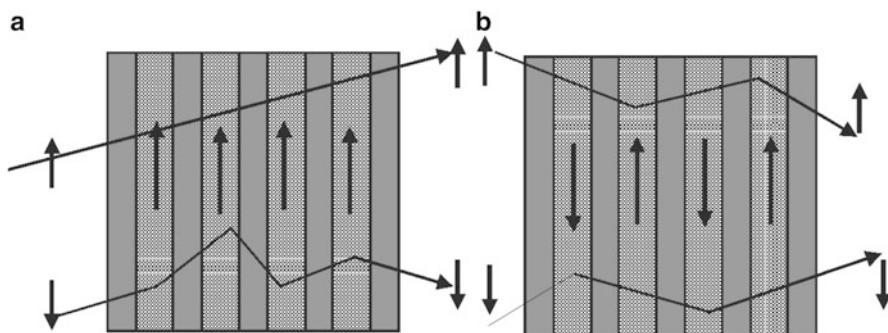


Fig. 10.12 Schematic of occurrence of magnetoresistance in magnetic multilayers with (a) ferromagnetic and (b) antiferromagnetic coupling. In Fig. (b) both up and down spins get scattered unlike in (a) where spin up suffer less scattering and the overall resistance in this case is less

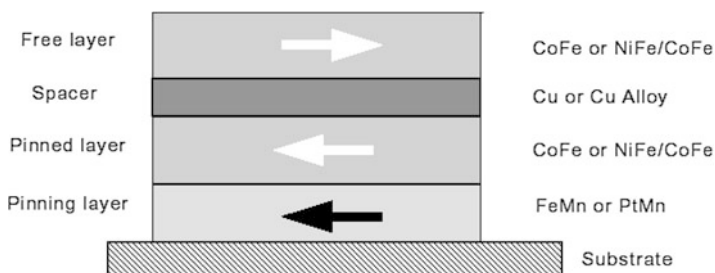


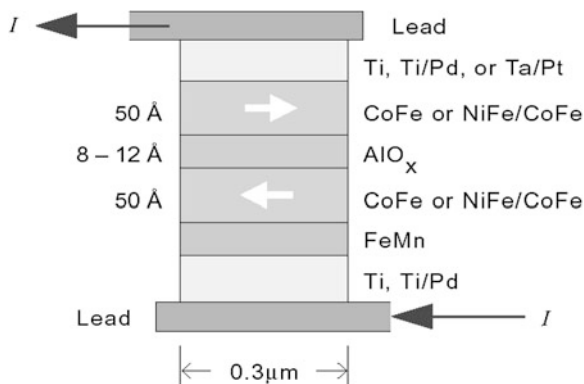
Fig. 10.13 Schematic of a spin-valve structure

upon the thickness of the interlayer (copper layers in this example), the magnetic layers can be coupled ferromagnetically or antiferromagnetically. There are many combinations of magnetic multilayers known now and have been used in computers.

10.4.2 Spin Valve

Based on the GMR effect, multilayer structures have been designed for various applications. Spin valve (Fig. 10.13) is a thin film made up of essentially tri-layers. One layer is a magnetically very soft material, meaning it is very sensitive to small magnetic fields. The other is made of a magnetically ‘hard’ meaning insensitive to fields of moderate strength. The central part of the sample consists of two magnetic layers (for example CoFe or NiFe), separated by a Cu spacer layer. One magnetic layer is pinned or exchange biased by an antiferromagnetic material (for example FeMn or PtMn). As the soft ‘free’ layer moves about due to applied field, the resistance of the whole structure can vary. Spin valves are commercially used in computer read heads. Their use has enabled to increase the data storage capacity of magnetic memory devices due to their ability to detect small magnetic fields.

Fig. 10.14 Schematic of a magnetic tunnel junction



10.4.3 Magnetic Tunnel Junction (MTJ)

In the early 1990s, high magnetoresistance (MR) was discovered for magnetic tunnel junction (MTJ) material. MTJ material is made of at least two magnetic layers (Fig. 10.14) separated by an insulating tunnel barrier. The current flows perpendicular to the film plane. The best results have been achieved with aluminum-oxide tunnel barriers. Since the initial experimental discovery of MTJ material with promising MR, the technique of producing these materials, as well as key properties, has been dramatically improved. Tunnelling MR values are in the 20–50 % range.

Instead of using inorganic insulating barrier layers like aluminium oxide, attempts are also made to insert organic insulating layers to make MTJ devices. This can also lead to the fabrication of flexible organic devices in future.

10.4.4 Spin Field Effect Transistor (SFET)

The Spin polarized Field Effect Transistor (SFET) was first proposed by S. Datta and B. Das way back in 1990. In this case the source and the collector (or drain) are ferromagnetic or half metal materials. The electrons from the source are injected in a planer nonmagnetic metal and are to be transmitted through a planer non-magnetic metal to the collector. In the metallic layer the spin polarization gets reduced which can be controlled by the gate voltage. When the electron spins in the non-magnetic 2-D, metal are parallel to both source and collector. They are able to pass into the collector (on-state) and if they are antiparallel (off-state) then they are not able to pass into the collector or carry current. In fact the electrons in the non-magnetic 2-D metal layer precess (through an effect known as Rashba effect) and need to be controlled through the gate voltage. Realization of SFET took about two decades after it was proposed. Polarized beam of light was used to obtain spin polarized electrons from the source. It is expected that this will help in faster and efficient data processing.

10.5 Nanophotonics

When nanostructures (quantum dots, nanowires or 2-D thin films) or nanocomposites are used to produce light or detect light we are already dealing with a branch in nanoscience known as Nanophotonics. It is expected that one day we will have nanophotonic chips just like semiconductor chips in which light production, propagation, manipulation like amplification, filter, detection etc. can be performed on a ‘nanochip’. Manipulation of signals would then be faster. The present research is towards achieving these goals. The interaction of light with wavelength smaller or comparable to the metal or semiconductor nanostructure sizes was discussed in Chap. 8. We already saw that some of the effects like localized surface plasmon resonance, surface polariton and exciton excitation occurred as a result of interaction of light with nanostructures. We also saw that, as a result of light confinement, there occurred Near Field Effect and consequently one could use it for overcoming the diffraction limit of microscopes. One also has possibility of using photonic crystals to manipulate light. The photonic crystals are abundant in nature through the peacock feathers, butterfly wings etc. The beautiful colours we see in these living animals are the result of periodic arrangement of some proteins which allows certain wavelengths of lights to pass through them and certain wavelengths are forbidden. This is similar to electron states in a semiconductor. We know that in a semiconductor, between the valence and the conduction band, there is an energy gap in which no electron states exist and is known as forbidden gap. It means that the electrons with energy between valence and conduction band are not allowed to propagate in the crystal. Similarly in a photonic band gap material, the photons of certain wavelength (or frequency) are forbidden to propagate. Thus a photonic band gap results. Such a gap can be created artificially by arranging nanoparticles of small uniform size in periodic lattice. The variation of size or dielectric constant is sensitive to optical gap and can be used to sense small variations in the photonic crystals.

Nanophotonics is currently a developing branch and there is plenty of scope to obtain novel materials and their structures to manipulate and propagate light through small structures for ultrafast communication systems.

Further Reading

- G.W. Hanson, *Fundamentals of Nanoelectronics* (Pearson Education, Upper Saddle River, 2009)
O. Manasreh, *Introduction to Nanomaterials and Devices* (Wiley, Hoboken, 2012)
S.M. Sze, *Physics of Semiconductor Devices*, 2nd edn. (Wiley, New York, 1999)

Chapter 11

Some Special Nanomaterials

11.1 Introduction

We discussed in last few chapters synthesis, characterization and properties of nanomaterials in general. Few examples were given from time to time. In this chapter we shall discuss some nanomaterials like fullerenes, graphene, carbon nanotubes, porous silicon, aerogel, zeolites, self assembled materials and core-shell particles, which form a large section of nanomaterials due to their novel properties.

11.2 Carbon Nanomaterials

Carbon is one of the very interesting elements which constitutes a major part of the living as well as non-living world. It also is the backbone of an important branch of chemistry viz. polymer chemistry. It is not surprising that clusters and nanomaterials that we know today provide a rich variety of carbon forms. We can get carbon in various forms as 0-D (small clusters and fullerenes as well as nano diamonds), 1-D (carbon nanotubes), 2-D (graphene and graphane) and 3-D (diamonds) materials. In fact it is often said that half of the community in science is working on carbon based, particularly fullerenes, graphene and carbon nanotubes, indicating the importance of these materials. In this section we shall briefly try to understand these nanomaterials.

11.2.1 Fullerenes

Crystalline carbon can exist in diamond, graphite, fullerene, carbon nanotubes and graphene forms. Out of these allotropes of carbon, fullerene, carbon nanotubes and graphene are relatively new. Fullerene was discovered around 1985. Fullerene

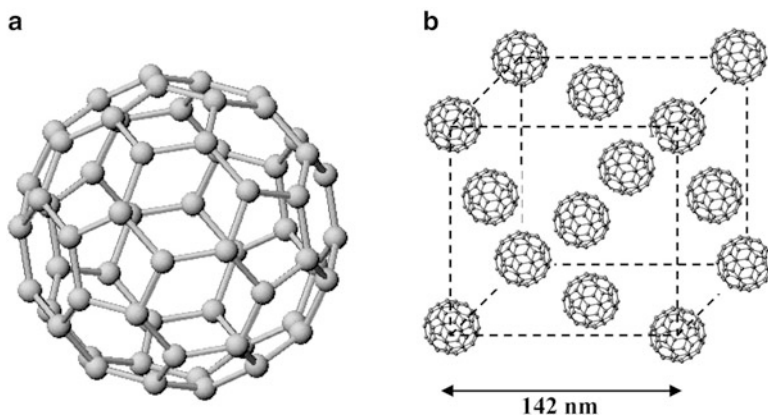


Fig. 11.1 (a) Fullerene C₆₀ and (b) Crystalline form of C₆₀

crystal is a molecular crystal having Face Centered Cubic (FCC) structure as illustrated in Fig. 11.1. At the corners and faces of a cube there are ball shaped carbon molecules which themselves belong to fullerene family. Each fullerene molecule has a cage-like structure. All the carbon atoms are located on the surface of a nearly spherical figure. Fullerenes can have 60, 70, 78 or more (specific or magic) number of carbon atoms on the surface, out of which 60-atom molecule (or cluster) is the most stable and spherical in shape.

11.2.2 Carbon Nanotubes (CNTs)

This 1-D form of carbon was accidentally observed in 1991 by S. Iijima under a transmission electron microscope. He was actually examining some sample of carbon clusters viz. ‘fullerenes’ synthesized using electric arc discharge method. Carbon nanotubes were so less in this sample that Iijima himself said somewhere that it was like observing ‘a needle in a hay stack’. However prior to this observation, Smalley, one of the Nobel laureates who made the discovery of Fullerenes had speculated that like fullerenes which were spherical or nearly spherical molecules, even tubes should be possible. This was followed by some elegant theories predicting the nature of electron and phonon spectra of carbon nanotubes. Discovery of nanotubes was a breakthrough in that sense. His experiments suggested that even a simple set up as used in producing fullerenes should produce even nanotubes under certain conditions. It was quite clear as more and more scientists started reporting the nanotubes that the conditions favoured for fullerenes were unsuitable for producing the nanotubes. Novelty and later the potential applications of nanotubes created a wave of excitement amongst the scientists which led to unfolding of many unique properties carbon nanotubes possess. Their potential applications in electronics, optoelectronics and energy saving systems have been well realized. Just in one

decade several groups all over the world have dedicated their research activities to synthesize and analyze the nanotubes. Following the original arc deposition method, it was found that chemical vapour deposition, laser ablation and some other methods could be employed to produce carbon nanotubes. Further it was found that not only carbon but many other materials like ZnO, TiO₂ and MoS₂ can have shape of nanotube (Box 11.1). They have their own applications but carbon nanotubes still remain the most important ones due to their technological potential. A number of excellent reviews and books have been published on carbon nanotubes.

Here we shall briefly discuss some of the unique features of carbon nanotubes, synthesis methods, properties and applications.

Box 11.1: Why the Name Fullerenes?

Why the name fullerene is given to such cage-like carbon clusters has an interesting story. H.N. Kroto of the Sussex University, U.K. and R. Smalley's group in the Rice University, U.S.A. were performing some experiments in the Rice University. Their objective was to simulate some interstellar molecules in the laboratory. They were trying to laser ablate a graphite rod and study its mass spectrum using a mass spectrometer. Evaporation from graphite was carried out in ultra high vacuum chamber in helium atmosphere. To their great surprise they repeatedly got a mass signal due to 60 carbon atom cluster. They tried to model it nearly for a week and suddenly Kroto remembered of an industrial exhibition he had once visited. There he had seen a dome-like steel structure constructed by a very famous architect Buckminster Fuller. This dome had some pentagons and some hexagons distributed on its surface. Kroto, Smalley and Curl succeeded in constructing 60 atoms cluster model with 12 pentagons and 20 hexagons with a restriction that no two pentagons touched each other and each pentagon shared an edge with hexagon. Later it was found that one could also have carbon clusters with more number of atoms having 12 pentagons but more number of hexagons. It was in the honour of Buckminster Fuller that the name Fullerene was given to the family of this newly found cage-like carbon cluster. It is normally written as carbon atom symbol C with suffix to denote the number of atoms present. For example C₆₀, C₇₀, C₇₈ etc. are 60 atoms, 70 atoms and 78 atoms carbon fullerenes. Kroto, Smalley and Curl received the Nobel Prize in 1996 for their discovery.

However it was not until 1989 that fullerenes were widely investigated. The original apparatus of Smalley group at Rice University where the discovery of fullerenes was made was quite expensive and not available at many places. Besides one could not get any samples out of this apparatus, it was only detection by a mass spectrometer. Huffman and Krätschmer when they heard of discovery of fullerenes wondered if they too were producing this form

(continued)

Box 11.1 (continued)

of carbon as they had found that they too in their experiments had some strange form of carbon. Their further work showed that indeed they too were producing fullerenes. To their great surprise with a simple set up they were producing large quantities of fullerenes. When their work was published, there started a big activity all over the world to synthesize fullerenes. It was also found that it is possible to trap some ions inside the fullerene cage or attach some functional molecules to fullerene from outside. It was realized that due to small size (diameter of C_{60} fullerene is just 0.7 nm; that is why we can consider fullerenes as 0-D material) of fullerenes, they can be used for even drug delivery. The fullerenes could crystallize in a solid form. Properties of fullerenes and their crystals are well investigated now. Smaller clusters of carbon atoms (<60 atoms) also have been investigated. They are detected in mass spectrometers but are not stable and not possible to collect. But most important is that investigations on fullerenes also led to the discovery of Carbon Nanotubes which are found to have a great technological potential.

11.2.2.1 What Are Carbon Nanotubes?

Carbon nanotubes can be considered as cylinders made of graphite sheets, mostly closed at the ends, with carbon atoms on the apexes of the hexagons, just like on a graphite sheet. Thus, as shown in Fig. 11.2, one can consider carbon nanotube as folding of a graphite sheet (it is only an imaginary sheet, actual growth can be different), just like one rolls a piece of paper into a cylindrical form. The difference however is that, a paper is a two dimensional solid material (area much larger, few cm^2 , as compared to thickness of $\sim$ few micrometres) and graphite sheet we are talking about has an area of few μm^2 and thickness just the atomic size of

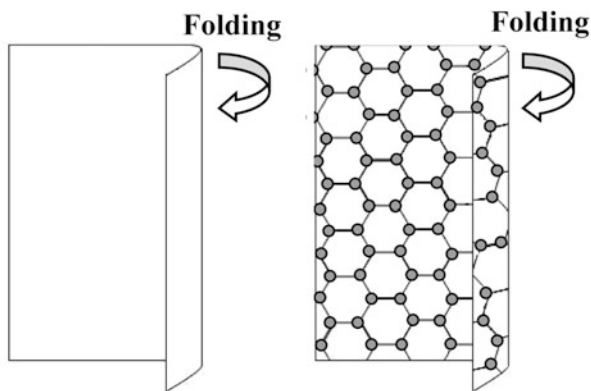
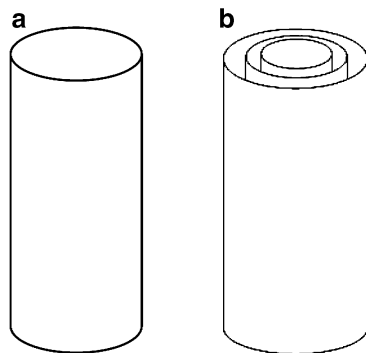


Fig. 11.2 Rolling the carbon sheet to obtain carbon nanotube

Fig. 11.3 (a) SW and (b) MW tubes



carbon atom. If we consider the rolling of graphite sheet, we can imagine carbon atoms being spread in hexagonal arrangement with some lattice strain. The lines showed connecting the filled spheres (carbon) are the bonds that exist between the atoms. Besides, during their formation, nanotubes get capped with hemispheres of fullerenes. It is also possible that many concentric cylinders may be formed as a nanotube. Such concentric nanotubes are termed as Multi Wall Carbon Nanotubes (MWCNT). The distance between their walls is 0.334 nm. This is similar to what one gets between two graphite layers in a single crystal. MWCNT are most common and easily formed. However under certain conditions, it is possible to obtain even Single Wall Carbon Nanotubes (SWCNT). Figure 11.3 illustrates the concept of both SWCNT and MWCNT.

MWCNTs can be turned into SWCNTs using some etching methods. SWCNTs have diameters ranging from 1 to 2 nm. MWCNTs have outer diameters ranging from 2 to 25 nm. The concentrically formed MWCNTs are, however, rotationally disordered (turbotactic). Both MWCNTs and SWCNTs have their own range of applications and studied rigorously.

As the carbon nanotubes can be imagined as folding of a graphite sheet, two things are obvious: (1) carbon atoms on nanotubes are sp^2 bonded like in graphite, although some strain would be expected due to curvature and (2) there should be more than one way of folding the graphite sheet.

Indeed three types of carbon nanotubes (we will consider here only the SWCNT for the sake of simplicity) are possible viz. armchair, zigzag and helical, under appropriate conditions of growth. In order to understand the differences in these three types, consider a graphite sheet as shown in Fig. 11.4.

It is indeed possible to uniquely identify each hexagon as (a, b) with $a = 0, 1, 2, 3 \dots$ and $b = 0, 1, 2, 3, \dots$ only (b, a) is not allowed. Position of any hexagon would be given by a vector $\mathbf{R}$ as

$$\mathbf{R} = a\mathbf{x} + b\mathbf{y} \quad (11.1)$$

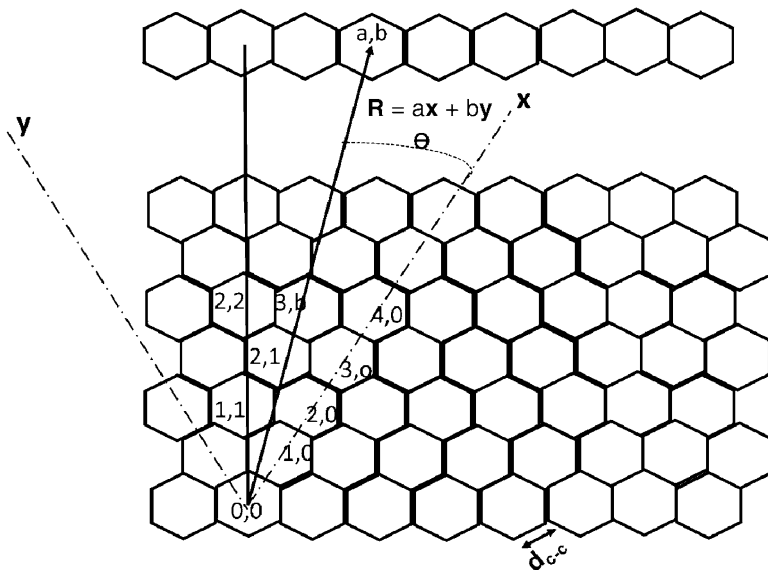


Fig. 11.4 Graphite sheet and generation of tubes

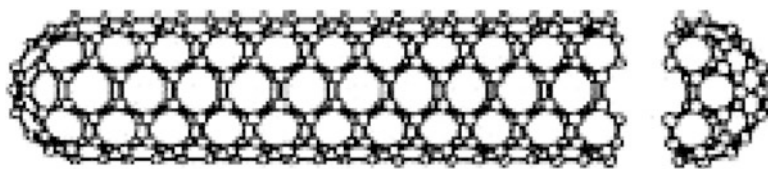


Fig. 11.5 Fullerene with end caps

Consider now the length of primitive vectors $\mathbf{x}$ and $\mathbf{y}$, in terms of distance ' d ' (d_{c-c}) between two nearest carbon atoms. It is clear from the simple geometrical considerations that

$$\mathbf{x} = \sqrt{3.d} \quad \text{and} \quad \mathbf{y} = \sqrt{3.d} \quad (11.2)$$

Vector $\mathbf{R}$ denoting the position of a hexagon is known as a 'chiral vector'. A tube obtained by folding the sheet along $\mathbf{R}$ (a, b) is called a chiral tube (a, b). An angle between x -axis and vector $\mathbf{R}$, θ also can be used to denote the folding. It is observed that all the angles between $0 < \theta < \pi/6$ are sufficient to uniquely define different types of tubes except (b, a) is not possible. The tubes whose mirror image is identical with its own image is known as 'achiral' tube and 'chiral' otherwise. The tubes are normally terminated with the hemispheres of fullerenes (see Fig. 11.5), as was predicted by Smalley.

These are simply called as 'caps' or 'end caps'. Caps contain six pentagons (half the number in C_{60} fullerene) and different number of hexagons so that they can fit

on the tubes properly. It should be remembered that in the notation used here, for a nanotube we need to have $a > b$. Diameter of the nanotube is obtained as follows

$$D = \frac{\text{circumference of tube}}{\pi} \quad (11.3)$$

$$\begin{aligned} \text{Circumference of tube} &= |\mathbf{R}| \\ &= \sqrt{\mathbf{R} \cdot \mathbf{R}} \\ &= \sqrt{(ax + by) \cdot (ax + by)} \\ &= \sqrt{3} \cdot d_{c-c} \sqrt{(a^2 + ab + b^2)} \end{aligned} \quad (11.4)$$

as

$$x \cdot x = y \cdot y = 3d_{c-c}^2 \quad (11.5)$$

$$D = \sqrt{3} \cdot d_{c-c} \frac{\sqrt{(a^2 + ab + b^2)}}{\pi} \quad (11.6)$$

where D is diameter, d_{c-c} – distance between two nearest carbon atoms, and a and b are chiral lengths of vector $\mathbf{R}$.

Angle θ in terms of chiral lengths a and b is obtained as

$$\cos \theta = \frac{\mathbf{R} \cdot a\mathbf{x}}{|\mathbf{R}| \cdot |a\mathbf{x}|} \quad (11.7)$$

$$= \frac{(a\mathbf{x} + b\mathbf{y}) \cdot a\mathbf{x}}{\sqrt{3} \cdot d_{c-c} \sqrt{(a^2 + ab + b^2)} \sqrt{3} d_{c-c} \cdot a} \quad (11.8)$$

$$\theta = \cos^{-1} \left\{ \frac{2a + b}{2(a^2 + ab + b^2)} \right\} \quad (11.9)$$

For the angles $0 < \theta^\circ < \pi/6$.

11.2.3 Types of Carbon Nanotubes

Depending upon their chirality or the way of folding as discussed above for a SWCNT, basically three types arise viz. zigzag, armchair and helical CNT.

Zigzag CNT: These are formed for $\theta = 0$ and chirality $(a, 0)$. i.e. by folding parallel to x-axis. The name zigzag has been given due to zig zag arrangement of carbon atoms that can be seen (Fig. 11.6) if cross section of the tube as shown in the figure is taken. This type of tubes are ‘achiral’ tubes i.e. their mirror images are similar as the original structure.

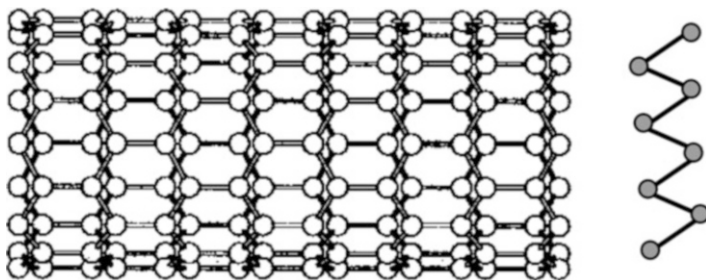


Fig. 11.6 Zigzag CNT

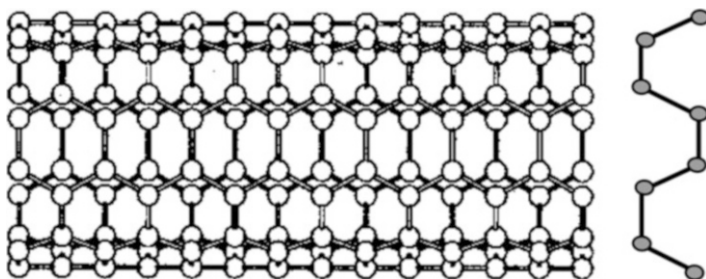


Fig. 11.7 Armchair CNT

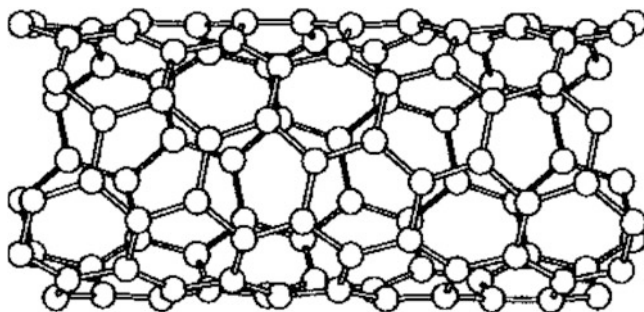


Fig. 11.8 Helical CNT

Armchair CNT: These are formed for the angle $\theta = \pi/6$ and chirality (a, a) . These also are ‘achiral in nature’.

Armchair structure with arrangement of carbon atoms can be seen in Fig. 11.7.

Helical CNT: These are obtained when angle θ is anywhere between 0 and $\pi/6$ and chirality is (a, b) . Helical structure of CNT is shown in Fig. 11.8.

Such tubes are ‘chiral’ and their mirror images appear to differ from their original structure. Table 11.1 shows the comparison of three types of CNT.

Table 11.1 Comparison between different types of CNTs

Type of CNT	R	θ°	Cross section
Zigzag (achiral)	a, 0	0	Trans
Armchair (achiral)	a, a	30	Cis
Helical (chiral)	a, b	$0 < \theta < 30$	Mixture of trans and cis

Besides these basic types, a variety of shapes like ropes, springs, stripes, bamboo structure, conical shapes etc. are observed to have been formed of CNTs under different experimental conditions.

11.2.4 *Synthesis of Carbon Nanotubes*

Iijima had detected carbon nanotubes in an electric arc discharge set up for synthesizing fullerenes. His nanotubes were multiwall type. The yield of nanotubes was very low compared to the fullerene content. Due to tremendous interest the scientific community took in CNTs, soon arc discharge and other methods like laser ablation and chemical vapour deposition were optimized to increase the yield and even to get SWCNTs. It is now understood that electric arc discharge mostly can produce MWCNT and laser ablation can be used for SWCNT. There are also some attempts to use ion beams to obtain CNTs. However such methods are uncommon. In the following we briefly discuss the parameters and few points regarding the CNT synthesis using electric arc discharge, chemical vapour deposition and laser ablation techniques. More details about these techniques are discussed in Chap. 3.

11.2.4.1 Electric Arc Discharge

Carbon nanotubes, in an electric arc discharge set up (see Fig. 11.9) between two graphite rods as electrodes, are formed under certain conditions as follows.

Electrodes: Graphite

Diameter of electrodes: 5–20 mm

Gap between the electrodes: ~ 1 mm

Helium gas pressure: 100–500 Torr (no CNT formed if $p < 100$ Torr)

Current: 50–120 A

Voltage: 20–25 V

Cooling of the chamber is preferred.

Yield of MWCNT: ~ 30 –50 %

Nanotubes are found to be deposited in the central region of cathode. For MWCNT it is not necessary to use any catalyst. Central region of cathode reaches a temperature close to 3,000 °C and nanotubes are aligned in the current direction

Fig. 11.9 Electric arc discharge set up

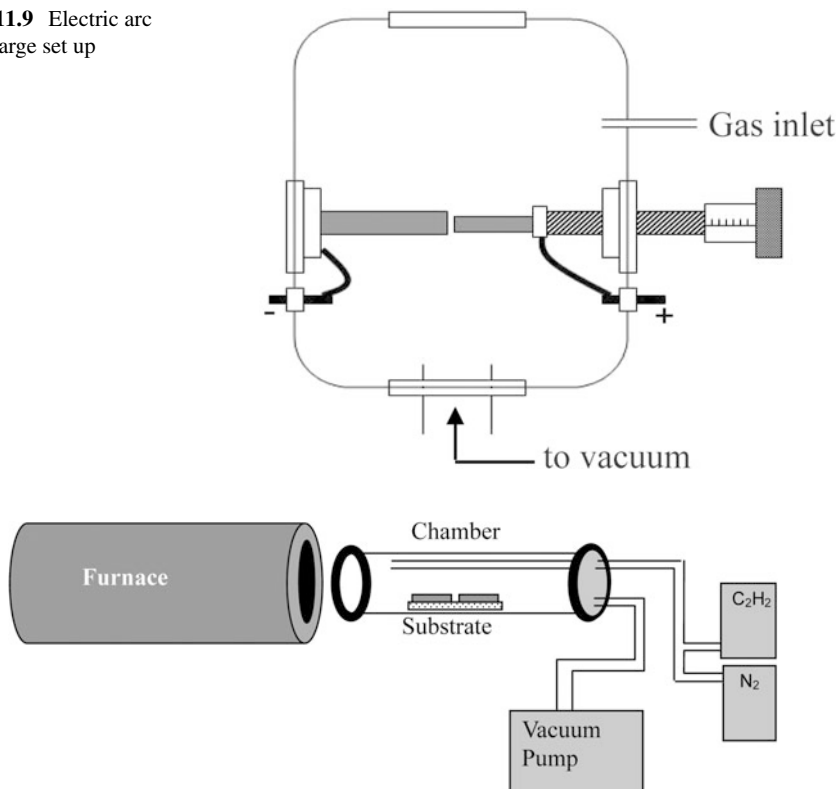


Fig. 11.10 CVD set up (schematic). The chamber and the substrates are shown outside the furnace just for the illustration

between the two electrodes. The central region of cathode is surrounded by a region in which nanoparticles, fullerenes and amorphous carbon are formed. It, therefore, becomes necessary to purify the MWCNTs to get rid of these other particles. No tubes are found to be deposited on the walls of the experimental chamber as is the case with fullerenes.

Diameters of nanotubes are in the range of 0.7–1.5 nm. Incidentally, 0.7 nm is the diameter of C_{60} fullerene. Length of the tubes is typically 1 μm . However lengths in the few mm to mm range also have been possible.

11.2.4.2 Chemical Vapour Deposition (CVD)

For large scale production of the nanotubes, this method is most useful. Here a hydrocarbon gas is cracked under certain conditions. Experimental set up is shown in Fig. 11.10.

As there are no graphitic hexagons present in some of the precursors used to deposit carbon nanotubes, catalysts play a very important role in carbon nanotubes formation. Advantage of this method is that aligned nanotubes can be deposited on some solid substrates so that they can be used for some electronics or other application. Both MWCNT and SWCNT are possible to obtain by this method. No nanoparticles or amorphous carbon formation takes place, making high purity nanotubes.

Gases: CH_4 , C_6H_6 etc.

Pressure in the chamber: $\sim 10^4$ Pa

Catalyst: Fe, Co, Ni or Pt

Furnace Temperature: $\sim 1,000^\circ\text{C}$

11.2.4.3 Laser Ablation

Schematic sketch of the laser ablation set up for CNT deposition is given in Fig. 11.11. Advantage of using laser in the synthesis of carbon nanotubes is that the nanotubes are invariable SWCNTs. Ropes of 10–20 nm diameter and lengths $\sim 100\ \mu\text{m}$ also are observed. Narrow size distribution of diameters of SWCNTs is an attractive feature of this technique.

11.2.5 Growth Mechanism

One may wonder at this stage as to how do the nanotubes grow at all. At the synthesis temperatures as high as used in CVD ($>800^\circ\text{C}$) and electric arc deposition ($\sim 3,500^\circ\text{C}$) or laser ablation, we really don't think that long sheets of graphite are really released and folded. It is probably *atom by atom* or *molecule by molecule addition* that under certain favourable conditions nanotubes are formed. Some scientists think that an initially formed fullerenes with 12 pentagons and hexagons would open up its cage to accommodate a dimer (C_2) and/or a trimer (C_3) as shown in Fig. 11.12.

Pentagons on a fullerene are more strained and can break to accommodate a C_2 dimer or C_3 trimer. This can continue and long nanotubes be formed. Some

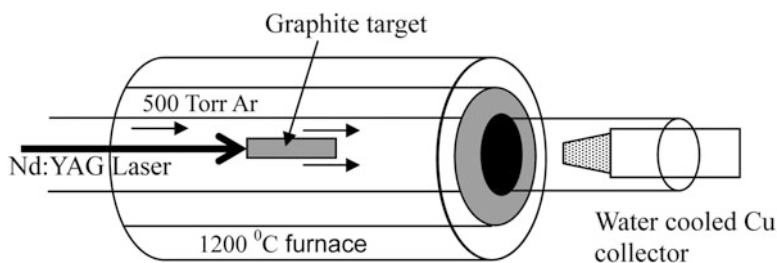


Fig. 11.11 Laser set up (schematic) to deposit CNTs

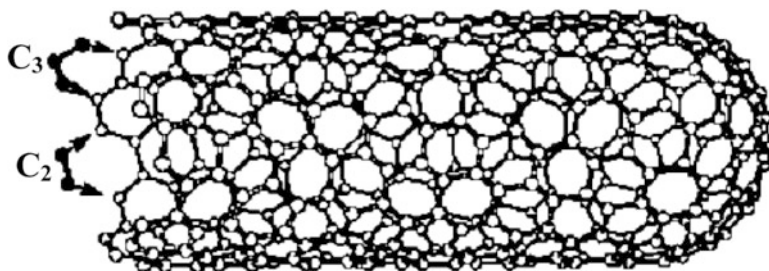
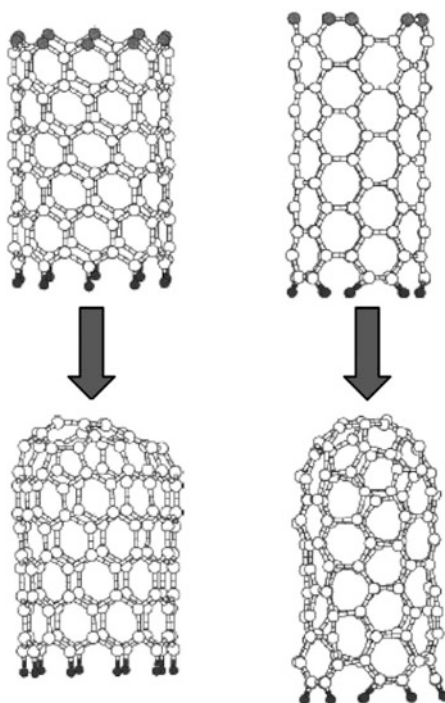


Fig. 11.12 Growth of CNT by absorption of C_2 dimers and C_3 trimers

Fig. 11.13 Carbon without end cap, dimer, trimer addition for growth and spontaneous closure



others are of the opinion that during the growth process, nanotubes are not capped but ultimately get end caps (Fig. 11.13). Other structures with concentric tubes and different shapes also have been explained. As there are a variety of methods, temperature, catalysts and pressure which dictate the growth, it is rather impossible that a single mechanism would be responsible for the nanotube growth.

Experiments to determine the conductor type of CNT includes Scanning Tunneling Spectroscopy, Electron Spin Resonance and ^{13}C nuclear magnetic resonance. Measurements are in agreement with the theoretical prediction that both semiconductor and metallic nanotubes should be formed with above mentioned conditions.

Table 11.2 Carbon polymorphs and their properties

Dimension	3D	2D	1D	0D
Polymorph	Diamond	Graphite/graphene	Nanotube	Fullerene
Bonding type	sp ³	sp ²	sp ²	sp ²
Bond length, Å	1.54	1.42	1.44	1.40
E_g , eV	5.47	0	1–0	1.9
Conductor type	Insulator	Metal/semimetal	Semicond. to metal	Semicond.
Density, gm/cc	3.52	2.26	~1.2–2.02	1.7 (for C ₆₀)

Superconducting transition temperature T_C in CNT is as low as $\sim 1.5 \times 10^{-4}$ K. Doping can alter T_C in superconductors as well as affect metallic or semiconductor behaviour (Table 11.2).

11.2.6 Graphene

Graphene is just a single layer of graphite crystal. As was discussed in Chap. 2, a graphite crystal has stacks of carbon layers weakly bonded to each other. That is why these layers can easily slip over each other. Each layer has hexagonally arranged carbon atoms. Such layers when appear as single, self standing material is graphene, when rolled into a tubular form it becomes carbon nanotube or even becomes something like a fullerene.

Graphene, which is an isolated sheet of carbon atoms, was considered to be difficult to achieve. However, in 2004 it was shown by the scientists from the University of Manchester, U.K. and Institute of Microelectronics Technology, Chernogolovka, Russia that it is possible to separate single carbon sheets (monolayer thick) by simply peeling them off using a scotch tape. The properties of graphene sheets could be measured which led to lot of interesting work on graphene. It is also possible to produce graphene by other methods like chemical vapour deposition. Electronic structure and transport studies have revealed a great deal of interesting properties of graphene. For example quasiparticles in graphene are found to be massless Dirac fermions and Quantum Hall Effect which is usually observed at very low temperatures is observed at room temperature in graphene. Discussion on these topics is beyond the scope of this book. We just end this topic here by remarking that it is a topic with tremendous scope in electronics, spintronics (spin coherence length as large as 1 μm is observed in graphene at room temperature), sensors as well as quantum computers. The electron mobility of graphene at room temperature is very large and its resistivity is even smaller than silver. Mechanical properties of graphene are superior to steel. It is ~ 200 times stronger compared to steel. The work on graphene has been considered to be a breakthrough which led Geim and Novoselov to receive prestigious Nobel Prize in the year 2010.

As there would be dangling bonds on carbon atoms, they can be passivated by hydrogen. This form of carbon is known as ‘graphane’.

11.3 Porous Material

There is a large variety of porous materials with pores of few nanometers to micrometres. Porous materials are classified as microporous, mesoporous and macroporous materials based on the pore sizes. Materials with pore diameters less than about 2 nm are called microporous materials and those with larger than about 50 nm size are known as macroporous. The pores of intermediate sizes i.e. between 2 and 50 nm are referred to as mesoporous materials. There exist some natural porous materials like soil, snow, sandstone, corals, bones, wood and so on. Several examples of manmade porous materials around us are bread, biscuits, sponge, brick, chalk etc.

In this section we shall discuss some of the artificially synthesized porous materials. These materials have emerged over a long period and have considerable applications from drug delivery, energy storage to space science. Due to their porous nature in general, all types of porous materials inherently have lower density and larger surface area (includes surface area of pore walls) than the corresponding bulk materials.

11.3.1 Porous Silicon

Silicon is the most widely used semiconductor material by electronics industry. It is abundant in nature and techniques to purify, grow single crystals economically and on large scale are well developed. The methods of doping and polishing also have been perfected. Oxide of silicon is stable and metal-silicon contacts can be made and understood to good extent. Thus microelectronics industry is well established to use silicon in making various components and systems like diodes, transistors, Field Effect Transistors (FETs), Metal Oxide Semiconductor FETs (MOSFETs), Integrated Circuits (ICs), Very Large Scale Integrated Circuits (VLSIs) and Mechano-Electrical Machines (MEMS). However, there is a basic drawback with silicon viz. unlike some other semiconductors like ZnS, CdS, GaAs, InAs and InP, light emission capability of silicon is extremely poor. Therefore it is not possible to make light emitting diodes (LED) or lasers using silicon wafers that are usually used in electronics. It was thus considered for a long time that silicon cannot be an optoelectronics material.

However in 1990, Canham showed that porous silicon can be very easily and routinely formed on silicon wafers, used by electronics industry, to emit photoluminescence in the visible range at room temperature. He further showed that silicon nanocrystallites formed, as a result of pores formation, were responsible for light emission and as the crystallite size reduced, the emission wavelength also decreased. This triggered the wave of enthusiasm in the scientists and technologists, as it was a major breakthrough for silicon industry to use silicon as an optoelectronic material. As a result, considerable work has been carried out since 1990 to make porous silicon by number of ways, use different substrates, etchants etc. and investigate the various structural, optical and electrical properties and demonstrate new applications. Even other semiconductors like GaAs also have been made porous and inves-

tigated. It should be remembered that scientists were using for quite a long time the procedures, specially the electrochemical polishing, similar to that used in porous silicon formation and some interesting properties were already noticed. However, prior to Canham's publication, emphasis of early work was on studying the properties of thin film observed during etching or electropolishing of silicon. Detailed investigations showed that such films were porous and since 1970 the word 'porous silicon' was being used. Canham showed the nanocrystalline nature of the pore walls and its role in luminescence that became interesting and decisive to consider porous silicon as an optoelectronic nanocrystalline material. Here we will discuss, how porous silicon can be made, its properties and briefly some applications (Box 11.2).

Box 11.2: History of Silicon Etching

- 1956—A. Uhler at Bell Labs. USA [ref. *Bell Syst. Tech. J.* **35** (1956) 333] observed while polishing silicon by electrochemical method that at high current density polishing occurred but at low current density surface looked as if there was some blackish, brownish or sometimes reddish deposit. He attributed this to the formation of some silicon suboxides.
- 1958—D.R. Turner [ref. *J. Electrochem. Soc.* **105** (1958) 402; *Surface Chemistry of Metals & Semiconductors* (H.C. Gatos, ed.), Wiley, New York (1960); *The Electrochemistry of Semiconductors* (P.J. Holmes, ed.) Academic Press, New York, Vol. 155 (1962)] attributed anodic film formation to $(\text{SiF}_2)_n$ polymeric network with predominance of silicon. He considered that the formation of anodic regions depends upon whether oxidizing or reducing species attack a certain region. When HNO_3 attacks, it removes an electron from the surface making it a local cathode. The positively charged hole is created as a result of loss of electron. Hole can wander around and towards Si where HF attacks. It is then ready to form an anodic region. Silicon complexed with fluorine, dissolves into the solution forming a pit.
- 1959–1956—Robins and Schwartz used HNO_3 and HF solution for etching silicon. In a series of papers [*J. Electrochem. Soc.* **106** (1959) 505; *ibid.* **107** (1960) 108; *ibid.* **108** (1961) 365; *ibid.* **123** (1976) 1903], they discussed that process of etching first involved the formation of silicon oxide and then its removal by HF. They, however, did not mention the formation of any porous film.
- 1960—R.J. Archer [*J. Phys. Chem. Solids* **14** (1960) 104] mentions porous silicon formation and observation of silicon hydride formation.
- 1965—K.H. Beckman [*Surf. Sci.* **3** (1965) 314] studied infra red spectrum of porous silicon and showed that polymeric network of $-\text{SiH}_2$ existed on the etched silicon surface. If electrochemically prepared SiH_2 gets connected through Si-Si bonds and if chemically prepared, the network is connected through Si-O-Si bonds.

(continued)

Box 11.2 (continued)

- 1972—M.J.J. Theunissen [*J. Electrochem. Soc.* **119** (1972) 351] clearly showed for the first time that porous silicon had crystalline nature.
- 1983—G. Bomchil, R. Herino, K. Barla and J.C. Pfister [*J. Electrochem. Soc.* **130** (1983) 1611] measured surface area and sizes of pores in various lightly and heavily doped p and n type silicon samples.
- 1984—C. Pickering, M.I.J. Beale, D.J. Robinson, P.J. Pearson and R. Greef [*J. Phys. C* **17** (1984) 6535] carried out optical studies of etched Si. Refractive index of p^+ -Si varied from 3.47 towards 1.0 with increasing porosity. They observed the presence of complex amorphous alloy of Si, H and O along with crystalline Si in the porous film. They also mentioned the observation of an intense band gap luminescence at 4.2 K as well as orange red luminescence. They attributed it to amorphous SiH_x formation. Further by changing H concentration, room temperature photoluminescence also was observed.
- Rutherford Back Scattering (RBS) and nuclear measurements of porous silicon showed that it was crystalline and surface was covered with 50 % H, 5 % O, 1 % F and 3 % C (mostly as contaminant). SiH formation was concluded.
- 1988—Work on porous silicon was reviewed by G. Bomchil, A. Halimaoui and R. Herino [*Microelectron. Eng.* **8** (1988) 293]. Even some applications of porous silicon for semiconductor on insulator (SOI) appear but optical studies were neglected.
- 1990—Canham [*Appl. Phys. Lett.* **57** (1990) 1046] proposed that quantum confinement took place as a result of nanocrystallites formation in porous silicon. Photoluminescence was observed from porous silicon at room temperature and its wavelength depended upon the size of nanocrystallites produced by etching. This observation of Canham can be considered as the most important step in porous silicon research.

11.3.2 How to Make Silicon Porous?

Surface of solid silicon can be turned into porous layers by a variety of methods like ion irradiation, spark erosion, chemical etching or electrochemical etching. However, many of the investigations show that most effective method of making porous silicon is electrochemical etching. Chemical etching involves dipping the crystalline silicon in some strong etchant solution like HF mixed with HNO₃. The reaction responsible for making pores begins autocatalytically. Reproducible and controlled porosity in relatively short time can be obtained using electrochemical method and is described here.

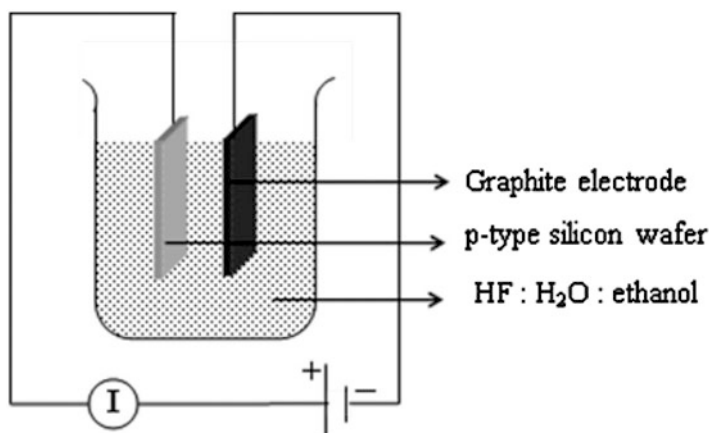


Fig. 11.14 Electrochemical cell

11.3.2.1 Electrochemical Etching

Porous silicon can be easily formed with appropriate conditions at room temperature in a simple electrochemical cell. There are various designs that can be used to obtain anodic etching of silicon but a schematic one is shown in Fig. 11.14. A crystalline piece from a silicon wafer is dipped in an electrolyte, usually a solution of HF, water and ethanol. By making silicon a working anode and a piece of graphite or platinum as cathode a current is passed through the solution. Typical parameters used to etch a *p*-Si are shown in Box 11.3. As HF (hydrofluoric acid) is used as an etchant, precaution has to be taken to use Teflon or polyethylene container as HF severely attacks glass. *It also should be remembered that although the experimental procedure is very simple, HF is extremely dangerous to use. If HF comes in contact with the skin, fluorine makes entry through the skin to the bones affecting the limbs. CaF_2 is formed which mixes with the blood, poisoning it and can turn out to be fatal. Even the waste electrolyte needs to be handled properly, otherwise due to contact with some chemicals it can even explode.*

Box 11.3

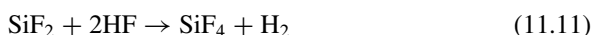
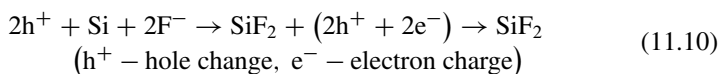
Typical parameters in electrochemical reaction used for making porous silicon

Substrate: <i>p</i> -Si (111) (B doped)	Voltage: 24 V
Current: 30 mA	Solution: HF + ethanol + water

11.3.3 Mechanism of Pores Formation

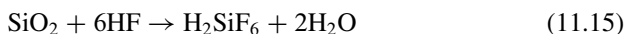
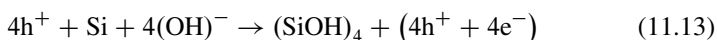
Mechanism of pores formation on silicon surface is a complex process. Attempts are made to understand them from the observed species formed on the surface.

For the formation of a porous network on silicon surface, it is necessary that local cathode and anode regions be formed on its surface. When an electron is extracted into the electrolyte from certain region of silicon, usually some sharp point acts like a local cathode. Release of an electron is equivalent to the production of a positively charged hole. The hole can migrate to some other region on the surface, which becomes an anode. If such a region is also attacked by HF present in the electrolyte, following reaction becomes feasible.



H_2SiF_6 dissolves in the electrolyte forming the initiation of pits. It can be seen from above equations that two holes are necessary for the formation of pores. At some low current value and HF concentration such conditions are created.

As the current density is increased beyond a certain value, it is possible to generate more holes at a time and chances of forming surface oxide are increased. If following reaction becomes possible then it is expected that polishing would occur rather than porous structure.



H_2SiF_6 and water dissolve in the electrolyte. Obviously the h^+ and e^- , already present in the silicon substrate as a result of doping, would greatly influence the pores formation. Specially, it would matter considerably, whether it is a p-type (excess h^+) or n-type (excess e^-) silicon substrate being used for pores formation. It should not also be surprising if doping level—whether low concentration, moderate concentration or extremely high concentration—would alter the nature of pores formed. Depending upon the concentrations of excess charges, either thermionic or tunnelling of charge carriers would be effective.

High porosity of a material means large area of the exposed material to the ambient. Therefore porous silicon can easily get oxidized. The optical properties of porous silicon are found to be unstable after their formation. However they are reversible too.

11.3.3.1 Factors Affecting the Porous Structure of Porous Silicon

From the vast literature available now on porous silicon, few trends are noticed as follows.

The morphology, optical and structural properties of porous silicon depend upon (1) Substrate type (n or p), (2) Doping level (low or high), (3) Concentration of HF in the electrolyte, (4) Current density used to etch silicon, (5) Effect of light, and (6) Duration of etching (Box 11.4).

Box 11.4: Porous Material

Materials become porous when some solid portion is removed from them leaving voids in them. The volume of the original material does not alter. Therefore density of a porous material is lower than the original material. Materials can be made even 98 % porous. There is a large variety of porous materials, natural (e.g. wood) as well as synthetic (e.g. thermocole). The pores can be of regular, irregular and different sizes and shapes. Pores can be closed, open, penetrating or ink bottle type. Internal area increases with pores (Fig. 11.15).

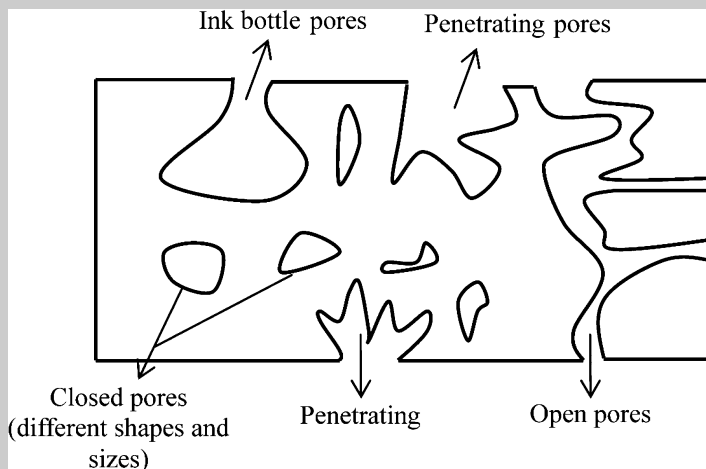


Fig. 11.15 Various types of pores

(continued)

Box 11.4 (continued)

In general, properties of porous materials are quite different than those of corresponding bulk materials. Porous materials are often divided into microscopic (less than 2 nm), mesoscopic (2–50 nm) and macroscopic (larger than 50 nm) depending upon the diameters or size of the pores.

1. **Substrate Type:** Pore structure is normally formed due to holes and silicon atoms interacting together with fluorine. As was discussed in the mechanism of pore formation, it is necessary to consider the substrate type viz. p or n type as they already have excess holes or electrons due to doping. For example it has been found in case of moderately doped n or p type silicon that columnar or sponge-like structure is respectively formed.
2. **Doping Level:** One can identify four types of silicon wafers for forming pores on their surfaces by electrochemical etching. They are: p^+ -Si (very high concentration of holes), p-Si (moderate or low concentration), n^+ -Si (high concentration of electrons) and n-Si (moderate or low concentration of electrons). It is found that p, p^+ and n^+ -Si substrates have similar trends in their pores formation, but n-Si needs large voltage to form pores.
3. **Concentration of HF:** If the current density is held constant, concentration of HF in the electrolyte is an important factor in determining pore formation. Although porosity has been found to increase with decrease of HF concentration, it also is a function of current density. At large current density, usually electropolishing takes place as schematically shown in Fig. 11.16. In the transition region, a rough surface is produced. At lower concentration of HF, the pH of the solution increases, which promotes pore formation.
4. **Current Density:** Figure 11.17 shows three different regions viz. a region of pores formation, an intermediate region and electropolishing region, similar to that in Fig. 11.16.

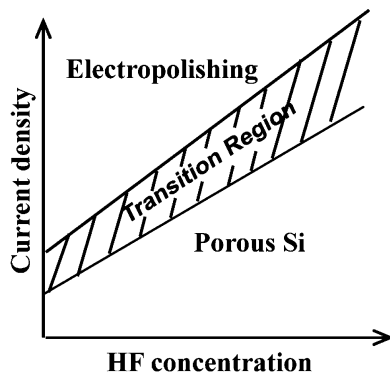
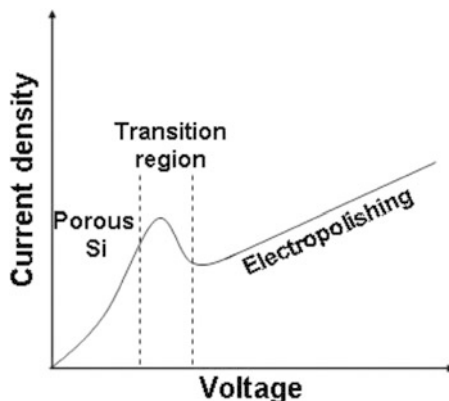


Fig. 11.16 Different regions of silicon surface morphology, which depend upon current and HF concentration in electrochemical etching of Si

Fig. 11.17 Region of silicon morphology which depends upon current density and the applied voltage in the process of electrochemical etching of silicon



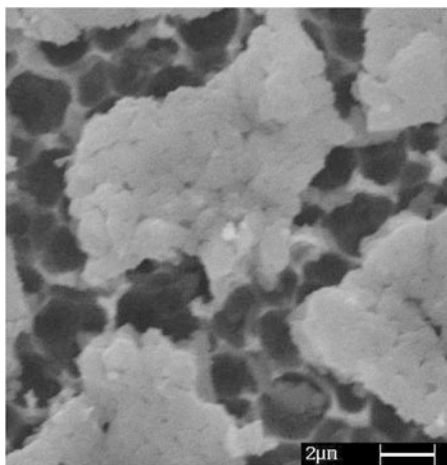
For a given electrolyte, it can be seen that porous structure can be formed only at low voltages and low current densities. The curve may differ for n-Si substrate quite considerably.

5. **Effect of Light:** Chemical reactions can be often controlled using light. When light of appropriate wavelength falls on a semiconductor, it generates electron-hole pairs. If intensity of light is adequate, a large number of charge carriers necessary for the pores formation can be generated using light. Thus the illumination of the silicon anode during electrochemical reaction can affect pores formation. It is found that p-Si can etch out in dark but for n-Si illumination by photons is necessary for pores formation. Sometimes, ambient light can be sufficient.
6. **Duration of Etching:** Etching creates pores which are surrounded by the silicon walls or pillars. These walls contain nanocrystallites of silicon as was first postulated by Canham. Longer the etching time, smaller are these crystallites. The photoluminescence which is observed in porous silicon is due to the formation of these nanocrystals. Therefore properties of the porous silicon are affected by the crystal size and the etching time.

11.3.4 Properties of Porous Silicon Morphology

Depending upon the substrates used for etching, the morphology of porous silicon can change. In general, it is found that 'nanopores' produced with sizes less than ~ 10 nm (difficult to image) when p-Si with low doping level are used. On the other hand if n-Si substrates are etched using large current, pores even larger than ~ 1 μm size are easily formed. Medium sized (mesopores) pores ~ 10 – 100 nm can be formed in heavily doped p^+ or n^+ type silicon. Porosity, of course, will go on increasing with etching time and typically 50–90 % porosity can be achieved. Shapes of pores also can be altered using different substrates. Often circular, square,

Fig. 11.18 SEM image of porous silicon



triangular, star or random shapes and non-homogeneous are observed. Heavily doped p-Si substrate has columnar structure but lightly doped p-Si has spongy surface structure. Heavily doped n-Si substrate also has columnar porous structure. Figure 11.18 shows an SEM image of porous silicon.

11.3.4.1 Structure

Structure of porous silicon is quite complex. Nanocrystallites of silicon are usually embedded in amorphous silicon. However, in order to produce porous structure, it is necessary to use crystalline silicon. Attempts to use amorphous silicon substrates have failed to produce porous silicon.

11.3.4.2 Chemical Nature

Surface of porous silicon alters chemically during the etching process. It is non-homogeneous on microscopic level and different researchers often obtain conflicting results. This is because of the use of different substrates and variations in electro-chemical parameters. In general, along with Si, one finds F, H and O to be present. The relative amounts, however, differ from experiment to experiment. Sometimes carbon also can be detected which arises as a contaminant from the solution or ambient. Samples left in the laboratory environment changes in H or O content with time, also altering the electrical or optical properties. Networks of SiH_2 , SiO_2 or complex Si-O-H are often reported.

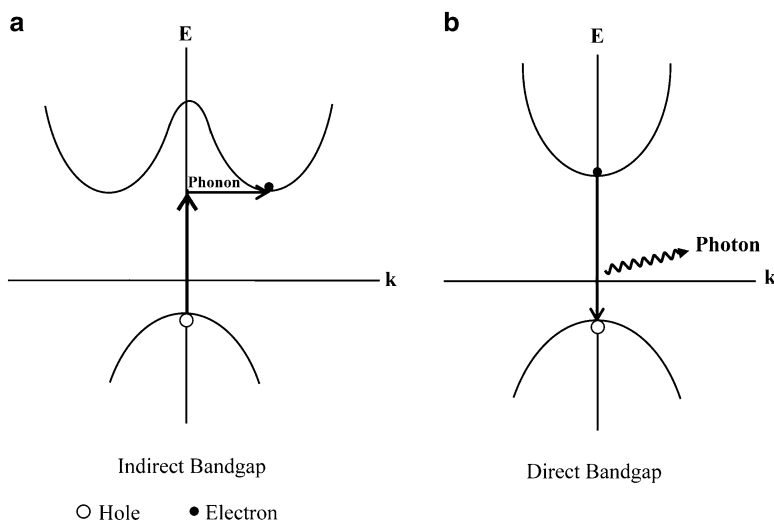


Fig. 11.19 Schematic band diagrams (a) silicon bulk and (b) Nanocrystalline silicon

11.3.4.3 Electronic Structure

Electronic structure of porous silicon differs considerably from that of the bulk, crystalline silicon. For the bulk silicon crystal the band gap is 1.17 eV at absolute zero and is 1.125 eV at room temperature. This band gap is an indirect band gap as schematically shown in Fig. 11.19.

This is because the maximum in the valence band occurs at $k = 0$ or Γ (gamma) point in Brillouin zone. The minimum in the band is closer to Brillouin zone boundary. An electron from valence band, even if excited using a photon with sufficient energy, cannot go to conduction band minimum, as both energy and momentum must be conserved in an absorption or emission process. Thus an electron in conduction band minimum at $k \neq 0$ also cannot make a transition to valence band state at $k = 0$. This is why Si is not an optoelectronic material. However, as soon as Si nanocrystallites are formed, the electronic structure alters dramatically.

As there is no long range periodic arrangement, the small momentum conservation rule breaks down and transitions become allowed. According to the Effective Mass Approximation (EMA) theory, energy gap increases compared to that in the bulk and depends upon the magnitude of the nanocrystallite size.

11.3.4.4 Photoluminescence

Depending upon the preparation conditions and the nanocrystallite size porous silicon can exhibit photoluminescence over the entire visible range as illustrated in Fig. 11.20.

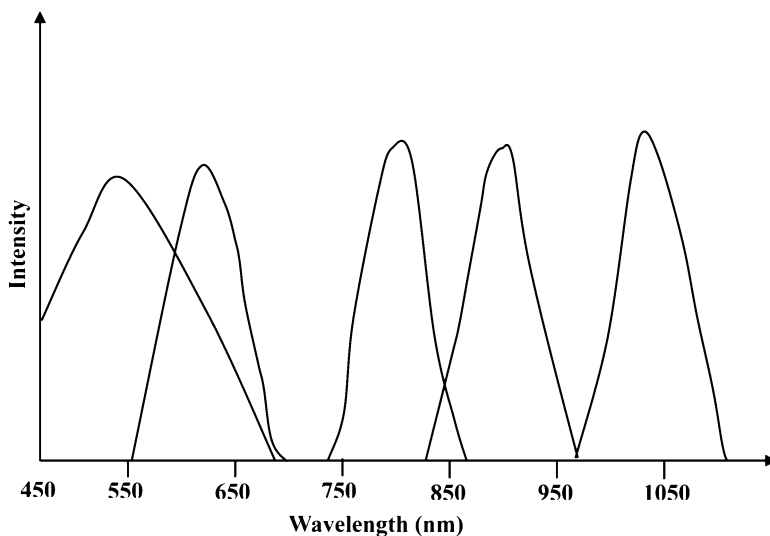


Fig. 11.20 Photoluminescence spectra for different porous silicon samples at room temperature

Due to wide distribution of nanocrystallites, luminescence peak widths can be large.

11.4 Aerogels

Aerogels constitute another class of highly porous materials. These ultra low density, monolithic materials are formed by interconnection of particles of nanometer size, to form a nanoporous solid. The pores themselves are usually nonuniform with sizes from ~ 10 to 100 nm. These materials are synthesized by sol-gel method (discussed in Chap. 4) and dried by special procedures to retain their porous structure. Mixture of reactants forming colloidal particles, which are dispersed in liquid, is known as ‘sol’ (Box 11.5).

Box 11.5: Brief History of Aerogels

Aerogels were first synthesized by an American scientist Samuel Kistler in 1930s. He was working in Stanford University on puzzle of cracking of gels while drying. It was found that gels were made up of network of interconnected solids dispersed in some solvent (water). Drying of gels led to evaporation of this solvent and surface tension exerted by evaporating liquid

(continued)

Box 11.5 (continued)

collapsed the solid network. It was realized that while drying gels, they shrink and develop cracks. Kistler developed, therefore, crack-free drying method by removing the solvent above its critical point. Dried materials thus obtained were crack-free and retained their original size and shape. Kistler replaced water in the gel (which has very high critical temperature and pressure) with alcohol and removed it by putting it in an autoclave (a vessel able to withstand high temperature and pressure, just like a pressure cooker, but equipped with temperature, pressure measurement gauges) and raising its temperature and pressure above critical point of the alcohol. Above critical point, alcohol is a gas-like fluid which can be removed from the solid network without any surface tension. He named that transparent, low density solid 'aerogel'. These first aerogels synthesized by Kistler required very long (weeks) time consuming process. In 1970s, S. Teichner at University of Lyon, France, wanted to use these highly porous materials for storing rocket fuel. Looking at lengthy and laborious process he tried to develop a simpler and faster process to synthesize aerogels. He used tetra methylorthosilicate (TMOS) and hydrolyzed it in methanol, in presence of some catalyst. He could get gel in one step and exchange of solvent with alcohol was not required. He could synthesize aerogels in 1 day (Fig. 11.21).

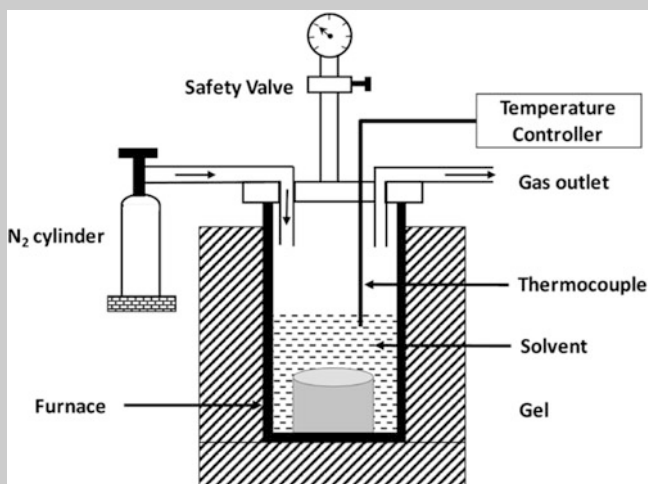
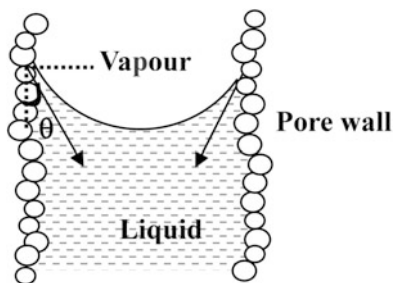


Fig. 11.21 Schematic diagram of an Autoclave (for the photograph, see Fig. 4.32)

These colloidal particles aggregate to form a continuous three dimensional network, which occupies total volume of the solution. The viscous semi-fluidic solid hence formed is known as 'gel'. Formation of gel from solution is dependent

Fig. 11.22 Surface tension at liquid vapour interface



on various parameters such as reactant concentration, temperature, pH value etc. In such gels, liquid is filled in the pores (empty space between the networks of solid particles). In simple evaporation, liquid and vapour coexist within the pore and surface tension of liquid at liquid vapour interface causes collapse of the network (as showed in Fig. 11.22). Magnitude of this interfacial force is given by Eq. (11.16).

$$P_s = \rho gh = \frac{2\sigma \cos \theta}{r} \quad (11.16)$$

where ρ is density of liquid, g – acceleration due to gravity, h – height of the capillary, σ – surface tension of the liquid, r – radius of the pore, θ is contact angle of the liquid with the pore wall.

Hence liquid has to be evaporated in such a fashion that network of particles does not collapse. One way to achieve this is by supercritical drying in which drying of gels is carried out above the critical point of liquid present in pores. Above critical point surface tension is zero. Whole liquid is present in vapour state and can be removed from gels without rupturing the network. In recent years it was realized that it is also possible to dry gels in ambient condition. This method is known as sub-critical drying. A comparison of sub-critical and super-critical drying procedure is shown in Fig. 11.23. In this case to overcome the interfacial pressure (surface tension), various techniques such as strengthening of gel by aging it in suitable solvent or use of some polymers or template to produce uniform size pores (to avoid pressure gradient while drying) etc. are used. Only drawback of this process is that it is time consuming and aerogels obtained by this method have higher density as compared to those dried by supercritical drying. Another way of drying is by freezing the liquid in the pores and sublimating by vacuum pumping. This method is known as freeze drying.

11.4.1 Types of Aerogels

Kistler had successfully demonstrated aerogel formation of various materials such as silica, alumina, nickel tartarate, stannic oxide, tungstic oxide, gelatin agar,

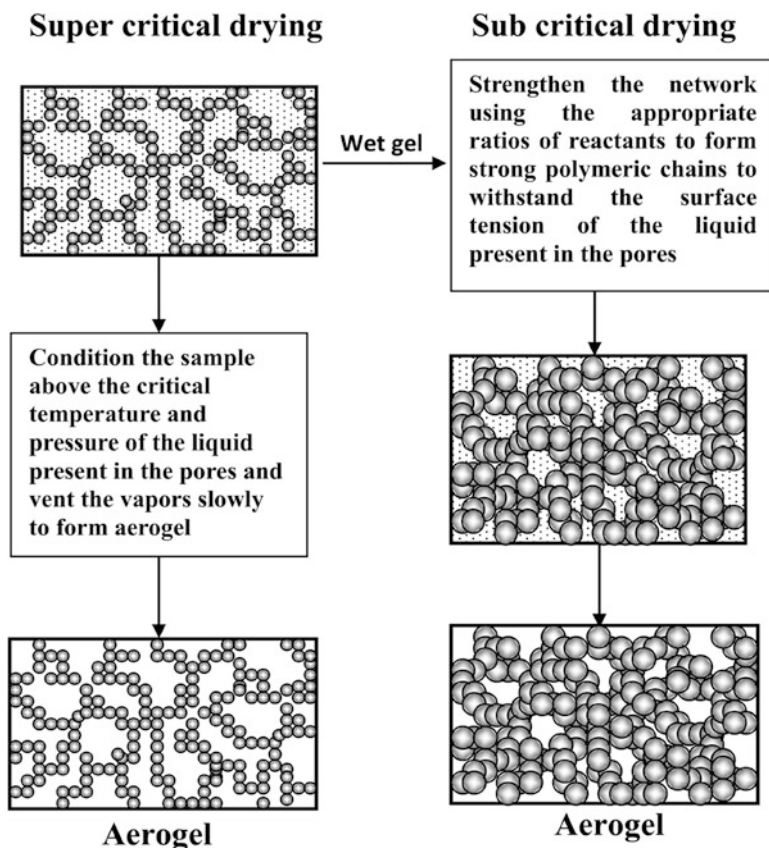


Fig. 11.23 Comparison of supercritical and sub critical drying

cellulose egg albumin etc. Inorganic aerogels such as SiO_2 , TiO_2 , ZrO_2 and Al_2O_3 have interesting properties and applications. It is also possible to synthesize mixed oxide aerogels such as $\text{SiO}_2\text{:TiO}_2$, $\text{Al}_2\text{O}_3\text{:SiO}_2$, $\text{SiO}_2\text{-ZrO}_2$ etc. Fascinating properties and applications of inorganic aerogels motivated also the development of organic aerogels such as RF (Resorcinol Formaldehyde) and MF (Melamin Furfural). They were synthesized by using different organic precursors such as Phenol-Furfural, Polyurethane, cellulose-toluene 2, 4-diisocyanate etc. RF aerogels were first synthesized by R.W. Pekala in United States in 1989. Organic aerogels can be heat treated with some specific heat cycle in inert atmosphere to convert them into carbon aerogels. It is indeed possible to dope aerogels with various metals such as gold, silver, manganese and even organic dyes for variety of applications. Among these materials silica aerogels are very commonly synthesized in many laboratories as well as commercially. Aerogels can be treated by various chemicals to make them hydrophilic or hydrophobic.

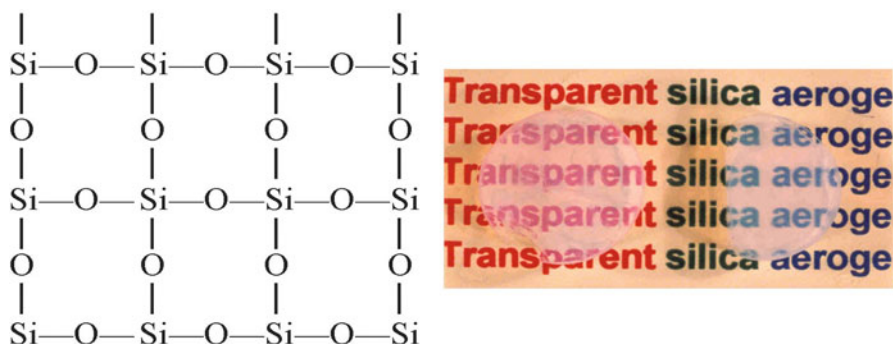
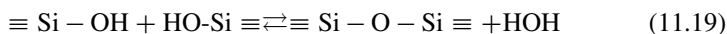
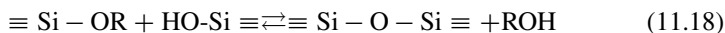
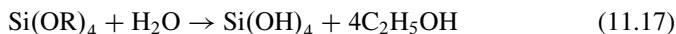


Fig. 11.24 Formation of 3D network of silica aerogel and a photograph showing transparent silica aerogel (diameter 2.5 cm and thickness 1.0 cm)

11.4.1.1 Silica Aerogel

Silica aerogels are historically the first ones that scientists have synthesized. They have been explored a lot due to their immense applications in various fields. Tetraethylorthosilicate (TEOS) or Tetramethylorthosilicate (TMOS) are typically used as primary reactants with ethanol or methanol used as solvents. Acid (HCl, HNO₃, NH₃) or base (NH₄OH, NH₄F) catalyst or combination of both is used for accelerating the reaction. Reaction consists of two steps viz. hydrolysis and polycondensation. For example, TEOS is first hydrolyzed using water. Two hydrolyzed molecules get condensed to form a dimer. The condensation process continues to form polycondensed silica gel, with Si-O-Si linkage as shown in Fig. 11.24. These gels can be further dried using supercritical drying. For carrying out this type of drying, autoclave is used. Ethanol or methanol can be used as solvent. Critical temperature for ethanol is 249°C and critical pressure is 78 bars. Depending upon reaction conditions aerogels can be transparent or opaque and density can typically vary between 0.01 and 0.8 g/cm³.



11.4.1.2 RF Aerogels and Carbon Aerogels

Inorganic aerogels such as silica have been explored a lot but organic aerogels such as RF (Resorcinol Formaldehyde) have been relatively less investigated. These aerogels are interesting, particularly because of the possibility to convert them

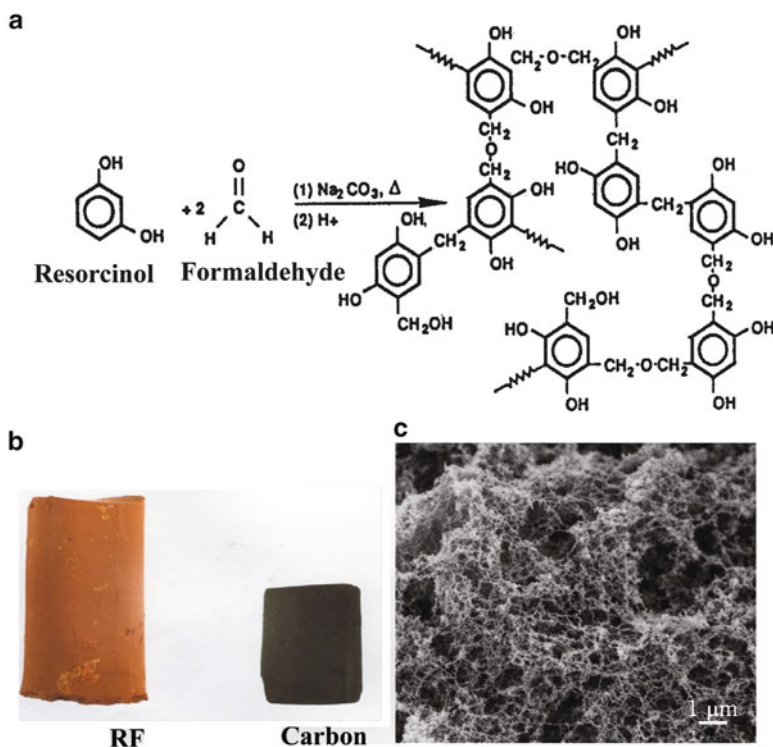


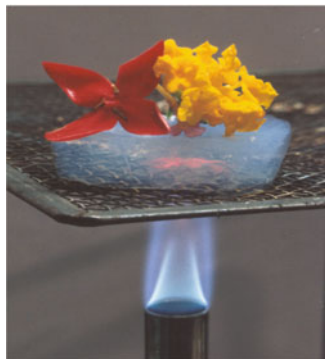
Fig. 11.25 (a) Formation of three dimensional network of RF aerogel; (b) photograph of RF and carbon aerogel and (c) SEM of RF aerogel

into carbon aerogel by pyrolysis (i.e. by heat treatment). Unlike silica aerogels they need high temperature (85°C) for gelation. For synthesis of these aerogels, resorcinol ($\text{C}_6\text{H}_5(\text{OH})_2$) and formaldehyde (CH_2O) are reacted in presence of sodium carbonate in aqueous medium. The reaction is shown in Fig. 11.25. The chemicals are mixed in a flask and stirred vigorously for 3 h and then left to themselves for 1 day at room temperature. The solution is then kept at 85°C for 6 days to form a gel. This gel can be dried supercritically or subcritically. To convert RF aerogel into carbon aerogel, it is necessary to heat RF aerogel under inert atmosphere at $\sim 1,100^\circ\text{C}$. A photograph showing RF and carbon aerogels along with SEM image of RF aerogel showing porous structure are as in Fig. 11.25.

Aerogels, however, due to their weak network of nanoparticles (bonded by Van der Waals interaction) have inferior mechanical properties. Therefore they can be easily broken with as small pressure as exerted by fingers. There are now attempts to increase the strength of aerogels by incorporating high strength nanomaterials like graphene or carbon nanotubes during gel formation process as strengths of graphene or carbon nanotubes are even superior than the steel of same dimension.

Table 11.3 Densities in g/cm^3 of some commonly known materials

Iron	Diamond	Glass	Graphite	C ₆₀	Paper	Charcoal	Aerogel	Air
7.87	3.5	2.4	2.3	1.7	0.7	0.57	0.003–0.8	0.0013

Fig. 11.26 Photograph of silica aerogel as best thermal insulator

11.4.2 Properties of Aerogels

Aerogels possess many interesting properties because of their highly porous structure. They are the lightest materials (comparison is made with densities of other materials in Table 11.3) man has ever synthesized having pore diameters of few tens of nanometers with 2–5 nm size particles. It is possible to have very high porosity (80–90 %), very low density and high surface area (500–1,500 m^2/g) for these aerogels. They have very low index of refraction (1–1.05) and very low speed of sound through them (~ 20 m/s). Young's modulus in these materials is quite low (10^6 – 10^7 N/m^2). Very low value of thermal conductivity (0.003 W/m.K) which is lower than most commonly used insulators make them best available thermal insulator. See Fig. 11.26.

11.4.3 Applications of Aerogels

For the decades since their first formation it was considered to be quite difficult to use aerogels commercially because of number of reasons like risky supercritical drying process, expensive synthesis procedure and their low fracture toughness etc. With the development of cost effective synthesis routes and technological developments and demands, interest in aerogels is increasing. They are nontoxic and biodegradable making them eco-friendly. Their unique properties such as very low thermal conductivity makes them best thermal insulators for various applications such as insulation of space vehicles, automobile engines etc. Aerogels have been commercially used in jackets and blankets to be used under extreme low temperature conditions. Large windows of aerogels are being used in houses and buildings

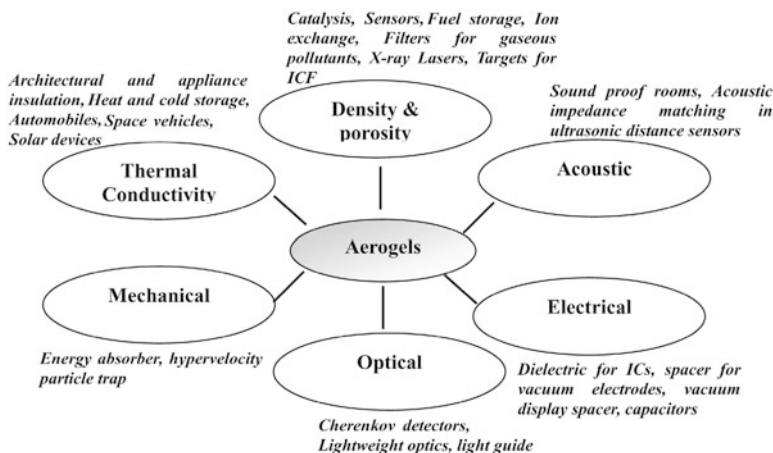


Fig. 11.27 Various applications of aerogels depending upon their properties

in some places to control the temperature in the interior. The transparency and heat insulation properties are used. Indeed it is possible to make aerogel-textile composites so as to make blankets, coats etc. by making them rollable. They can be obtained in the form of thin films, molds and tiny balls. Various applications of aerogels depending on their properties are shown in Fig. 11.27.

11.5 Zeolites

We discussed in the previous two sections two different types of porous materials viz. porous silicon and porous aerogels. Both the materials have disordered pores of irregular shapes and sizes. There is another class of materials known as zeolites, which has pores usually smaller (<2 nm) than aerogels or porous silicon but the pores are highly ordered. The zeolites have not only the pores of uniform size but are periodically arranged to have long range order. The material is crystalline. Some zeolites occur naturally. However due to their technological importance as catalysts and highly sorbent material, they are also synthesized on large scale.

The word 'zeolites' origins from 'zeo' and 'lithos' meaning 'boil' and 'stone' respectively. Boiling away water from some materials makes them porous zeolites. They have a crystalline structure with translational long range order for their unit cells in three dimensions. Unit cell itself is often the result of some quite complicated subunits. Most common zeolites have the building blocks of tetrahedral units of Si and Al connected together by Si-O-Al types of bonds. There can be Si-O-Si bonds also in zeolites but no Al-O-Al bonds exist. This implies that the number of silicon atoms will be always equal or larger than Al. The network of zeolites can be viewed as shown in Fig. 11.28.

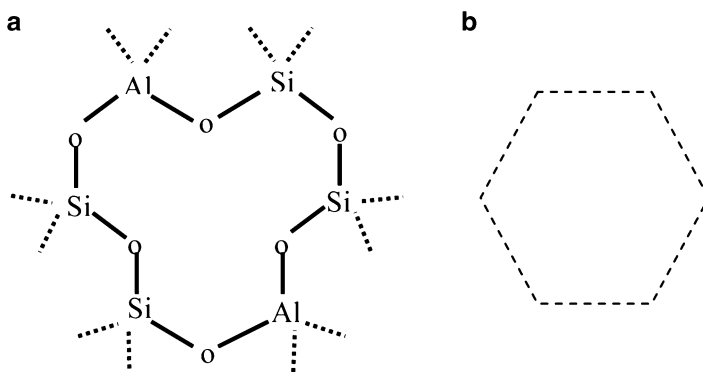


Fig. 11.28 (a) Al–Si–O subunit and (b) its schematic subunit

However, valence of silicon is 4 and that of aluminum is 3. This would mean that with Al substitution, there would be excess negative charge on the network, which needs to be balanced. This is done by adding some cations or electron acceptors in the material. A general formula of a zeolite can then be written as

$$[\text{Al}_x\text{Si}_y\text{O}_{2(x+y)}]\text{M}_x\text{zH}_2\text{O} \quad (11.20)$$

known as aluminosilicate. Besides this there can be some other tetrahedral units such as (GeO_4) , (ZnO_4) , (PO_4) etc. which can be used instead or in addition to Al. There are around 40 naturally occurring zeolites and much larger variety can be synthesized.

11.5.1 Synthesis of Zeolites

Synthesis of zeolites, however, is quite a delicate balance of additives as well as processing conditions. Often some organic additives of some ammonium or amines have to be added in some suitable quantities. Both metal cations and organic additives are believed to act as templates for building various structures. Variation of Si/Al, Si/H₂O, cation/Si and organic additives can give rise to numerous zeolites exceeding the number of natural zeolites. Some of the chemical species are responsible for achieving particular structure in zeolites. Such chemical species play the role of so called molecular templates. However particular structure can be obtained using more than one organic molecule and one molecule can create different structures for different precursors forming zeolites. Silicon and aluminium precursor solutions are hydrolyzed and polycondensation takes place similar to that in sol-gel process. Organic and inorganic reactants are added and a gel thus formed is hydrolyzed in an autoclave at high pressure (several tens or hundreds of bars) and temperature (usually 100–350 °C).

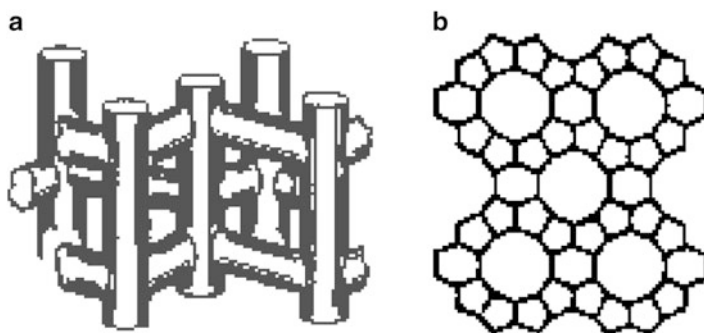


Fig. 11.29 Zeolite ZSM-5 (a) side view and (b) top view

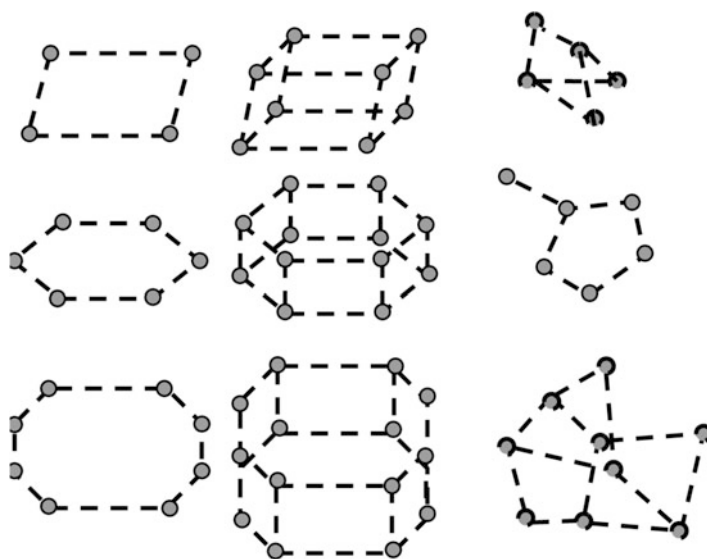


Fig. 11.30 Structural sub units in zeolites

Pore diameters can be controlled with a narrow size distribution in the size range of 0.3–2.0 nm. The pores can be connected in 1-D, 2-D or 3-D. Such porous materials find numerous applications. A widely used ZSM-5 zeolite structure is as shown in Fig. 11.29.

11.5.2 Properties of Zeolites

Zeolites are crystalline, highly ordered porous materials. The pore sizes can be precisely controlled. As illustrated in Fig. 11.30, there can be a large number of subunits forming zeolites of different constituents.

Pores in zeolites are useful to carry out different reactions. Zeolites are useful catalysts or sorption material. By controlling the pore sizes to desired size, it is possible to synthesize and organize nanoparticles inside zeolites.

11.6 Porosity Through Templates

Porosity in nanomaterials can be obtained through various procedures. Here we shall discuss the method in which micelles are used as the templates. Another method involves using core-shell particles in which core is partially or completely removed and the shell becomes porous in the whole process. This part we shall discuss in the core-shell particles.

11.6.1 *Micelles as Templates*

This also is a good example of self assembled materials. Large (2–10 nm) and ordered pores in three dimensional solids can be obtained using micelles as the templates. However major difference between zeolites and M41S and its family as they are known, is that the pore walls are amorphous in this case. These are only synthetic materials with no naturally occurring material of this type. They were first synthesized in 1992 by some scientists in Mobil company. These are synthesized as 1-D ordered (MCM41), 2-D ordered lamellar (MCM-50) and 3-D ordered (MCM48) arrangement of pores.

11.6.1.1 Synthesis of Ordered Porous Materials Using Micelles

Very interesting, novel route of synthesis has been adopted to synthesize these materials. Amphiphilic organic chain molecules, water, inorganic silica precursor and some catalyst are used in making these materials. Reaction temperature can vary over a wide range from -14 to 180 °C depending upon the materials of interest. However the underlying idea is as follows. As discussed in Chap. 3, when amphiphilic molecules are mixed in water, depending upon the molecules above the critical concentration, one can spontaneously obtain micelles. With increasing concentration, one can even obtain different structures like hexagonal, cubic or lamellar at certain temperature. If silicates are condensed on preformed structures of micelles, it is possible to obtain ordered composites as shown in Fig. 11.31.

Alternatively, the formation of observed periodic pore structure can be viewed as follows. Before the self assembly of micelles takes place forming some ordered structures, each micelle may acquire a silicate shell around it and then form the ordered structure. The assembly always occurs irrespective of the stage at which the silicate precursor is added to the micellar solution. Once the reaction is over, micelles can be removed by plasma etching, calcination or solvent exchange processes.

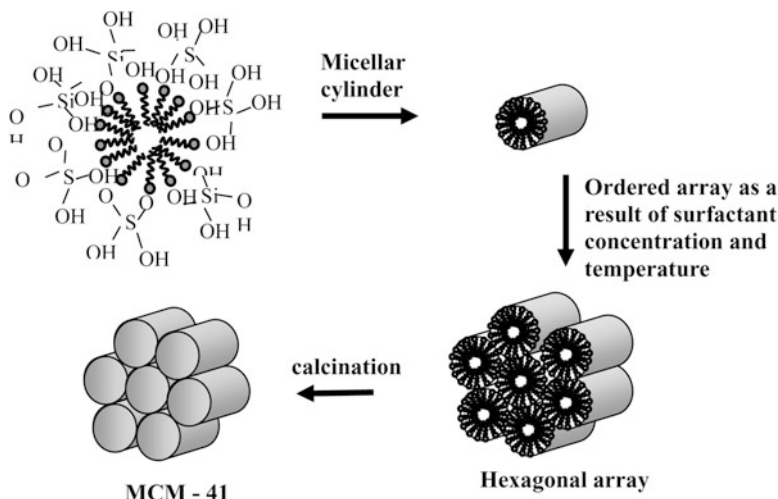


Fig. 11.31 Micelles covered with silica

Such ordered pore structures themselves are nanomaterials and can be also used as templates to synthesize nanoparticles in their cavities.

11.6.2 Metal Organic Frameworks (MOF)

Metal organic frameworks are under intense investigations from the point of view of finding new MOFs and their applications. MOF frameworks are co-ordination polymers in which metal ions make a one-dimensional, two-dimensional or three-dimensional pore network through organic molecules. They resemble zeolites to some extent but they are different from zeolites. Zeolites are mainly aluminosilicates. The densities of zeolites are quite high and can be $\sim 50 \text{ g/cm}^3$ as against the densities of MOFs which can be $\sim 0.5 \text{ g/cm}^3$ viz. two orders of magnitude smaller than zeolites. The pore sizes in MOFs can be as low as $\sim 5 \text{ nm}$ diameter. This results into extremely large internal surface area ($\sim 6,000 \text{ m}^2/\text{g}$). In MOFs this large surface area is the most important parameter as this is important in deciding the choice of a material as the potential candidate for H_2 storage application or drug delivery. Some of the organic molecules which have been used in the synthesis of MOFs are oxalic acid ($\text{HOOC}-\text{COOH}$), malonic acid ($\text{HOOC}-\text{CH}_2-\text{COOH}$), succinic acid ($\text{HOOC}-(\text{CH}_2)_2-\text{COOH}$), glutaric acid ($\text{HOOC}-(\text{CH}_2)_3-\text{COOH}$), phthalic acid ($\text{C}_6\text{H}_4(\text{COOH})_2$), citric acid ($\text{HOOCCH}_2\text{C}(\text{OH})(\text{COOH})\text{CH}_2(\text{COOH})$), trimeric acid ($\text{C}_9\text{H}_6\text{O}_6$), 1,2,3 triazole ($\text{C}_2\text{H}_3\text{N}_3$) and squaric acid ($\text{C}_2\text{H}_2\text{O}_4$). Some of the metal ions used are Mg, Fe, Zn, Mn and Cu.

The synthesis of MOFs is carried out mostly by hydrothermal or solvothermal and mechanochemical techniques. Hydro or solvothermal reaction of precursor

chemicals is carried out in an autoclave at high temperature and high pressure similar to that in the synthesis of zeolites. In mechanochemical synthesis metal acetates are mixed with organic molecules and ground in a ball mill.

The biodegradability of MOFs is excellent. Their functionality makes them useful in H_2 and other gases like CO_2 , O_2 storage. Additionally MOFs can be useful in catalysis, drug delivery, data communication and replacement of liquid crystals which are under investigations.

11.7 Core-Shell Particles

Core-shell particles form a novel class of nanocomposite materials in which a thin layer of nanometer size is coated on another material by some specialized procedure. The core can be just a nanoparticle (few nanometers to tens of nanometers) with a nanometer thick coating or it can be a large core (few tens to hundreds of nanometer diameter) with nanometer thick coating as schematically shown in Fig. 11.32. The properties of core-shell particles are different from core or shell material. Their *properties depend usually upon core to shell ratio*. These particles are synthesized for a variety of purposes like providing chemical stability to colloids, enhancing luminescent properties, engineering band structures, sensors, drug delivery etc. These materials can be of economic interest also as precious materials can be

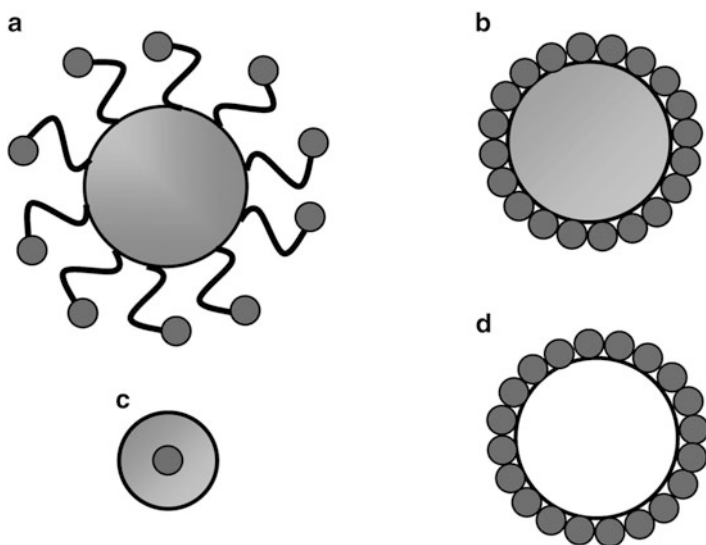


Fig. 11.32 Variety of core-shell particles: (a) Surface modified core particles anchored with shell particles, (b) Smooth coating of dielectric core with shell, (c) Encapsulation of very small particles with dielectric material and (d) Quantum bubble

deposited on inexpensive cores. Core particles of different morphologies such as rods, wires, tubes, rings, cubes etc. also can be coated with thin shell to get desired morphology in core shell structures.

Core-shell materials can be synthesized practically with all the materials, like semiconductors, metals and insulators. Dielectric materials such as silica and polystyrene are popular materials to use as core because they are soluble in water and hence can be useful in biological applications. Core-shell particles can be synthesized using variety of combinations such as dielectric-metal, dielectric-semiconductor, dielectric-dielectric, semiconductor-metal, metal-metal, semiconductor-semiconductor, semiconductor-dielectric, metal-dielectric, dye-dielectric, dielectric-biomolecules etc. Although core shell particles have novel properties, these can be further assembled and utilized for creation of another class of novel materials like colloidal crystal or quantum bubbles (i.e. hollow spheres with thin shells). It is indeed possible to create novel core shell structures having multishells and tuning optical properties from visible to infrared region of the electromagnetic spectrum.

Synthesis of core-shell particles requires highly controlled and sensitive synthesis protocols to ensure complete coverage of core particles with shell. There are various methods to fabricate core-shell structures which involve precipitation, polymerization, micro emulsion, reverse micelle sol-gel condensation etc. Although these methods themselves may appear to be simple, it is rather difficult to control the thickness and homogeneity of the coating. If reaction is not controlled properly, eventually it leads to aggregation of core particles, formation of separate particles of shell material or incomplete coverage.

Here we shall discuss some silica based core-shell particles as an example. Preparation of core-shell particles is a multi-step synthesis procedure. One can make coating of silica on nanoparticles or grow silica particles of large size and then anchor nanoparticles or coat thin shell around, of few nm thickness. As mentioned earlier one can use metals, semiconductors or any other dielectric material as shell or core with silica.

11.7.1 Synthesis of Silica Cores

Silica is a very widely used material to form core-shell particles because of its stability against coagulation. It is also chemically inert and optically transparent. For various purposes it is desirable that particles remain well dispersed in the medium which can be achieved by coating silica on them to form an encapsulating shell.

Silica particles having very good monodispersity can be synthesized by what is known as Stöber method. This method involves hydrolysis and successive condensation of TEOS (tetraethylorthosilicate $\text{Si}(\text{C}_2\text{H}_5\text{O})_4$) in alcoholic medium in presence of ammonium hydroxide (NH_4OH) as a catalyst. By varying relative concentration of TEOS to solvent (dilution) and amount of catalyst, one can synthesize these particles in various sizes ranging from 50 nm to 1 μm .

11.7.1.1 Synthesis of Metal, Semiconductor or Insulator Particles (<20 nm)

Synthesis of metal, semiconductor and insulators with very small size (<20 nm) and small size distribution can be performed using some of the chemical techniques described in Chap. 4 (as well as Chap. 14).

11.7.2 Core-Shell Assemblies

Core-shell assembly is a multi-step process in which core and shell particles are separately synthesized and then shell particles are anchored on cores by specialized procedures or shell is directly synthesized on preformed cores.

In the first method, surface of core particles is often modified with a surfactant or bifunctional molecules to enhance the coverage of shell material on its surface. Surface of core particles such as silica can be modified using bifunctional organic molecules such as APS (3-aminopropyltriethoxysilane). APS molecule has an -OH group at one end while -NH group at the other end. APS forms a covalent bond with silica particle through -OH group and -NH group is available for interaction. Some nanoparticles having affinity for nitrogen can be attached with silica particles through NH group. There are several other modifiers such as APTMS (3-aminopropyltrimethoxysilane) and AEAPTMS (N-(2-aminoethyl)-3-aminopropyltrimethoxysilane) which produces surface terminated with amine group, MPTMS (3-mercaptopropyltrimethoxysilane) which produces surface terminated with thiols, DPPETES (2-(Diphenylphosphino)ethyltriethoxysilane) leaves the surface terminated by diphenylphosphine group and PTMS (propyltrimethoxysilane) gives surface terminated with methyl group. It is thus possible to make variety of core-shell materials with dielectric cores such as silica, titania and polystyrene can be anchored with different materials.

In another approach, synthesis of shell particles is carried out in presence of preformed cores (may be single nanoparticles or many nanoparticles). These core particles act as nuclei and shell material hydrolyses and gets condensed on cores forming core-shell particles. Reactant concentrations and amount of added core particles play important role in deciding shell thickness.

In other technique, which is known as layer by layer technique, alternate layers of anionic particles and cationic polymer are deposited on surface-modified template molecule by heterocoagulation.

Successive removal of core material (either by calcination or dissolution in suitable solvents such as toluene or HF) yields hollow particles consisting of shell material, popularly known as quantum bubble. This is known as template based synthesis of porous or hollow particles. For example SiO_2 nanoparticles are coated with SnO_2 thin film and etched mildly in controlled way using HF. Then with repeated deposition of SnO_2 and mild etching, core SiO_2 which mixes with SnO_2

layers is slowly etched out. This leaves SnO_2 particles with hollow core as well as porous SnO_2 shell. The process is known as galvanic replacement (of SiO_2 by SnO_2 in this case).

Selection of suitable pair for core and shell assemblies requires understanding of individual properties of core and shell materials. Core particle should withstand the process used for the coating of shell material. Core and shell particles should not interdiffuse and surface energies of core and shell particles must be similar so that probability of heterogeneous nucleation is more than homogeneous nucleation.

11.7.3 Properties of Core-Shell Particles

Coating of colloidal particles with shell offers most simple and versatile way of modifying its surface chemical, reactive, optical, magnetic and catalytic properties. Silica particles coated with gold shell have been studied, which show changes in the Surface Plasmon Resonance (SPR) position of which depends upon core/shell optical properties. CdSe nanoparticles (<5 nm) coated with ZnO, CdS/ZnTe and CdTe nanoparticles coated with CdSe also have been studied for enhancement in luminescence of core nanoparticle. Use of higher band gap material as a shell for lower band gap material as a core increases the probability of photoexcited electron of core particle being trapped and increases the photoluminescence of core particle. Magnetic particles of iron oxide also have been coated with dye incorporated silica shell. Such particles show magnetic properties arising from core as well as luminescent optical properties arising from shell. Thus functional materials with novel properties can be synthesized by using various combinations of core shell materials.

11.8 Metamaterials

The literary meaning of ‘metamaterials’ is ‘beyond Nature or Supernatural’. These materials have been fabricated artificially as a consequence of advances in the fabrication of nanomaterials in the desired shapes and sizes as well as improved understanding of materials, particularly their interaction with the electromagnetic waves. Metamaterials are also known as ‘negative refractive index materials’ or ‘left handed materials’.

These materials are subwavelength periodic artificial structures capable of producing negative refraction. It has been realized after about year 2000 that having such materials would help making ‘superlenses’, band pass filters, novel photonic devices and even cloaking (invisibility of objects).

The concept of negative refractive index is not really new in physics. Lamb and Schuster had come across it in their theoretical work. However they had dismissed the idea thinking that it would not have any physical meaning or application.

Fig. 11.33 Illustration of usual (positive) refraction of a ray and negative refraction of the ray when ray of light traverses from medium of low density into higher density medium

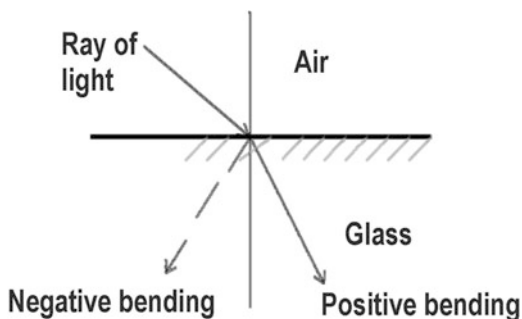
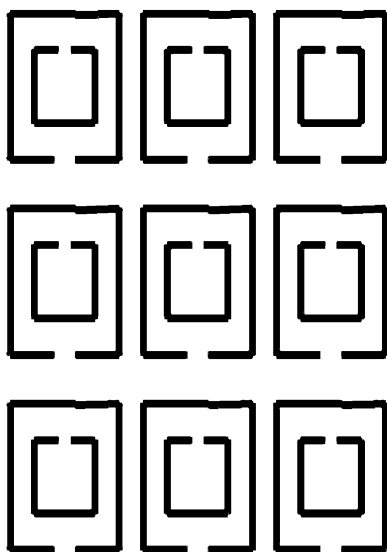


Fig. 11.34 Split resonators. Such structures are obtained in gold and form three dimensional arrays



In 1967, Victor Veselago revisited the concept and showed that negative refractive index was possible provided dielectric constant ϵ and magnetic permeability μ were simultaneously negative. Dielectric constant ϵ is frequency dependent and becomes negative at certain frequencies (this was seen in the discussion on plasmonic materials, Sect. 8.4). However, permeability μ can be achieved negative with some artificial structures (usually some resonating circuits are required). Around 2002, several structures were theoretically suggested and experimentalists proved that indeed it is possible to obtain negative refraction, as shown in Fig. 11.33.

One such resonator arrangement section is illustrated in Fig. 11.34. Usually these are lithographically constructed by electron or ion beams in gold as gold does not get easily oxidized and has negative dielectric constant in the large frequency range.

With considerable research progress in this branch, it has been found that negative refractive index materials can be Double Negative (DNG) Refractive Index materials in which both dielectric constant ϵ and permeability μ are negative or only ϵ (epsilon) is negative (ENG) or only μ (mu) is negative (MNG).

Superlenses and cloaking also have been demonstrated. In superlenses the focussing can be achieved without the usual distortions due to curvature and other aberrations of the optical lenses. Cloaking has been demonstrated in the microwave region. In object kept inside the 'cloak' is invisible to microwaves as they are completely reflected. Based on this idea even 'plasmonic' cloaking has been proposed which would make the object invisible to particular wavelength but seen at other wavelengths.

11.9 Bioinspired Materials

We had discussed in Chap. 5 that the nature has mostly taken the 'bottom up' approach in building the animal and plant kingdom. Once the nanostructures are formed they may assemble (self assembly) into hierarchical structures to give functionalities. The scientists after getting the powerful microscopes at hand have gone into the details and found that many interesting phenomenon like cleanliness of water leaves, crawling of lizards on walls without falling or beautiful colours of peacock feather or butterflies or fish, birds are the results of 'Nature's Nanotechnology'. In this section we will try to understand some of the effects and the reasons behind them.

11.9.1 *Lotus Effect (Self Cleaning)*

In recent times there has been new understanding about how the hydrophobic (water hating) and hydrophilic (water loving) surfaces work. This effect has been there for millions of years and is now recognized as the 'lotus effect'. Lotus is a beautiful flower known to all civilizations and is being praised over generations in many cultures, religions and languages. How it appears beautiful and its leaves do not become dirty even by staying in muddy water without letting water stick to it has been a point of wonder. However, it has been investigated and found that this is due to the hierarchical micro-nano structure of the lotus leaf and many other leaves as well as body of animals.

A lotus leaf when observed under the microscope exhibits some bumps of micrometre size. The bumps are decorated with nanometer sized structure. This can help a water drop stay on the lotus leaf without spreading on it. The waxy substance on the lotus leaf also helps the water drop to form the droplet and not spread on the surface. When the water drop rolls on the leaf, it collects the dirt on the surface and makes the surface clean. Lotus effect can be understood from Fig. 11.35.

The phenomenon of wettability in lotus leaf is also common in many other plant leaves, butterfly wings, skin of fish and many others in plant and animal kingdom. It is understood in terms of hydrophobicity/hydrophilicity and the contact angle a water drop makes with the surface. A hydrophilic surface is water loving and

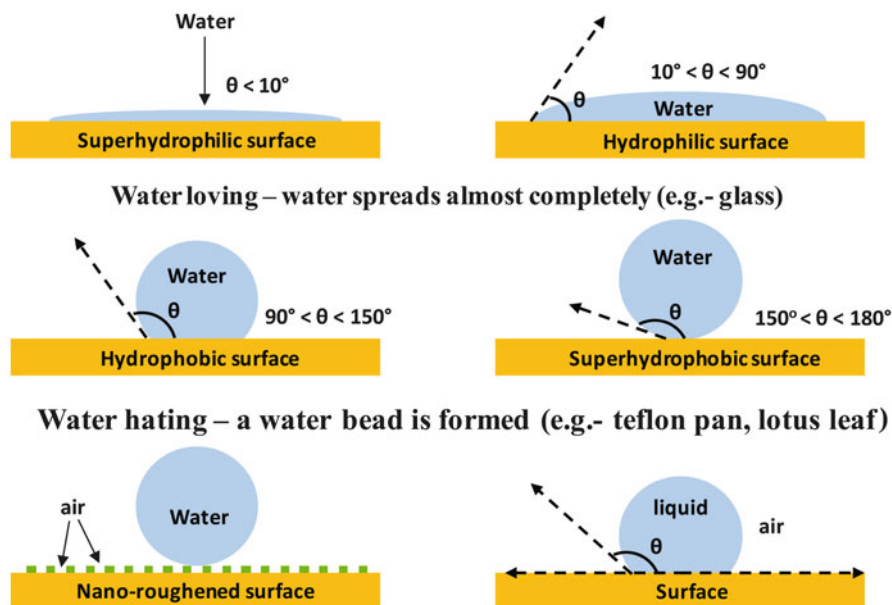


Fig. 11.35 Formation of water drop on the surface

hydrophobic surface repels water. Depending upon the contact angle of water drop, one can quantify the amount of hydrophilicity or hydrophobicity. As illustrated in Fig. 11.35, water contact angle is defined as the angle between solid surface and the liquid-surface. Contact angle depends upon surface roughness, surface condition and surface material. Interaction between liquid and solid surface is very important. When the contact angle is less than 90° the surface is known as the hydrophilic and if more than 90° it would be hydrophobic. If the angle is above $\sim 150^\circ$ then it is superhydrophobic and if less than $\sim 20^\circ$ then it is superhydrophilic. The angles which distinguish between superhydrophilic and hydrophilic or hydrophobic and superhydrophobic are not very precise and could vary within 5° – 10° .

Scientists have tried now to make use of this phenomenon to artificially obtain the hydrophilic/hydrophobic surfaces in various materials of various shapes and sizes. One can even make coatings which would not only make a surface just hydrophobic or hydrophilic but also can change the nature in the controlled way using external stimuli such as application of an electric field or irradiation with UV light or X-rays.

Another interesting and more dramatic natural phenomena, which often goes unnoticed is known as 'petal effect'. It can be easily seen in common flowers e.g. a rose. If water is sprinkled on flowers, very tiny droplets can be observed. These droplets cling to the petals so strongly that even if the flower is turned upside down, the droplets do not fall under gravity. This effect is due to micropapillae of $\sim 16 \mu\text{m}$ diameter and $\sim 7 \mu\text{m}$ height with micropapillae folds $\sim 730 \text{ nm}$ size.

11.9.2 Gecko Effect (Adhesive Materials)

Somewhat similar interesting phenomenon observed in nature over centuries and wondered about is so-called ‘Geko (family name for lizard) effect’ or crawling of a lizard on wall. How does a lizard balance its weight? Does it secrete any fluid to stick or what else? Now the scientists believe that the peculiar construction of the tips of its feet on which millions of nano hair exist which help lizard in crawling. Interestingly the force that each hair strand exerts between wall and itself is just weak Van der Waals force. However, due to millions of hairs total force is sufficient to balance its weight and remove the force in controlled way to crawl. This effect may be useful to make robots which would climb the walls or do some scientific operations without manual aid. Some simple application of geko effect would be a stick tape which can be used number of times without falling like a lizard. One may also be perhaps able to make seals without nuts and bolts.

Basically understanding even natural phenomenon around us can lead to new discoveries and progress in science and technology.

Further Reading

F. Caruso, *Adv. Mater.* **13**, 11 (2001)

R.T. Collins, P.M. Faushet, M.A. Tischler, Porous silicon: from luminescence to LEDs. *Phys. Today* **24** (1997)

L.M. Liz-Marzan, M.A. Correa-Duarte, I. Pastoriza-Santos, P. Mulvaney, T. Ung, M. Giersig, N.A. Kotov, Nanostructured materials, micelles and colloids, in *Hand Book of Surfaces and Interfaces of Materials*, ed. by H.S. Nalwa, vol. 3 (Academic Press, San Diego, 2001)

J.B. Pendry, D.R. Smith, Reversing light with negative refraction. *Phys. Today* **57**(6), 37 (2004)

K. Tanaka, T. Yamabe, K. Fukui, *The Science and Technology of Carbon Nanotubes* (Elsevier, Oxford/Amsterdam, 1999)

S. Ulrich, H. Nicola, *Synthesis of Inorganic Materials*, 2nd edn. (Wiley-VCH, New York/Weinheim, 2005)

Y. Xia, B. Gates, Y. Yin, Y. Lu, *Adv. Mater.* **12**, 693 (2000)

Chapter 12

Applications

12.1 Introduction

The ability of materials to dramatically change their properties at nanoscale has opened up the possibility of making new devices, instruments and consumer goods to function in a much better way than was possible earlier. We have seen in Chaps. 10 and 11 that nanomaterials have enabled us to design new products which were not possible using bulk materials. Rapid progress in the synthesis and understanding of nanomaterials in just a few years has led them to enter the world market in a big way. Figure 12.1 shows an overview of various fields in which nanomaterials have entered or are about to enter. In this chapter we shall briefly discuss some of these applications.

12.2 Energy

Nanotechnology is expected to play an important role in the field of ‘energy’ due to the availability of high efficiency and cost effectiveness of nanomaterials. We all know that natural energy resources like coal, oil and natural gas used in transportation, communication, agriculture, industry, houses and many other human activities are limited and depleting very fast. Future generations will have to look for alternative sustainable energy sources like nuclear, geothermal, wind, solar energy or hydrogen-based fuel cells to satisfy their requirements.

The world energy consumption is ~ 13 TW per annum. The distribution of resources is roughly as follows: 4.66 TW from oil, 2.98 TW from coal, 2.89 TW from gas, 0.92 TW from nuclear, 1.24 TW from biomass, 0.28 from hydro and 0.28 TW from other renewable sources like solar, wind or geothermal. This can be seen in Fig. 12.2 as percentage distribution.

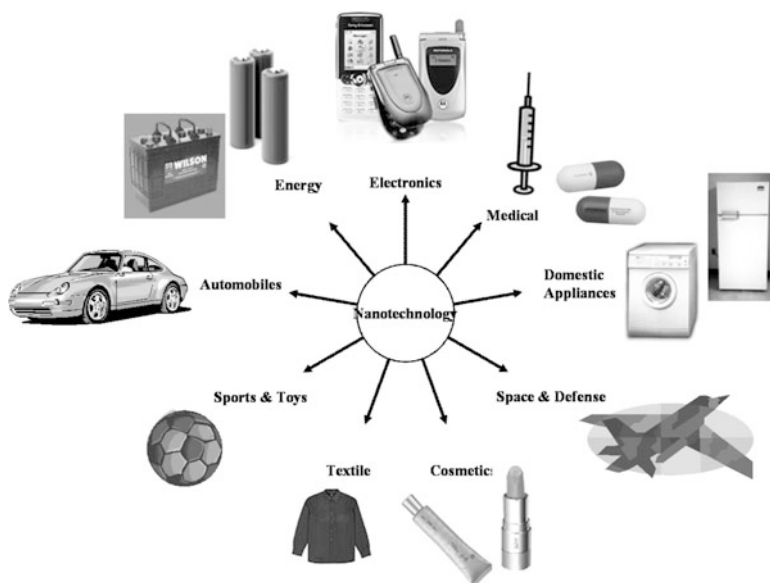
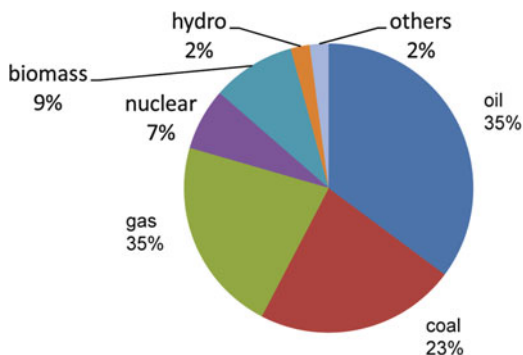


Fig. 12.1 Applications of nanotechnology

Fig. 12.2 Pie diagram showing the energy derived from different sources



There is a considerable amount of research going on to tap hydrogen fuel by splitting water (H_2O) using sunlight in the presence of nanomaterials (photocatalysts). Available hydrogen can indeed become a good source of fuel for automobiles and other transportation purposes. However storing hydrogen is not an easy job as it can readily catch fire. Material like carbon nanotubes, a special class of nanomaterials, is being investigated for its potential use as hydrogen storage material. Current cost of carbon nanotubes is very high, but scientists are trying to find inexpensive ways of making them on large quantities, which would help in future to use hydrogen

fuel without risk. There are also attempts going on to increase the efficiency of solar cells for energy production using nanoparticles.

Numerous gadgets like laptops, cellular phones, cordless phones, portable radios, CD players, calculators etc. need rechargeable, light weight batteries or 'cells'. Presently, the batteries for such gadgets need to be either replaced with new ones or recharged quite frequently due to their low energy density or storage capacity. Attempts are being made to increase their energy density, for example by replacing the electrode materials. Some metal hydride nanoparticles like nickel hydrides or high surface area, ultra light weight materials like aerogels are found to be better options than the conventional materials in improved batteries.

Although there have been numerous efforts to tap various alternative energy sources than coal, oil and gas, main hurdle is the cost effectiveness. The solar energy, although spread over a large area of the earth's surface, is so much abundant on the earth that if one could harness it, the energy received on earth in one hour from the Sun is sufficient to satisfy the needs of entire human population on the earth for 1 year. About 120,000 TW of radiation from the Sun strikes the earth every year. However, so far all the attempts to tap the solar power (or other sources like wind, hydro etc.) have not turned out to be cost effective as compared to oil, gas or coal—the conventional energy sources (Box 12.1).

Box 12.1: History of Solar Cells

- **1839** – French scientist A.E. Becquerel understood the photovoltaic effect i.e. generation of voltage in a material due to incident photons.
- **1883** – Charls Fritts fabricated the first solar cell by using gold contacts with selenium. Efficiency of $\sim 1\%$ was achieved.
- **1946** – Russel Ohl patented junction solar cell.
- **1954** – D.M. Chapin, C.S. Fuller and G.L. Pearson made diffused p-n junction solar cell in silicon.
- **1991** – Grätzel developed dye sensitized titania nanoparticles based solar cell with $\sim 11\%$ efficiency. This marks the search for nanomaterials based solar cells other than silicon.

Globally the scientists and technologists are making efforts with domestic and international financial support to tap alternative sources of energy. In many nations there is a move against the establishment of new nuclear plants and even shut down the existing nuclear plants due to risk of fatal accidents. Nuclear waste, leakages and accidents can be dangerous for generations. Therefore in future we expect that the use of nuclear power plants will be on the decline.

Besides their depletion, the conventional energy sources also turn out to be problematic due to emission of carbon dioxide. Carbon dioxide once emitted does not dissociate easily and its equilibration by biomass and ocean takes place after several years. It is realized now that this has resulted into global warming, causing the melting of polar caps or huge glaciers at the north and the south poles of the earth. Glaciers while melting also release methane adding to the pollution. Already the pH of oceans is changing and at some places has resulted into de-colouration of the corals. Temperature rise by $4-5^{\circ}$ would extinguish many species on the earth. This has already started disturbing the eco-balance and the whole world is quite concerned about it due to the obvious reasons. However, the modern life style also depends upon the increasing consumption of power. It is speculated that the 13 TW annual consumption of power by the earth population (earth population also is increasing) would increase the demand to 20 TW by the year 2050. From where is this energy going to come? One needs to find clean, safe sources of energy which are sustainable as well as cost effective.

Scientists are seriously considering solar energy as one of their best options (contribution to the energy generation by wind, geothermal, ocean tides etc. is not expected to be very high/cost effective). Solar energy spectrum extends from UV to infra red. One way to tap energy is to collect the thermal energy in concentrated manner to heat water, cooking purposes and so on. This is being done quite profitably in many places now. However, it is quite a marginal relief to the energy stress on the global scale. Another option is to use photovoltaic panels and store the energy in batteries and use it whenever required like usual switching of electricity. Advantage with photovoltaic panels is that they can work in remote places like in villages in isolation or in urban areas. The surplus energy can be transferred to the electric grid. Such installations have proved to be cost effective even with the crystalline silicon photovoltaic panels which are a bit expensive.

Currently, the photovoltaic (solar cells) are divided as first generation, second generation and the third generation. First generation solar cells are made using single crystalline silicon wafers and are nothing but p-n diodes. Although the efficiency of these cells in some configurations has reached as high as $\sim 40\%$, they are too expensive. The second generation of solar cells is based on the thin films of crystalline silicon, amorphous silicon, CuInSe_2 based cells and many other thin film solar cells. They are cost effective to some extent but not efficient. The third generation of solar cells is based on the nanocrystalline materials and technologists are hopeful that it would be efficient as well as cost effective.

Currently two types of solar cells viz. inorganic and organic are competing with each other in terms of research. Attempts are being made to improve the performance (efficiency increase) and reduce the cost of production/installation. Some hybrid designs are also being demonstrated.

Besides solar cells, efforts are made to improve the fuel cells. In the following sections we will discuss some of the photovoltaic solar cells, fuel cells and issues related to hydrogen generation and storage for energy supply.

12.2.1 Dye Sensitized Photovoltaic Solar Cell (Grätzel Cell)

The first solar cell based on nanomaterials is probably the dye-sensitized Grätzel cell, named after a scientist from Switzerland. Instead of using just one p-n junction like in the first or second generation solar cells it uses multiple junctions at each nanoparticle of TiO_2 and a dye molecule. Although the efficiency of such a cell is only $\sim 11\%$ (claims upto $\sim 15\text{--}16\%$ also are made), the cost effectiveness has been very attractive and hoped that even with lower efficiency the energy problems can be solved to some extent (Boxes 12.2 and 12.3). Besides, Grätzel cell has generated a new concept of utilizing efficient and cost effective nanoparticles for photovoltaics. Various hybrid solar cells using nanoparticles, CNTs, graphene, polymers and other materials are under intense research. It has been shown that there is no problem for the nanoparticle solar cells to reach the thermodynamic limit of $\sim 90\%$ efficiency.

Box 12.2: Solar Spectrum

Energy received by the earth is in the form of an electromagnetic radiation (Fig. 12.3), in the range of ~ 250 to $\sim 2,300$ nm. Energy radiation from the sun results due to nuclear fusion as hydrogen gets converted into helium. In this process $\sim 6 \times 10^{11}$ kg hydrogen gets converted per second ($E = mc^2$, Einstein equation) to helium with $\sim 4 \times 10^3$ kg mass difference. This amounts to $\sim 4 \times 10^{20}$ J. Assuming that the mass of the Sun is 2×10^{30} kg, estimated life of Sun is more than 10 billion years. Thus practically solar energy is inexhaustible.

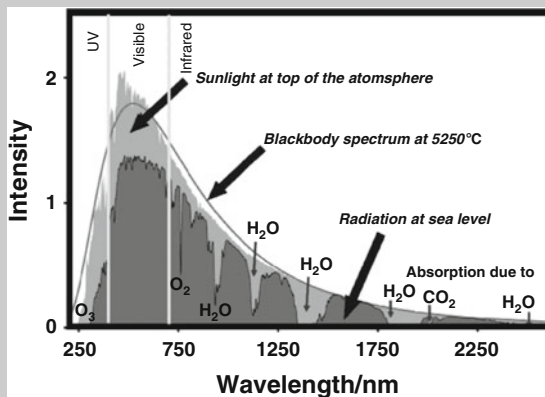


Fig. 12.3 The solar spectrum

(continued)

Box 12.2 (continued)

Solar energy spectrum, although continuous, is absorbed/scattered in certain regions of wavelengths due to some molecules like water, ozone or dust particles before reaching the surface of the earth. Therefore it is necessary to consider the solar intensity in free space, where there is no air. Usually intensity is measured (see Fig. 12.4) at different ‘air mass’ and referred to as AM0, AM1 or AM2. The angle between Sun and the Zenith is used to determine the atmospheric length the rays have to travel relative to the minimum path length. Minimum path length is when Sun is overhead at a point on the earth (at the sea level).

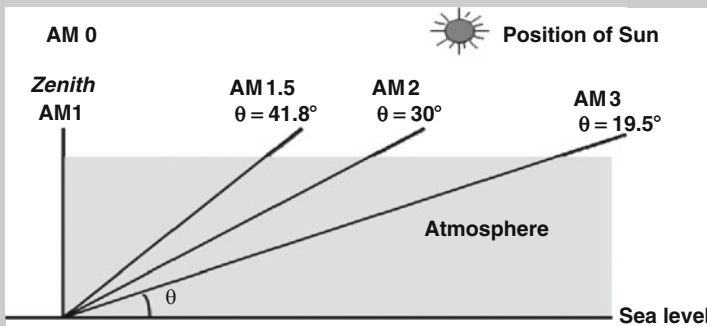


Fig. 12.4 Path of solar rays through atmosphere

Box 12.3: Solar Cells

Solar cells can obtain the energy from the sunlight to produce electricity. It is a sustainable, clean form of energy without consumption of any fuel. Large area solar cell panels can be constructed in order to produce large power or with smaller modules power can be locally delivered. Usually solar cell based panels or devices can last for 30–40 years without much maintenance.

Working principle of a solar cell can be understood as follows. Solar cell in its simplest form is a p-n diode (see Fig. 12.5) in a single material. The diode is formed by diffusing n-type dopant from one side of a p-type semiconductor or diffusing p-type dopant in n-type semiconductor. Electrons from n-side diffuse to p-type and combine with the holes and the built up charges on both

(continued)

Box 12.3 (continued)

the sides of the junction sets up an electric field by which a depletion region is formed which is devoid of any charges.

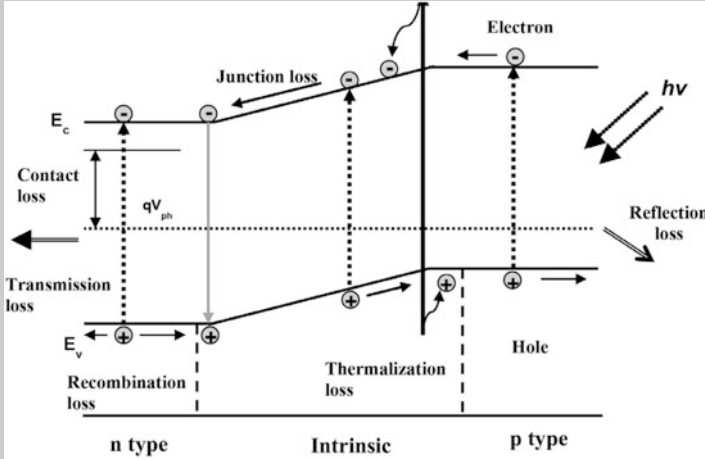


Fig. 12.5 Principle of a conventional photovoltaic solar cell

When a photon of energy larger than the energy gap is incident on the cell, electrons are excited to the conduction band and various processes as depicted in Fig. 12.5 take place. It can be seen that there are different sources of losses of photons and charges. There are many approaches adopted to minimize these losses using various materials and architectures.

Efficiency of a solar cell is defined as

$$\eta = \frac{I_{sc} V_{oc} FF}{P_s} \times 100 \quad (12.1)$$

where

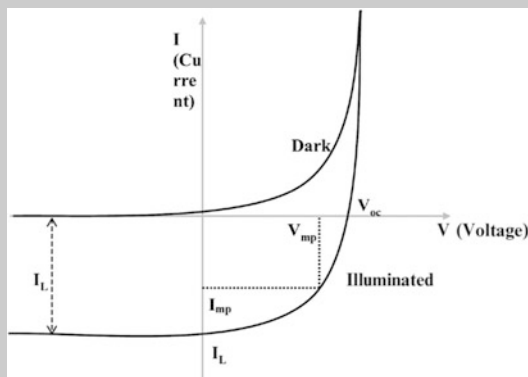
$$FF = \frac{I_{mp} V_{mp}}{I_{sc} V_{oc}}. \quad (12.2)$$

The various values of current and voltages given in the equations are as shown in Fig. 12.6.

(continued)

Box 12.3 (continued)

Fig. 12.6 I-V characteristic of a typical solar cell. I_{SC} is the short circuit current i.e. when $V = 0$, V_{OC} is the open circuit voltage i.e. when $I = 0$ and $P = I \times V$ i.e. the power output of the solar cell



For more details about the solar cell efficiency determination and terminologies the interested readers may refer to the literature or recommended books at the end of the chapter.

As illustrated in Fig. 12.7, TiO_2 particles are deposited on a glass substrate with a conducting coating like FTO (fluorinated tin oxide). The resistivity of a typical FTO coating, few nm thick is few ohm.cm. Usually ruthenium complexes are used to sensitize the titania particles. An iodine layer acts as an electrolyte. Another glass substrate coated with platinum serves as the second electrode. The two electrodes are connected to each other externally through a load to obtain the current.

The TiO_2 nanoparticles in a dye sensitized solar cell are randomly deposited. This is one of the obstacles for the generated charge carriers to reach the appropriate electrodes, reducing the efficiency of the cell. Attempts are therefore made to obtain organized nanoparticles and increase the efficiency. Nanotubes and some other porous media are being considered for this purpose.

Further, the solar cells using quantum dots of CdS, CdSe, PbS, and PbSe are being investigated with varied success. The reason for using these nanomaterials is that they can absorb the visible and/or IR part of the solar spectrum and can replace the liquid dye used in the dye sensitized solar cells. This kind of solar cell is often known as 'Quantum Dot Solar Cell' and is depicted in Fig. 12.8.

Lot needs to be done in order to obtain yet high efficiency quantum dot solar cells. Considering multiple junctions in a solar cell, interfaces need to be improved so that no charges are lost during the transport across the cell. Good n-type, p-type semiconductors compatible with the energy level diagram of quantum dots is a must. Addition of an anti-reflection coating is a challenge. Multiexciton generation in quantum dots could be very useful. Attempts are also made to use plasmonic structures (nanoparticles, rods or other shapes) in order to profitably use their strong visible light absorption capabilities.

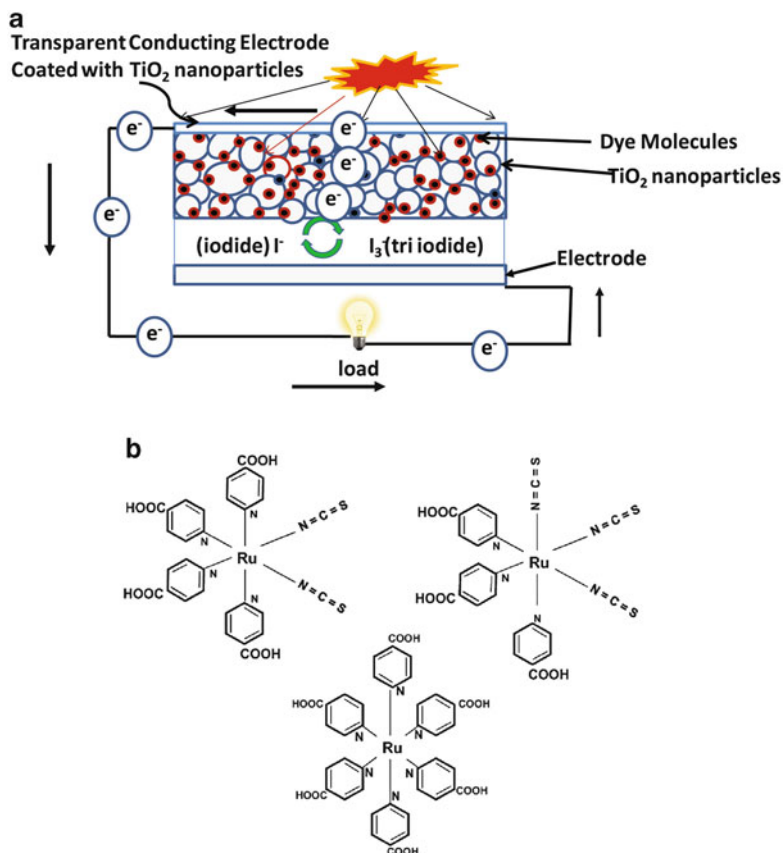


Fig. 12.7 (a) Schematic of a dye sensitized solar cell; and (b) Some dye molecules used in a dye sensitized solar cell

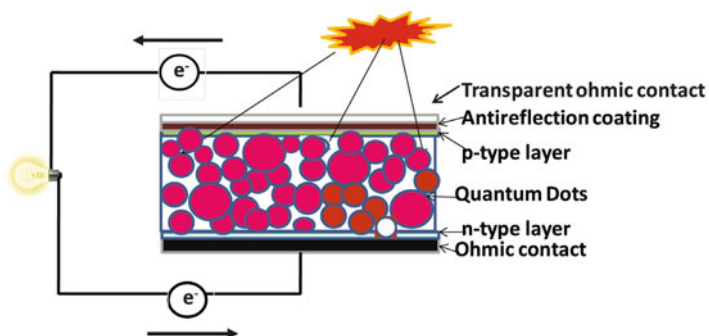


Fig. 12.8 A quantum dot solar cell

12.2.2 Organic (Polymer/Small Organic Molecules) Photovoltaic Cells

Although lots of efforts are being made to improve the efficiency and reduce the cost of the dye sensitized and quantum dot solar cells parallel efforts are going on to fabricate low cost organic solar cells using either small organic molecules or polymers. The most important advantage of use of organic/polymer is their low cost and the possibility of easy processing. Organic molecules can be processed by solution spinning, dipping or drop casting. They are easily deposited on flexible substrates like plastics.

There is the possibility of roll-to-roll process by which, just like newspaper printing (using inks), organic molecules can be printed. This should help in large scale production. Figure 12.9 shows some of the commonly used organic molecules in the organic solar cells. However, only few commonly used molecules or molecular units of polymers are shown here. The scientists are trying to use a large variety of low band gap semiconducting molecules. Often multilayer structures are also designed with different polymers. Inorganic nanoparticles (plasmonic materials like Au and Ag) also are used to increase the functionality like photon absorption efficiency of the organic cells.

The organic cells can be divided into three types: (1) A semiconducting polymer simply sandwiched between two electrodes as in Fig. 12.10a. (2) Two polymers, one electron donor and other electron acceptor are sandwiched between two conducting electrodes (Fig. 12.10b). (3) A bulk heterojunction (BHJ) in which a polymer capable of donating electrons is thoroughly mixed with (usually) an electron acceptor (fullerene based) polymer or molecules. This is illustrated schematically in Fig. 12.10c.

It is important that the excitons be efficiently generated in the semiconductor polymer and electron-hole (e-h) be separated. In a single layer polymer solar cell this is rather difficult. Once generated by absorption of photons, excitons just drift (excitons are charge neutral) and should not recombine. However, if there is an interface of acceptor and donor semiconductor polymer at the interface, they can be separated. Therefore, two polymer layer solar cell as in Fig. 12.10b is better than a single layer polymer solar cell (Fig. 12.10a). In order to absorb sufficient light the polymer thickness needs to be typically ~ 100 nm which is much larger than the typical exciton diffusion length. Thus photogenerated e-h pairs recombine before reaching the interfaces and result into an inefficient solar cell. A better solution is then to blend the donor and acceptor polymers forming a large number of interfaces where the excitons may get formed and electron-hole get separated. The BHJs, therefore, have been widely investigated (Fig. 12.10c) with the hope of getting reliable, reproducible, inexpensive and efficient, flexible solar cells for the future.

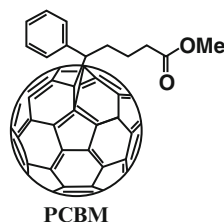
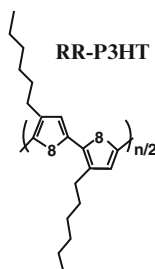
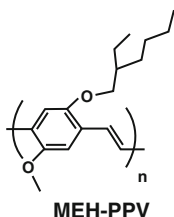
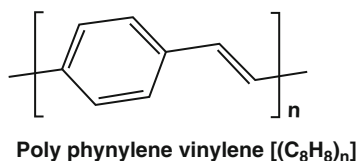
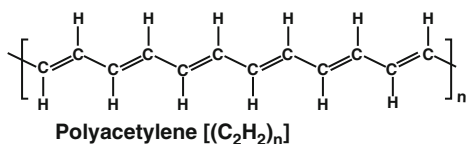


Fig. 12.9 Variety of organic molecules/polymers used in the organic solar cells

12.2.3 Fuel Cell

A fuel cell like photovoltaic cell is a strong candidate to provide electricity. It is an alternative source of energy. The first fuel cell was fabricated by Sir William Grove in 1839 but the name 'fuel cell' was coined in 1889 by Ludwig Mond and Charles Lanser. Grove used four cells in which he used hydrogen and oxygen gases to produce electricity which in turn was used to split water. The commercial application, however, came in 1960 when NASA started the use of fuel cells in the spacecrafts. Indeed in Gemini and Appollo spacecrafts the fuel cells were used

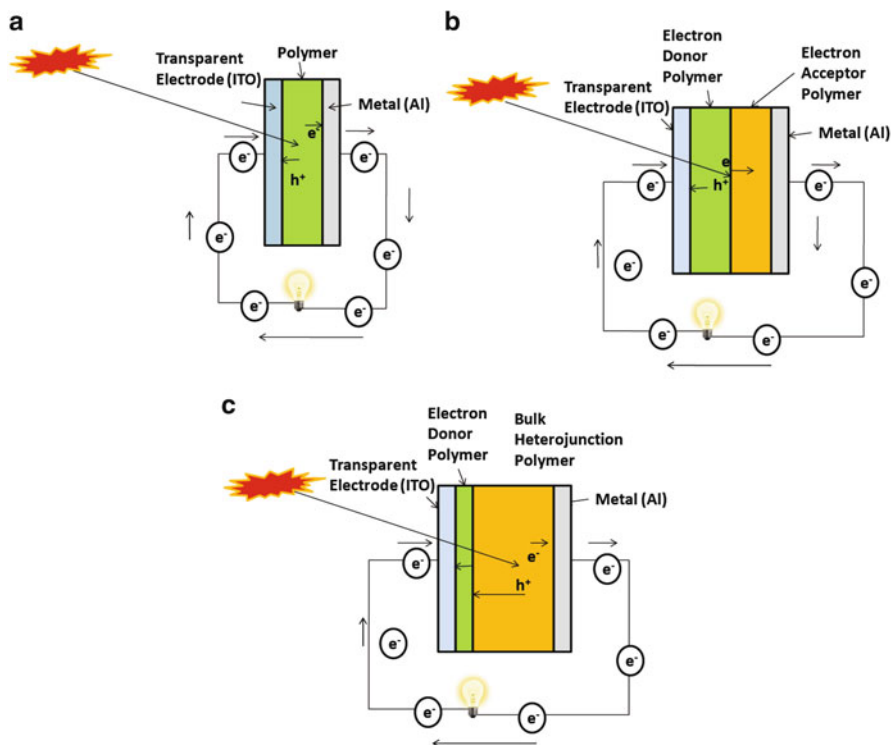


Fig. 12.10 (a) A single layer organic solar cell, (b) A bilayer organic solar cell and (c) A bulk heterojunction organic solar cell

to produce electricity (and water as the byproduct). Since then different fuel cells are being produced and some are in various installations in commercial buildings, hospitals, restaurants, airports and many other public places. They are also used for onboard applications in public transport or personal cars in few countries. The main advantage of the fuel cells is that most of them when used with hydrogen as the fuel provide electricity and heat with water as the byproduct. There is no emission of green house gases, and environment is not polluted. It also does not have moving parts so there is no acoustic pollution. Many automobile companies are trying to use fuel cells in their cars. The main hurdle being the use of hydrogen gas, its supply (stations) and associated risk of carrying high pressure gas onboard. Besides, the cost of fuel (H_2) and expensive electrode materials needs to be reduced. Here nanomaterials would become important. The expensive platinum electrodes can be replaced with porous, non-corrosive alternative electrodes and efforts are being made towards it. Some nanomaterials also are being considered as catalysts for the water splitting as a source of generating hydrogen. The photocatalysis using nanoparticles forms an important branch in nanoscience. Let us discuss now what are fuel cells, their types and how they function.

Currently there are six types of fuel cells in practice, viz.

1. Proton Exchange Membrane Fuel Cell (PEMFC)
2. Alkaline Fuel Cell (AFC)
3. Direct Methanol Fuel Cell (DMFC)
4. Phosphoric Acid Fuel Cell (PAFC)
5. Molten Carbonate Fuel Cell
6. Solid Oxide Fuel Cell (Box 12.4)

As mentioned above they are conceptually similar but vary in their use of electrolytes and sometimes fuel. This may lead to differences in their operating temperatures and also in efficiency. Without going into design and too many materials aspects (which would finally determine the efficiency as well as cost of the cell), a basic outline of these cells will be given.

Box 12.4: Nanotechnology and Fuel Cells

Use of nanomaterials in fuel cells is a hot research topic. Although conceptually simple, cost of fuel cells as well as hydrogen storage/transport and supply are daunting problems. Efforts are, therefore, being made to replace some of the expensive components in the fuel cells like platinum electrodes (catalyst used to dissociate or oxidize hydrogen and reduce oxygen) as well as to improve the electrolytes or polymer membranes using nanomaterials. Nanomaterials due to their size reduction can act as better catalyst sites as well as decrease the cost of the electrodes due to the small amounts needed. Attempts to use platinum nanoparticles and nanowires show some success in this direction. There are also attempts to replace even the costly platinum itself using nickel particles, carbon nanotubes and graphene in the form of composites. These can even be doped with suitable atoms to alter their electronic structure to make them effective catalysts. It is interesting that suitably developed nanocatalyst materials can work equally well for H_2 or methanol as a fuel, giving more flexible fuel option in the same cell. Moreover, attempts are made to use 'self cleaning' electrodes so that H_2 can be replaced with low grade and cheaper hydrocarbon gases as a source of fuel. The problem faced while using inexpensive hydrocarbon gases to replace H_2 is that after dissociation they deposit carbon on the electrodes making them inefficient in short time. With the use of 'self cleaning' nanoparticles the contaminated electrodes can be cleaned at operating temperature of the cells. This would lower down the cost of the cells. There are experimental demonstrations of such ideas at laboratory scale now. Graphene coated indium-tin oxide (ITO) nanoparticles are shown to split hydrogen and oxygen equally well. In the quest for reducing the cost of the fuel cells successful demonstrations are given by developing ultra thin (<100 nm) polymer electrolyte membranes.

(continued)

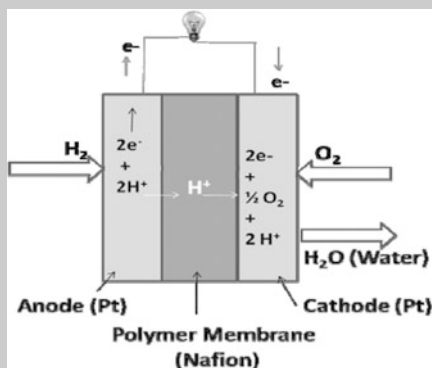
Box 12.4 (continued)

The electrolytes (Nafion) are also blended with sulphonic acid functionalized single wall carbon nanotubes. This improves the transport of protons from the anode to the cathode in the fuel cell.

Thus there are several attempts to reduce the cost of fuel cells as well as make them efficient using a variety of nanomaterials.

Similar to a battery, a fuel cell also transforms chemical energy into electricity. However unlike in a battery, a fuel cell needs a continuous supply of fuel flow viz. hydrogen gas. As illustrated in Fig. 12.11, a fuel cell consists of a solid polymer (usually Nafion) electrolyte sandwiched between two electrodes (usually platinum). Hydrogen gas is flowed in the anode region which catalytically splits in proton and an electron. The proton is conducted through the electrolyte and goes towards the cathode where it combines with the oxygen to form water molecule. The electron flowing in the external circuit assists to complete the reaction.

Fig. 12.11 A fuel cell concept

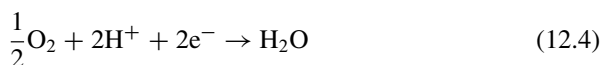


12.2.3.1 Proton Exchange Membrane Fuel Cell (PEMFC)

This cell works on hydrogen gas and oxygen as depicted in Fig. 12.11. It uses (usually) platinum electrodes and solid Nafion Membrane as the electrolyte. The reaction at the anode can be written as



and the reaction at the cathode as



Total reaction is



This results into an external voltage equal to 1.23 V. The operating temperature of the cell is 350–400 K. The efficiency of the cell is 35–60 %.

12.2.3.2 Alkaline Fuel Cell (AFC)

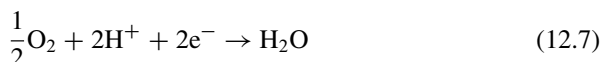
The alkaline cell is similar to the PEMFC cell described above except the electrolyte. In this cell liquid KOH is used as the electrolyte for the conduction of protons from anode to the cathode. The fuel is same as in PEMFC. Its operation temperature is ~360–420 K and efficiency is ~35–55 %.

12.2.3.3 Direct Methanol Fuel Cell (DMFC)

As we will see below, often the fuel hydrogen gas is produced from methanol gas. Avoiding this step of hydrogen production and then using it for a fuel cell, in this cell methanol is directly used as the fuel. The reaction at the anode can be written as



And that at the cathode as



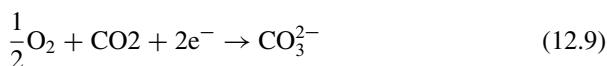
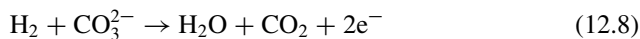
The electrolyte and electrodes are as in PEMFC cell. This cell produces 1.18 V output. Emission of CO_2 is undesired.

12.2.3.4 Phosphoric Acid Fuel Cell (PAFC)

In this cell the phosphoric acid HPO_3 is used as the electrolyte. It is corrosive in nature and can cause related problems. Other things remain the same as in PEMFC. The cell produces 1.23 V. The operating temperature of the cell is ~475–500 K and the efficiency is 35–45 %.

12.2.3.5 Molten Carbonate Fuel Cell

The electrolyte used in this cell is lithium and potassium carbonate. Cheaper nickel can be used as a catalyst. The fuel can be H_2 , CO , CH_4 or hydrocarbon gas. The reactions at the anode and cathode can be



The operating temperature of the cell is rather high ~ 950 K. Efficiency of the cell is ~ 45 – 55 % or even better. The drawback is the emission of CO_2 .

12.2.3.6 Solid Oxide Fuel Cell

This cell utilizes various solid oxides like zirconia (ZrO_2) or yttria (Y_2O_3) or yttria stabilized zirconia (YSZ) as the electrolyte and is a low cost, high efficiency cell. The operating temperature of the cell is $\sim 1,300$ K and efficiency is 50 – 60 %. The electrode reactions are as in PEMF. The cost can be reduced here by replacing the platinum electrodes with other electrodes.

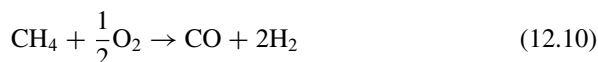
In fact most of the research efforts to improve the efficiency of the fuel cells are directed towards the replacement of the electrode materials with efficient, high temperature stable, non-corrosive materials.

12.2.4 Hydrogen Generation and Storage

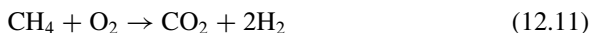
Hydrogen is one of the most abundant elements in our solar system. However it is not the most abundant element on the earth although it is a part of living beings, water, hydrocarbon gases, oil, polymers and many organic molecules. Hydrogen as a fuel does not emit harmful gases and is environmental friendly. Hydrogen gas chemical energy is basically converted into heat, water and electricity. In order to use hydrogen as a fuel it needs to be separated mostly as a gas and then stored either as pressurized gas in strong metallic cylinders, metal hydrides or some new storage systems like carbon nanotubes. It can then be transported and used as per demand. Here we discuss various sources of hydrogen production. It should be remembered that these reactions are accelerated in the presence of suitable catalysts and nanoparticles provide a good platform for this due to their large surface to volume ratio.

12.2.4.1 Partial Oxidation

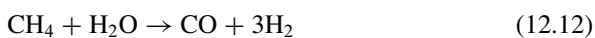
In this case a hydrocarbon gas like methane is oxidized as follows



or



12.2.4.2 Methane Steam Reforming



Above reaction needs heat to be supplied for it to take place (endothermic reaction).

Additional H_2 is produced from CO as



Here nickel nanoparticles are used as catalyst.

12.2.4.3 Electrolysis of Water

The reaction of producing H_2 from water is simple



However availability of electricity can be a hurdle.

12.2.4.4 Photo Electrochemical Cell (PEC)

Water splitting using photons (or solar energy) according to (12.14) is a possible way of getting hydrogen. Use of solar energy to produce H_2 and its use in fuel cells is the best way of getting energy particularly for transportation purposes. The major requirement is that the anode material should have a suitable band gap (semiconductor) to absorb visible radiation, should be robust, reusable, have a long life and should have large catalytic activity. Spherical nanoparticles, nanowires, nanotubes, nanoflakes and some other morphologies are being investigated as they can offer large surface area and catalytically favourable sites. In order to increase the efficiency of hydrogen production the catalytic anodes are modified by using TiO_2 nanoparticles doped with carbon as an anode and platinum as cathode for oxygen generation. Doping of TiO_2 is essential to utilize visible and the IR part of the solar spectrum, as the bulk band gap of TiO_2 itself is 3.2 eV and can increase in nanomaterials. Doping modifies the electronic structure and greatly accelerates the reactions. However, efficiency can be improved by using TiO_2 nanotube array which provides a diffusion pathway without considerable scattering. When UV light

was used on such an anode array the photoconversion efficiency as high as $\sim 16\%$ was obtained. TiO_2 itself needs to be coated on some transparent electrode through which the solar radiation can be incident on the PEC.

The transparent thin films ($<100\text{ nm}$) of Indium Tin Oxide (ITO), F doped ITO and F doped SnO_2 are the usual substrate materials for the anode nanomaterials deposition. The interface of transparent electrodes and catalytic materials can be a source of reducing the efficiency and need attention. Similar to TiO_2 , ZnO nanostructures are suitable candidates for cathodes in PEC for hydrogen production. Aligned ZnO nanorods or other suitable ZnO nanostructures (also doped with Al, Ni or other transition metal ions) have been widely studied and used to improve the light absorption capability of their electrodes. Another suitable nanoparticle or nanorod material as potential anode candidate is $\alpha\text{-Fe}_2\text{O}_3$. Nanorods of $\alpha\text{-Fe}_2\text{O}_3$ doped with W or Cr ions serve as good anode materials. It is found that WO_3 itself in the form of nanowires, nanorods or nanoflakes is a potential candidate for anode in a PEC for hydrogen production. $\text{WO}_3/\text{BiVO}_4$ composite nanoparticles also show a great promise in hydrogen production due to their catalytic activity towards hydrogen. However, more research is necessary to fully exploit the nanostructures in the hydrogen generation and the field is far from maturity.

12.2.4.5 Biological Methods

Human waste, garbage and plant scraps are the biological sources of H_2 . Fermentation or photosynthesis by bacteria and algae can produce hydrogen. Some organisms in water also can split water into H_2 by photosynthesis. However, tapping such resources is still under research. Nanomaterials also should be used to assist the degradation of bio-waste.

12.2.5 Hydrogen Storage (and Release)

The world wide pollution problem is caused mostly due to the emission of gases like CO_2 , CO, NO_x etc. which are not easily absorbed by sea or environment easily. Use of hydrogen has been considered as a good option particularly for the on-board application in the vehicles as H_2 fuelled vehicles emit water molecules as their byproduct in a fuel cell. However storage of H_2 is challenging as it occupies much larger volume compared to gasoline for the same amount of energy release. Storing hydrogen as a pressurized gas and carrying it on-board is not a suitable option for the vehicles. There are also problems with gas pumping stations availability at suitable locations. Therefore the scientists are since long looking for the suitable options for hydrogen storage. Some of the materials considered so far can be metal hydrides, complex hydrides, nanotubes, nanofibres, nanoparticles, polymer nanocomposites, aerogels doped with TiO_2 , Al_2O_3 , MgO or Fe_3O_4 and Metal Organic Framework (MOF).

Ideally the hydrogen storage material should be light weight, inexpensive, available, and simple to use in the desired storage tank; hydrogen releasing temperature should be low, should get released fast and should have long life to undergo number of adsorption-release cycles. The storing material should have at least 6 wt% H storage capability and operating temperature should be less than 400 K. In this respect the research on materials shows that many nanoparticles as well as porous materials are suitable. MgH_2 nanoparticles with Al, Ni nanoparticles, LiBH_4 , $\text{LiBH}_4 + \frac{1}{2} \text{ZnCl}_2$ doped with 3–10 nm Ni nanoparticles, alloys of Mg-Ni-Al, Ca-B-Ti alloys in the form of nanoparticles are investigated as H_2 storage materials. Various Aluminates, particularly NaAlH_4 and Na_3AlH_6 , in the nanoform are also under investigations for their suitability as hydrogen storage materials. Carbon in various forms like carbon nanotubes, fullerenes or graphene along with other materials like NaAlH_4 , MgH_2 or $\text{MgH}_2 + \text{CNT} + \text{transition metal nanoparticles}$ (<10 nm) hold a great promise as hydrogen storage materials. Fullerene C_{60} molecule alone is able to store 60 H atoms on its surface and some inside. It shows ~ 7.7 wt% hydrogen storage capacity but the hydrogen release temperature is larger than about 500°C . However, scientists are hopeful to produce some nanocomposites in future for the suitable hydrogen storage for the on-board application with the capability of recycling or refueling the fuel cells.

12.2.6 Hybrid Energy Cells

We have seen in Sect. 12.2 that there are various types of photovoltaic (solar) cells using inorganic or organic materials. There are trials to improve the efficiency not just by using the nanomaterials of different shapes and sizes but even using organic-inorganic semiconductor layers or blending the materials with plasmonic and oxide materials/layers to improve the solar radiation absorption. Thus various functionalities of the nanomaterials are put together to achieve high efficiency and low cost solar cells to produce electricity.

However as seen from the previous section, the electricity can be obtained even from a battery or a fuel cell using chemical energy and the scientists as well as technologists are trying hard to tap from all the possible sources as reliability on the fossil fuel has not reduced. This is not only polluting our planet but making it difficult to satisfy the energy needs of the human population. In the present scenario, therefore, it is necessary that we tap all the possible resources. The energy demands have given rise to an idea that one should make some hybrid cells which make use of the mechanical, acoustic, thermal, chemical or solar energy to get electricity as per the availability of these resources. For each purpose separate cell is fabricated but can be used in isolation or simultaneously. Such cells need not be large but they are able to run cell phones, light emitting diodes, laptops etc. and they can take care of large load off of energy requirement. This is a new path in energy research using particularly nanodevices.

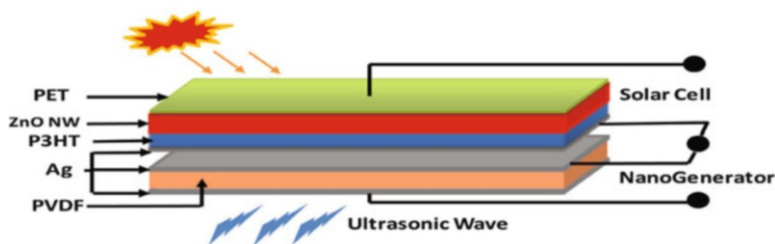


Fig. 12.12 A hybrid cell combining solar, mechanical and thermal energies

This has given rise to make ‘nanogenerators’. Figure 12.12 illustrates schematically a hybrid cell which integrates thermal, mechanical and solar energies. The cell is flexible due to use of polyester (PET) substrate and flexible material layers.

In this cell, a polyvinylidene fluoride (PVDF) polymer film $\sim 100\ \mu\text{m}$ has been sandwiched between two thin silver layers acting as the electrodes. This forms a pyroelectric/piezoelectric nanogenerator. The solar cell is composed of ZnO nanowires grown on the PET substrate and the poly (3-hexylthiophene) or P3HT. The PET and Ag act as electrodes for the solar cell here. The silver layers for NG and Sc are separated by a small gap. The PVDF layer showed the generation of pyroelectric current and voltage ($\sim -24\ \text{nA}$ and $-2.5\ \text{V}$ respectively at $\sim 314\ \text{K}$ and $31\ \text{nA}$ and $3.2\ \text{V}$ at $295\ \text{K}$). The same NG could use mechanical energy to generate piezoelectric current and voltage under compression. The ZnO and P3HT diode could produce with AM 1.5 with $100\ \text{mW}/\text{cm}^2$ light intensity $0.4\ \text{V}$ and $31\ \mu\text{A}$ current. By connecting NG and SC in parallel it was possible to show that there was always an output enough to drive an LED. Interestingly some scientists have made a hybrid cell on a few micrometre diameter sized single polymethylmetacrylate (PMMA) fibre using ZnO nanowires and graphene nanogenerator, supercapacitor and solar cell. It is expected that such a concept of hybrid cell which is still in its initial stage would be used in future.

12.3 Automobiles

A car is made up of large number of parts and materials. Its body and various structural parts are made up of steel, some alloys, rubbers, plastics etc. Body structure should be strong, non-deformable or rigid, of desirable shape and size. It is known that nanotube composites have mechanical strength better than steel. Attempts are made to make composites that can replace steel. Currently the synthesis of nanotubes is not economically viable but attempts are being made so that they can be used on large scale. Cars are spray painted with fine particles. Nanoparticle paints provide smooth, non-scratchable thin attractive coating. Some

research is going on to explore the possibility of applying a small voltage to change the colour of the car as desired.

It is possible to use special window glass materials. One can use 'self cleaning' glass so that it is not necessary to wash the windows with water. Self cleaning glass can be made by dissolving small amount of titania (TiO_2) nanoparticles in it while manufacture by melting together its other ingredients like silica (SiO_2), CaO , Ba_2O_3 etc. Titania is able to dissociate organic dust in presence of UV light available in the sunlight. Once dissociated it may fall down or simply evaporate. Even drops of waters on glass give hazy look. But TiO_2 -containing glass can spread water evenly giving clear sight.

Besides main engine there are large number of small motors in a car. For example wipers, window glass movements, removing CD's from player all need motors. Their operations need one kind of motor or the other. Very powerful electric motors are made using what is known as shape memory alloys using nanoparticles of materials like Ni-Ti. They perform better and are less power hungry than other motors. Such motors are finding their way in automobiles.

Tyres of cars consume considerable amount of rubber which not only increase in price but also add to its weight. Increase in weight is related to reduction of speed and increase in fuel consumption. By using nanoparticle clay—better, light weight, less rubber consuming thinner tyres are possible. Newer cars are expected to employ such tyres. Emission of particles and poisonous gases like CO and NO from car exhausts is one of the biggest concerns in cities. Increasing number of vehicles means increased air pollution affecting a large portion of world population. Use of efficient nanomaterial catalysts is one solution to convert harmful emission into less harmful gases. Large surface area of nanoparticles is useful to produce better catalysts.

Another solution to overcome the pollution problem is to use hydrogen as a fuel. There are numerous advantages of using hydrogen as a fuel. First of all hydrogen as a part of water molecule is abundant on earth as compared to depleting oil used as petrol or diesel after refining. Dissociation of water into H_2 and O_2 is not a difficult process. Therefore abundant H_2 fuel can be made available. When H_2 fuel is burned it can only produce harmless water vapour. However main problem of using hydrogen fuel is its storage. Hydrogen gas is normally stored in a metal cylinder under high pressure. Carrying metal cylinders under high pressure not only can add to the weight of the vehicle but is also dangerous. Hydrogen in contact with air can catch fire. A solution to this problem has been suggested that it be saved in some other forms like metal hydride. Another solution is to store it in 'nanocylinders' of carbon nanotubes. In Sect. 12.2 we have already discussed fuel cells and hydrogen storage issues. Currently many improvements in techniques are necessary to make CNT synthesis economically viable.

Thus nanotechnology may turn out to be one of the indispensable technologies for automobile industries. What is being discussed for cars may equally be applicable to other automobiles.

12.4 Sports and Toys

Nanotechnology has already been introduced into sports equipment and toys. Tennis balls using nanoclay are able to fill pores in a better way and trap the air pressure inside. This increases the life of balls. Some of the international organisations have accepted such balls for their tournaments.

Good quality tennis racquets are made of carbon. Light weight and toughness or strength is necessary for such racquets. It is possible that carbon nanotube composites would serve as a high strength, light weight material for racquets.

In other games too, balls and other equipments using nanomaterials can be employed.

Toy industry also has been well geared to embrace nanotechnology. Eye movements of dolls, robot movements etc. are enjoyed by children but appear quite brisk. Nanotechnology-based motors are being used by toy industry now making the body part movements very smooth, swift and natural looking.

12.5 Textiles

Textile industry is also quite excited about nanomaterials. There are some clothes produced which would give pleasant look of synthetic fibre but comfort of cotton. Special threads and dyes used in textile industry are products of nanotechnology. The clothes made using nanotechnology would not require ironing or frequent cleaning. In fact some companies are using silver nanoparticles in washing machines which make clothes germ-free. Use of silver in either washing or directly in textile assures germ-free environment necessary for bandages, surgical purposes and child care items. The use of highly fluorescent colours from nano semiconductors is one possibility but even just by changing the distance between the particles or changing the size of the nanoparticles woven in the fibre one can change the colour of the clothes. This is because the optical properties of nanoparticles at wavelengths smaller than the wavelength of the incident light depend on their sizes. There are also proofs of the concept that using polymer threads one can weave or integrate solar cells in the clothes so that enough power is generated for charging cell phones or playing MP3 or some such devices. One can also weave in polymer transistors and other passive devices to fabricate a wearable computer. The masks or even fashionable clothes can be made which would either capture pollution or release insecticides when needed to kill e.g. mosquitoes. Thus there is plenty of room to design novelty of wearable textile. The 'self cleaning' carpets or tapestry washing away coffee or tea stains are already marketed. Thus the field of textile research is coming out with novel concepts and useful nano research.

12.6 Cosmetics

Nanoparticles are also important in cosmetics. Besides gold, silver, copper, platinum, and metal oxide nanoparticles like zinc oxide, alumina, silica and titanium oxide, liposomes, solid lipid nanoparticles and nanoemulsions are found to be used in the formulations of various cosmetics like face creams, lipsticks, body sprays, hair care products, sunscreens and so on. Due to their small size nanoparticles-based creams are preferred as they can be used in small amount and do not leave any gaps between them. This gives a smooth appearance. The small particles in some of the creams scatter light in such a way that appearance of the wrinkles is diminished. Some of the nanoparticles can also penetrate deep inside the skin and help repair the skin damages like wrinkles. Usually white sunscreens were used but TiO_2 and ZnO based sunscreens are transparent, much thinner than white screens yet serve in a better way by blocking the UV radiation and protecting the face. Silver nanoparticles in some of the creams are able to kill the bacteria and are able to maintain protection. Nano-based dyes and colours are quite harmless to skin and can be used in hair creams or gels. Nano-based cosmetics are becoming quite popular; however some research on the effect of nano-cosmetics on human bodies shows that some nanoparticles can harm human bodies.

12.7 Medical Field

A great revolution is taking place in biotechnology and medical field due to nanotechnology. The small size of nanoparticles, huge variety of nanomaterials from noble to highly reactive makes it possible to use them for the diagnostic purposes in the laboratory as well as treatments as nanoparticles can be injected, inhaled or digested for different treatments. It is also possible, as we have seen earlier, that they can be designed in different shapes like spheres, wires, rods, tubes or core-shell so that some functional molecules can be attached inside or outside as desired so that drug delivery or molecular recognition can be achieved. Traditional medicines have used gold and silver formulations. Nanoparticles can be good heat or light absorbers or magnetically active and used appropriately. The nanotechnology finds application in medical imaging, drug delivery, cancer and tumour detection and destruction in the early stages, surgery, wound healing, cardiology, Alzheimer and Parkinson treatment, diabetes treatment, dentistry, vision and hearing prosthesis, in body implants and so on. In this section we try to understand how nanotechnology is helping achieving imaging, drug delivery, cancer therapy, tissue welding and bones and muscles treatment areas.

12.7.1 *Imaging*

Conventionally X-rays and computed tomography (CT) are used for imaging. Iodine and gadolinium (Gd) or radioisotopes are often used as contrasting agents. Iodine compounds are injected in the body intravenously which in small quantities can get cleared out of the body quickly. The large quantity has the toxic effects and the patients have side effects. Gd delivered using dendrimers (branched polymeric supermolecules of specific size and shape having highly functional surface) are able to load very small quantities. Gold with large number of electrons is preferred in imaging. In order to make it economically viable, gold is used as a thin coating on silica, alumina or carbon nanostructures making it non-toxic, biocompatible, low cost material. For the Magnetic Resonance Imaging (MRI) silica core-shell particles with conducting metal coatings are developed.

Semiconductor nanoparticles or quantum dots (discussed in details in Chap. 8) are highly fluorescent materials and can be used for biological labelling. When surfaces of such particles are functionalized using some specific antibodies or molecules they can be targeted towards specific receptors in biological cells. This kind of targeting has conventionally also been done using some organic dyes. But they get easily photobleached and also require depending on the dye, particular excitation wavelength for emission to occur. Therefore the analysis is sequential and time consuming. However fluorescent semiconductor nanoparticles have a wide excitation spectrum and can be excited with any wavelength in the UV or blue range. The emission wavelengths depend upon the particle size and can be tuned in controlled way. Semiconductor nanoparticles are resistant to photobleach and preferred. Figure 12.13 shows the fluorescence obtained from CdSe nanoparticles of varying size. They have been simultaneously excited using a hand-held UV lamp. This type of excitation is not possible for simple dyes.

It is also possible to make ‘microbeads’ or ‘core-shell’ particles for biological labelling. Often CdSe or CdTe nanoparticles with high fluorescence ability are coated with thin layer of ZnS shell. The overall size of such a core-shell particle may be <10–15 nm. The shell may be water soluble, hence useful for biological application, whereas the core can be synthesized by a non-aqueous process and may not be biocompatible but highly fluorescent. Fluorescent nanoparticles can be trapped inside the silica or some polymer beads, encoded with appropriate molecules on the surfaces and used to target the specific recognition sites. By taking appropriate amounts of differently fluorescing nanoparticles an exhaustive colour library can be generated. Such microbeads are also useful as barcodes.

There are also imaging applications using silver and gold nanoparticles which have excellent photothermal and optical properties. By targeting the Ag or Au nanoparticles to the cell nucleus the molecular signals within the plasmonic field could be detected. This in turn enables to distinguish cancerous and healthy cell. The effect of drugs on the cancerous cell till its death also can be monitored. Such a technique known as ‘targeted plasmonically enhanced single cell imaging spectroscopy (T-PESCIS)’ has been reported.

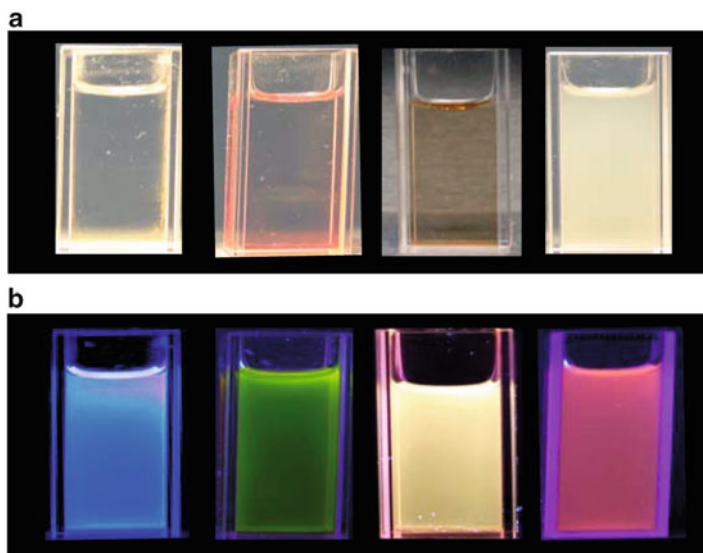


Fig. 12.13 In (a) the solution containing CdSe particles is under normal day light. In (b) the fluorescence produced by CdSe nanoparticles as the size of particles reduces (*red to blue*) can be seen. These are same solutions as in (a) but are illuminated with a hand-held UV lamp

Briefly, nano semiconductor as well as metal (plasmonic) particles are useful in imaging. However, it should be remembered that choosing the nanoparticles of appropriate size can be very challenging. Large nanoparticles can get accumulated in some body parts and lead to dysfunction of the organ. Very small size particles may simply not get detected or not stay in the body for long enough time to get detected. Semiconductor nanoparticles also show blinking. Thus there are still many challenges ahead.

12.7.2 Drug Delivery

In most of the conventional drug delivery procedures the drug circulates in the body. As a result, only small part of the loaded drug reaches the part which needs the treatment. Therefore the new approach is to do the 'targeted drug delivery' i.e. to use some molecular recognition systems or some coatings so that the drug is released only to the desired section of the body. Targeted drug delivery is quite important as it would avoid killing of healthy cells unlike in chemotherapy. There is considerable nanotechnology-based research going on to help diabetic, HIV etc. patients.

Use of appropriate nanoparticles like micelles, liposomes, nanospheres, core-shell particles, albumin-based nanoparticles or dendritic polymers enables to achieve these goals in the modern nanomedicine techniques. They are coated

with small molecules, proteins or peptides. Often the nanoparticles themselves may be effective as they may store drug inside them and release when they get to the targeted site. Initial tests of various drug delivery systems, cancer or tumour therapies or detection have been successful using nanotechnology. Nanoparticles being very small are easy to inject and target towards specific portion in a body. Particle surface can be modified using some functional molecules which can then go to specific receptor site. Some of the important examples of drug delivery system are in cancer therapy, insulin delivery and treating Alzheimer as well as Parkinson disease.

The scientists have developed novel nanocarrier systems to deliver cancer treatment drug to cancerous cells without affecting the healthy cells. This avoids excessive drug circulating in the body causing toxic effects. Some of the drugs are also not easily soluble in many solvents. In such cases nanocarriers turn out to be effective means of drug delivery.

Diabetes mellitus or simply diabetes is a disease which increases the level of blood sugar in the body. There are Type I and Type II (and some more, less frequent) diabetes. The main cause of diabetes is that the insulin producing β -cells in pancreas either stop producing insulin or an insulin resistance (Type II) is developed which can be controlled by regular medication. In Type I insulin needs to be regularly injected in the patient's body. Insulin is a peptide which is broken down by enzymes and helps to lower the glucose level in the body. However, controlled dose of insulin is important otherwise it leads to hypoglycemia. Regular injections are also troublesome to the patients. Insulin pumps (external wearable or implanted inside the body) regularly supply insulin to the body. Yet, attempts are made to deliver controlled insulin in nanocapsules made of polymers, polyanhydride or polyester by injecting it in soft tissues. Unfortunately their oral formulation is not yet developed as polymeric forms are not very stable in the gastrointestinal track. Chitosans are also found to be useful in insulin encapsulation. Chitin is a natural polymer which gives strength and protection to the cells in animals. It is basically a polysaccharide $(C_8H_{13}O_5N)_n$. However, it is not very successful as its initial burst is difficult to control.

In order to protect the delicate network of central nervous system of our body, there is so-called blood-brain barrier which does not allow large molecules to enter our brain. This is done by epithelial cell lining of blood vessels in the brain. The epithelial cells are lyophilic through which small molecules can pass due to their transporter proteins. Some large molecules can be transported across the barrier to brain through such transporters if they are able to bind with them. Yet, many large molecules which can be potential drugs to cure Alzheimer and brain cancer are not able to pass the blood-brain barrier. The suitably modified polymeric nanoparticles in the form of liposomes loaded with drugs are able to pass the barrier and make drug delivery. Radiolabelled Fe^{3+} and Cu^{2+} metal chelator clioquinol has been transported across the blood-brain barrier through lyophilic nano drug carriers.

At present the drug delivery using nanostructured materials is an active research area.

12.7.3 Cancer Therapy

Genes control the growth and division of cells. If genes are damaged the cell growth is uncontrolled and obstructs the normal body functions leading often to death. There are various types of cancers but the main cause for cancer is the cell growth. If not treated in early stages, it is fatal. The early detection of cancer has been claimed to be possible due to nanotechnology based diagnosis. This would enable to treat the patients before it is too late. As discussed earlier the drugs can be encapsulated in nanocapsules and targeted towards desired parts of a body. Drug can then be delivered at desired rate, by opening the capsule using some external stimulus like magnetic field or infrared light or under some physiological conditions. Targeted drug delivery is quite important as it would avoid killing of healthy cells unlike in chemotherapy. In fact targeted chemotherapeutic drug delivery has been successful. In this the cancer cell target is exploded and destroyed when encountered.

Fluorescent, nanoporous, water soluble and biocompatible silica nanoparticles loaded with camptothecin, taxol and other hydrophobic, toxic drugs (such as β -lapochene) are successfully delivered in animal models to cancer cells and destroyed. As the cancer recognizing ligands are attached to the silica nanoparticles healthy cells are spared which can eliminate the side effects due to the treatment on patients. By this nanotherapy it would be possible to cure breast, stomach, bladder, neck and colon (greater part of large intestine) cancer. To successfully treat cancer β -lapochene is delivered to cancer cells in polymer implants which gets slowly released destroying cancer.

Phototherapy (see Fig. 12.14) is another method which is becoming popular. One can use gold coated silica nanoparticles and target them to the cancer cells. The scattered light from such particles enables the location where the particles are attached. In fact this enables imaging combined with treatment. The silica-gold core-shell particles sizes can be tuned in such a way that there is maximum absorption of the laser power. This is possible (see Chap. 11, Sect. 11.7) due to Surface Plasmon Resonance of gold nanolayer. The IR laser beams are available. The advantage of using IR is that it penetrates inside the body without affecting the body. Moreover, the laser power required to destroy the cancerous cells is half than that needed to destroy the healthy cells, so the laser is operated at low power anyway. The heat generated by the core-shell particles raises the local temperature which is sufficient to kill the cancer cell.

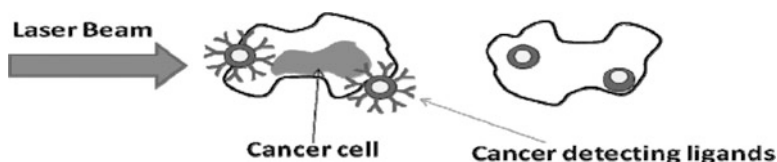


Fig. 12.14 Use of core-shell silica-gold nanoparticles to destroy cancer cells. Note that ligand are attached to the core-shell particles which identify the cancer cells and then with laser beam they can be destroyed

Instead of making gold core-shell particles, scientists have also coated gold nanoparticles with peptide lipid which encapsulated 'silencing ribonucleic acid (siRNA)' as the drug which were then directed towards cancer cells and then exposed to the laser beam. This also was an effective way of destroying cancer.

Another type of material useful in cancer therapy is magnetic nanoparticles. Magnetic particles with extremely small size can become superparamagnetic as discussed in Chap. 8. Every magnetic nanoparticle has a characteristic small size in nanometer range below which it becomes superparamagnetic. Such superparamagnetic particles (usually Fe_2O_3) when subjected to an AC magnetic field increase the local temperature around them by $\sim 4^\circ\text{--}5^\circ$ than the normal body temperature. This can destroy the cancer cells. This type of treatment is known as 'hyperthermia' and has a great application in the real world.

However, the above mentioned treatments are yet not available in the hospitals as they await the government approval.

Ideally a nanocarrier needs to be developed which will detect cancer in the early stage, destroy and even report that the cell death has occurred. Research in this direction has long way to go.

12.7.4 Tissue Repair

Tissues comprise of similar type of cells working in co-operative manner to give some functionality to body. There are four types of tissues viz. connective tissues, nervous tissues, muscle tissues and epithelial tissues. Connective tissues are responsible for body structure and support. They include fat and bones. Nervous tissues control the body processes. Muscle tissues enable movements of body parts. However, in body all actions are not voluntary but some muscles are continuously working on their own, involuntarily for body function like breathing. Epithelial tissues give inside and outside cover throughout the body.

Joining of tissues is important to repair the cuts, cartilage in joints, blood vessels, skin, cornea etc. For very long time this was done using chromophores or dyes as absorbers of laser power. With a small incision in the body an optical fibre can be inserted in the body for surgery using a laser. Specific lasers are required for specific dyes. Conventionally in the indocyanine green (ICG) dye an organic molecule is used in conjunction with infrared laser. However absorption is limited compared to now available gold nanoparticles (absorption is about 4–5 orders of magnitude more than dye molecules). Dyes also bleach fast or disintegrate. Gold nanoparticles on the other hand are quite stable. They can be tuned to desired/available laser wavelength in IR region, particularly if one uses different shapes like nanorod to tune the wavelength of absorption. Here the phenomenon of surface plasmon resonance (SPR) discussed in Chap. 7 is responsible for the strong absorption of light. The absorbed heat by the gold particles is used to heat the tissues which get sutured. This is joining of the tissues without stitches which reduces surgery related pain and infections. Although laser welding of tissues using gold nanoparticles has

been successfully demonstrated in quite a few laboratories in the world, it also has not come in practice for human beings.

The scientists are developing better body implants than available so far. The body implants should be strong and biocompatible. Body cells should be able to grow on them and hold in place. Growth of artificial muscles also is being one of the aims of nanotechnology. Much work is going on in understanding how the nature does it all and mimic it.

12.8 Agriculture and Food

May be it is urban or rural population in developed, developing or undeveloped country, food is the basic requirement. Importance of agriculture, therefore, needs no explanation. The agricultural raw products after cultivation and production need processing and proper packaging before they reach the consumers. Nanotechnology in more than a decade now is emerging not only for higher yield but also to lower the various inputs like pesticides, herbicides or reduce the requirement of water. Pesticide DDT was very popular in the twentieth century but due to its toxicity is banned now. Traditional method of increasing the yield of the crop has been to grow different crops alternately.

Most of the nanotechnology research in agriculture is related to sensing the soil condition, growth of plant parts and provide adequate amount of fertilizers, pesticides or herbicides as per the requirement. This can help proper and economical growth of the crop. Efficient and smart sensors are obviously the requirement (a smart sensor would sense the presence of a disease and release the pesticide/herbicide in adequate quantity). This is through what is termed as '*precision farming*' in which use of computers, global position sensing (GPS) and remote sensing can be made. For this variety of sensors will have to be spread in the fields. Some of these ideas are in their infancy but scientists dream of bringing them into reality.

Some groups are trying to investigate plants of tomato, soybean, rice, wheat, corn and so on which are consumed on a large scale. The nanoparticles of Au, Ag, Cu, CuO, SiO₂, Fe₂O₃ and ZnO, fullerenes, and carbon nanotubes are under investigation to use them catalyze fertilizers or pesticides more effectively or use as sensors. Interaction of synthetically made nanoparticles and naturally present nanoparticles in soil, however, is not known yet. Thus a major thrust in agriculture research using nanomaterials is expected to find out ways and means to produce healthy and large quantity of agriculture products.

Healthy raw crop is the first important step that needs to be achieved. However, lot of agriculture products if not processed and delivered properly, can result into waste of energy, man-power and money. Lot of major food product producing companies have turned their attention to use nanotechnology in packaging, increasing the shelf life of products, maintaining odour and taste of products, increasing food nutrition value etc.

The food companies are using some innovative ideas of nanotechnology. It is known that omega-3 fatty acid reduces blood clotting, decreases platelet aggregation, improves insulin response in diabetic patients, reduce obesity and help prevent cancer growth. Fish (tuna and others), soyabean, tofu etc. are rich in omega-3. Tuna fish oil odour smells bitter and unpleasant making it difficult to take it in food. The scientists have encapsulated tuna fish oil in a biocompatible capsule which bursts only after reaching stomach.

Ice-creams with 1 % fat as compared to usual ice-creams with 16 % fat are produced by reducing the size of particles which give ice-creams a texture.

Bioavailability for nutrition can be increased using nano self-assemblies of micelles. The nutrients enter blood stream from stomach and are made available. The technology is used to reduce cholesterol and this would be useful also to deliver vitamins. A company has produced the nano product which not only reduces the energy needed for cooking but extends oil life used in deep frying.

Packaging is very important in preserving food for long time as well as transporting it to long distances. A company has used SiO_2 nanoparticles in making air-tight wrappers which preserve food for longer time. Nanoplastic wrappers are produced using zinc oxide nanoparticles as catalysts which are antibacterial, UV and temperature resistant and fire-proof in nature. Nylon-based nanocomposites are used to bottle the beer to make a barrier for CO_2 and O_2 to keep the clarity, flavour and increase the shelf life. Emphasis of food preserving and packaging industry is to use nanocapsules inside so that flavour, odour are triggered only when used. Colours and nutrients also get added only when product is either opened or is being consumed. Capsules are dormant until then. It is also possible to attach colour strips to qualify if the food is good to eat or already started becoming stale. This would be convenient for food packages in big malls. Briefly, lot of revolution using nanotechnology is underway in agriculture and food supply.

12.9 Domestic Appliances

Use of silver nanoparticles is made in refrigerators, air purifiers or air conditioners and water purifiers. It is well known for long time that silver has antibacterial property. That is why it has been used as utensils material since long. But recent research has shown that silver nanoparticles are much more effective and only small quantity of them is required. Therefore some manufacturers have special nanoparticles lining in refrigerators, air conditioners or even in washing machines.

Food in refrigerators can remain fresh and prevent fungal growth for longer time than ordinary refrigerators. The clothes washed in silver nanoparticles lined washing machines are claimed to stay sterile for about a month. This should be quite useful in hospitals. Air conditioners or water filters with silver nanoparticles also are claimed to have advantages and are being marketed with very little additional price.

Some of the building blocks like window materials can be based on nanomaterials. One can maintain the inside temperature of the houses reducing heating/cooling

effects due to outside weather using appropriate window materials like aerogels which are highly insulating but can be made transparent for window purpose. Self cleaning glasses also are useful for windows. One can also use some window coatings to adjust the interior lighting of the houses. Dye sensitized solar cells can be an integrated part of modern architecture increasing the aesthetic appearance as well as satisfying local power requirement of the house or a building. Thus nanomaterials are quite useful.

12.10 Space, Defense and Engineering

Space and defense scientists also are trying to adopt nanomaterials as alternative materials and replace the conventional materials. Very low density materials known as aerogels (Chap. 11) are nano porous materials (i.e. materials having small nanosized pores in them). Aerogels can be of various materials. Their density is typically $0.01\text{--}0.8\text{ g/cm}^3$. One can compare this density with that of iron which is $\sim 7.8\text{ g/cm}^3$. So one can get an idea of how light aerogels are. Naturally it is very good to use for various applications in spacecrafts and defense to reduce the weight. Even some special light weight suits, jackets etc. can be made using aerogels as they can be made such that they are poor conductors of heat. Even some special high temperature materials which are otherwise difficult to make can be made at lower temperatures using nanomaterials.

In space, mainly solar energy is used to power the satellites or space crafts. Some of the solar cells currently in use have reached an efficiency of $\sim 30\%$ at air mass zero (AM0). Some solar cells are approaching an efficiency of $\sim 40\%$. However there is still a concern in reducing the weight of materials used in solar cells. It is speculated that future spacecrafts would be powered with luminescent dye sensitized nanoparticle based or nanoparticle-based solar cell arrays.

Space vehicles also need high performance, multifunctional materials which can withstand harsh and extreme environments during launching and in space. Materials should also sustain high or low temperature and high or low pressure. For some parts light weight polymers are quite attractive. Low processing temperature, possibility of having them as fibres, coatings or thin films makes polymers attractive as internal insulation in solid rocket motors. Polymer composites using silica fibres and nanoparticles have larger Young's modulus, low temperature coefficient of expansion and high impact strength.

It can be seen from Fig. 12.15 that nanoparticles in polymer composites are better radiation protectors as compared to microparticles-based composites.

In satellites better ignitors of nanocrystalline materials are being considered. Alumina particles are used in solid propellants. Nanoparticles of alumina are at least six times better than the conventional particles. A nanocomposite of Fe_2O_3 and aluminium burns much faster and is more sensitive than conventional thermites.

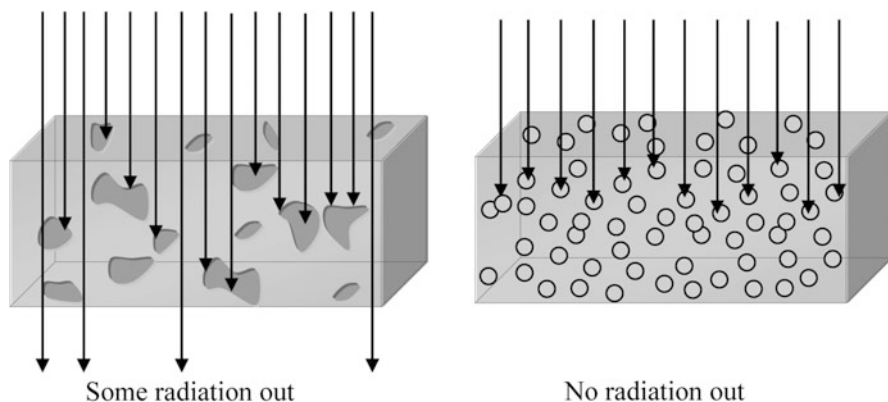


Fig. 12.15 Nanoparticle-Polymer composites

In aircrafts, superior property, specially fatigue resistant materials are required. Fatigue strength usually decreases with time. Some nanomaterials have better fatigue strength and life is increased by $\sim 200\text{--}300\%$.

There is always a fear that terrorist or enemies may use dangerous microbes or viruses as weapons. Quick detection of biological weapons is the urgent need of defense departments and other public offices to 'on the spot check' the dangerous traffic of harmful biological weapons or epidemic. For example it becomes necessary to detect anthrax or SARS patients immediately. Nanotechnology is expected to play an important role in these areas.

Some of the chemicals used in warfare are VX, HD, GD and GB. Some nanoparticle oxides like CaO , Al_2O_3 and MgO interact with such chemicals much faster than microparticles and are ideally suited for fast decomposition of warfare chemicals.

As was discussed in Chap. 11, the new kinds of structures based on metamaterials will be useful in cloaking or making the objects invisible. This will have considerable defense applications.

Further Reading

- M.A. Green, *Third Generation Photovoltaics, Advanced Solar Energy Conversion* (Springer, Berlin/New York, 2003)
 G.L. Hornyak, H.F. Tibbals, J. Datta, J.J. Moore, *Introduction to Nanoscience and Nanotechnology* (CRC Press, London, 2009)
 A.W. Miziolek, S.P. Karna, J.M. Mauro, R.A. Vaia, *Defense Applications of Nanomaterials*. ACS Symposium Series 891 (2005)

Chapter 13

Nanotechnology and Environment

13.1 Introduction

Whenever any new technology emerges there is rightfully a concern about its impact on social life, health and environment. This was seen when first trains were run. People were worried if there would be effect of travelling in a train. In the twentieth century many new technologies were developed and even a common man knows the benefits of new developments in science and technology. However, along with comforts achieved by adopting various technologies, mankind is now facing numerous problems. The problems are not limited to any small country or community, to rich or poor but by all. Major problems being faced are due to increased population in developing and underdeveloped countries, old and new diseases, depletion of natural resources like fossil fuel, oil and water, increased number of wars, terrorism and so on. Some of the problems are also related to global climatic changes, scarcity and quality of food and water as well as increased pollution in big cities. Therefore some often asked questions are: whether nanotechnology would help to solve or increase some of these problems? Do nanomaterials harm human bodies, animals or plants, may be by inhaling or by contact? Will nanomaterials pollute water, air and food?

We have seen from the earlier chapters that nanotechnology is capable of giving wonderful products at lower cost due to small amount of material requirement enabling their accessibility to large population. This is true for health care also. Therefore it is difficult to resist nanotechnology any more. Some scientists estimate that although there can be harmful effects of nanomaterials on environment and human bodies if proper care is not exercised, Nanotechnology can be so powerful that it would outweigh its own negative aspects. In this chapter we shall discuss (1) environmental pollution and role of nanotechnology and (2) effect of nanotechnology on human health.

13.2 Environmental Pollution and Role of Nanotechnology

Environmental pollution includes pollution or contamination of air, water and soil. When the level of chemicals, metal (ions), smoke, bacteria, viruses or pathogens increases beyond some tolerance limit for living animals it is pollution. The pollution of environment is an age-old problem and began when man learnt to make fire and started cutting wood for houses and other activities. When human settlements began and various civilisations flourished, the pollution problems, mainly water contamination, started. However until probably the nineteenth century the population on earth was not large and pollution problems were not severe. On one hand the advances in medical field accompanied by sophisticated electronic equipment, increased agriculture production, thanks to fertilizers, pesticides and improved technologies available for agriculture have resulted into human life span, it has also increased demands on houses, water, transport etc. to make life comfortable. There is a major threat of fossil fuel getting depleted to an alarming level on which we are dependent for our energy resources for all our activities. We continue with use of fossil fuel which contaminates air, as still it is cheaper than any other alternative energy. This in turn also gives rise to global warming i.e. possibility of increasing the overall temperature of the earth, resulting into melting of large icebergs on poles.

The industrial effluents, increased automobiles, trains and aircrafts add to water and air contamination. It is hoped (as well as believed) by the scientists that the nanotechnology will be able to reduce the stress on energy demands from fossil fuel by providing clean, alternate energy sources like photovoltaics and fuel cells at lower, affordable cost. As seen in the previous chapter, solar energy being abundant and almost inextinguishing source of light, if harnessed at low cost will solve our lighting problems. Nanotechnology-based solar cells may be dye-sensitized solar cells, quantum dots solar cells or organic solar cells and are possible to make in large quantities at lower price. Current obstacle is the high efficiency expected from these solar cells to become economically viable. However, future is bright and it is hoped that technological improvements will take place. Same is with fuel cells which is source of energy long awaited by the automobile industry for the onboard application. This is also a clean source of energy. When hydrogen gas is used as a fuel it only generates, along with the electricity, water and heat as the byproducts. The use of fuel cell powered vehicles will dramatically reduce the air pollution and demand for fossil fuel. Thus in both these energy sources the contamination will be eliminated unlike fossil fuel, at the source itself. We then need not worry as to how to reduce the pollution.

It has also been shown that nanocatalysts due to their increased surface activity are able to reduce the toxic products from exhausts of the vehicles running on petrol or diesel as the fuels. Therefore palladium, platinum and rhodium, in spite of their high cost are used for many decades. Some cheaper alternatives to these noble metal catalysts are nanoparticles of metal or oxy carbides in use. It should be remembered that these nanocatalysts were not purposely made 'nano' but many conventionally used catalysts happened to be nanoparticles.

Interestingly, gold, which is not a conventionally considered as a catalyst material (probably due to its non-reactive nature), in the nano size turns out to be a very good catalyst material. Gold nanoparticles impregnated in magnesium silicate hydrate clay, catalytically destroy odours in presence of ozone. Gold nanoparticles and even better gold-platinum alloy nanoparticles are able to dissociate unwanted trichloroethylene in ground water. Gold nanoparticles also interact with pesticides and are useful in removing them from water. It is also possible to use gold nanoparticles to remove one of the important toxic element mercury as a water effluent from coal mining industry. In view of the success of gold nanoparticles as a catalyst they were also used in the airconditioners to successfully remove CO from air in the rooms. It is also claimed by some scientists that the use of gold nanoparticles could reduce the hydrocarbon emission by 40 %.

Water and soil pollution is mainly due to various human activities, industrial effluents in which organic molecules (like dyes), inorganic ions of metals like Cu, Hg, Cd etc., pesticides, fertilizers, septic tank seepage are mixed. Although ground water (from larger depths from the earth surface) should be free from most contaminants due to industrial, agriculture and other human activities, it often contaminates due to seepages, mines, agricultural activities as well as natural minerals contamination. It can also have pathogens like bacteria and viruses and different salts. Therefore to provide fresh drinking water to humans and animals is quite challenging. This is mostly done by municipal organisations at some central water distribution system. The dirty water is usually treated with alum to remove large dirt particles by sedimentation and then filtered using sand and charcoal. This is followed by ozone, chlorine or UV radiation treatment to remove the bacterial contaminants before it is circulated to the village, town or city population. It is found that these treatments do not remove most of the nanoparticles like carbon nanotubes.

The problem with nanoparticles in this regard is that they often have organic ligands attached to their surfaces and some of them like carbon nanotubes (CNT) are hydrophobic in nature. They do not get removed by conventional treatments and are a potential health threat if they enter human (animal) body. It has been reported that CNTs and other fibre-like nanoparticles can lead to a lung disease known as fibrosis which was reported few decades back caused by asbestos fibres. As we shall see below some of the nanoparticles either due to their small size (small $<10\text{--}20$ nm size particles can easily penetrate the cells) or their large surface activity can be harmful. Thus it is necessary that nanoparticles should not mix with water. On the other hand some research shows that carbon nanotubes and even graphene are best water filters. Silver nanoparticles are being used in water filters, airconditioners, bandages and washing machines. It is known that silver is traditionally proven antibacterial material. Although use of silver nanoparticles in the above applications is justified, there is a concern that when it goes into water treatment tanks it also kills some useful bacteria/ingredients in water. Therefore there is a concern about the use of silver nanoparticles in these applications. Similarly CNTs used in inverse osmosis cartridges used to remove cations in water need to be studied carefully.

Interestingly, nanomaterials-based photocatalysts are developed which can dissociate organic pollutants and remove them. Some such catalysts are nano ZnO, TiO₂, SnO₂, Fe₂O₃, CdS, MoS₂, ZnS, CdS and PbS. Thus, some estimates claim that nanotechnology will in fact reduce the air and water pollution that we have today.

It may be added here that along with the synthetic nanomaterials like metal nanoparticles of gold, silver, copper, platinum, transition metal nanoparticles, quantum dots of CdS, ZnS, CdSe, ZnSe, PbS, SnO₂ etc., magnetic nanoparticles, fullerenes, carbon nanotubes, graphene etc. we also have a huge variety of naturally occurring nanoparticles like silica nanoparticles, zeolites, iron oxides and various organic particles. Such particles are produced mostly in volcanoes, fires, erosion or by marine waves.

Important question to be discussed here is which are the techniques available to us to detect the pollutants in air, water or soil. In fact there is a huge list of methods or sensors available today. Some of them are gas chromatography, high performance liquid chromatography, ion and ion exclusion chromatography, atomic absorption or emission spectroscopy, electrode methods, fluorescence and biosensors. Depending upon the analyte, the sensitivity may vary from ppm to even ppb level. However, gold and silver nanoparticles (in different morphologies) based Surface Enhanced Raman Spectroscopy (SERS) are considered to be best amongst all for their extremely high sensitivity (single molecule sensitivity). The requirements of any good detector/sensor are its high sensitivity, selectivity, and real time measurement. Simple design and portability are additional advantages. SERS sensors are proving themselves not only for heavy metal ion, CNT, graphene and organic molecule detection but explosives like TNT (2,4,6-trinitrotoulene), and DNT (2,4, dinitrotoulene) in water and soil. SERS sensors have additional advantage that they do not require complicated sample preparation for testing.

Nanotechnology is still in its early stages. It has demonstrated its ability to sensitively detect the pollutants as well as remove them if they occur. The future efforts would be to avoid or reduce the pollution as much as possible with the use of clean energy resources like solar cells and fuel cells for the global energy needs. Some of the experiments show that unintentionally the nanoparticles can mix into our environment and their removal would be challenging.

13.3 Effect of Nanotechnology on Human Health

We saw in the previous section that nanotechnology may help reduce the air, water and soil pollution through its nanocatalysts like silver, gold and some other particles as well as produce electricity/energy using nanotechnology like in solar and fuel cells which promise to provide clean energy at low cost. As fossil fuels are depleting as well as polluting our planet to a great extent, we do not have any other option but to embrace a technology which promises us the products we need to cater the needs

of very huge population i.e. newly emerged nanotechnology. No other option is available. Yet, we also have to be aware of pros and cons of using this technology so that we can extract benefits from it and avoid its disadvantages.

When nanomaterials will be produced in large quantities and used by mankind, it is expected that some part of it could unintentionally return to environment. Some nanomaterials would sweep in water or air while they are synthesized or when the used products get thrown after their use, as garbage. They are then bound to partially return to humans and animals through water they drink or food they eat. Hence some of the nanomaterials in water are found to be difficult to remove. What happens to the nanoparticles we may eat, drink or inhale? Do they harm our body? What kind of health problems would they create? Some nanoparticles intentionally will be injected or transmitted in bodies to cure some diseases like cancer, Alzheimer or carrying out surgeries. What would happen to them? Is it safe to have treatments based on nanotechnology? As many of the animal trials are successfully carried out, there is possibility that we may see many medicines, medical imaging and surgeries based on nanotechnology taking place. Therefore the scientists are now quite concerned about toxic effects of nanotechnology. Although much more needs to be done about the toxicity problem, few experimental results (on rabbits, fish or mice) are available now and are briefly discussed below without getting into the actual details.

Small size of nanoparticles increases their uptake as well as their interaction with the body tissues or cell. This can release free radicals which in turn can give rise to oxidative stress, damage of protein, DNA or cell membrane. It is quite easy for nanoparticles to get mixed in the blood stream due to their small size and reach various body parts. They can then damage the organs in short time inhibiting the new cell growth or causing cell death. Kidneys are found to be affected by the carbon nanotubes. CNTs are also known to cause fibrosis—a lung cancer.

Titania (TiO_2) nanoparticles are known to be very efficient photocatalyst material. TiO_2 nanoparticles also interact with bacteria. It is found that TiO_2 nanoparticles in the water treatment plants destroy even the useful bacteria in water. Same is true about fullerenes in water. This is not desired. The most popular application of TiO_2 nanoparticles is in the dye-sensitized solar cells. They are also used in cosmetics. It has been reported that TiO_2 nanoparticles have size selective effects on DNA. While particles of TiO_2 larger than 500 nm did not affect DNA but ~20 nm size particles interacted with DNA and broke it completely. The TiO_2 particles of even much smaller size i.e. ~2–5 nm showed inflammation in the body.

There is a strong drive in using formulations with ZnO, TiO_2 , gold, silver, iron oxide, silica etc. in some cosmetic formulations, sunscreens, moisturizers, sprays etc. Cosmetics usually have liposomes, nanoemulsions, solid lipid particles, nanocapsules, nanoparticles or metal oxide nanoparticles in their formulations. Nanoparticle-based sunscreens are popular because the face looks more natural than with white sunscreen lotions (older sunscreens). Some of the nanoparticles penetrate the skin quite fast. This is sometimes desired in some treatments like in anti aging treatments. The particles are expected to repair some cells at depth. However it

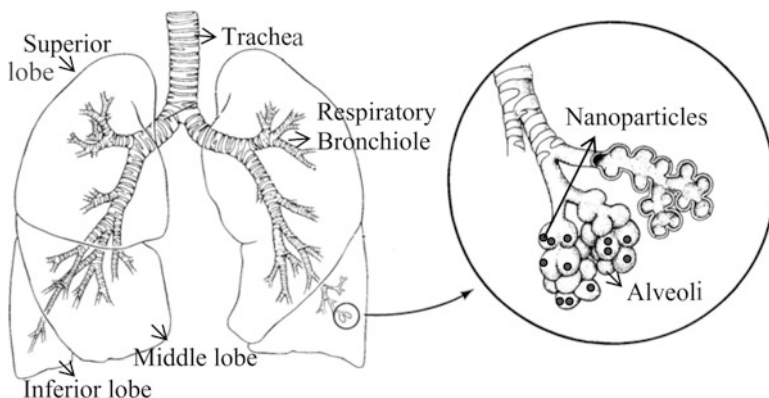


Fig. 13.1 Human respiratory system

was found that some nanoparticles could cross the blood-brain barrier (in mice) and penetrate the cells in brain and damage them. Inhaling of nanoparticles or spray also can let the particles reach the brain.

Liposomes (bilayered vesicles) also have fast skin penetration. They can also be potentially harmful. Fullerenes also are used in some cosmetics and one needs to know its effect.

Particles of size around 50 nm can affect cells and ~ 30 nm can affect central nervous system. Endings of our respiratory tree consist of small packets known as *Alveoli*, which are $0.3\ \mu\text{m}$ in size. They get affected by particles of size around 70 nm. In general small particles with size $<0.1\ \mu\text{m}$ affect our respiratory system (see Fig. 13.1) and other organs.

Details about nanomaterials and size dependent effects are not fully known yet. More work is necessary to understand the effects of nanomaterials on environment and human/animal bodies. On one hand nanotechnology may solve many of our energy, health diagnosis, medical treatment, pollution detection and removal problems and yet on the other hand it may threaten our health due to our ignorance. Perhaps the extensive research in future will teach us how to profitably use the nanotechnology in controlled way, which particles (material compositions and sizes) should be used for the benefit and which should be avoided to protect our environment and ourselves.

Further Reading

- G.L. Hornyak, H.F. Tibbals, J. Dutta, J.J. Moore, *Introduction to nanoscience and nanotechnology* (CRC Press, Boca Raton, 2009)
- B. Karn, T. Masciangioli, Wei-xian Zhang, V. Colvin, P. Alivisatos, *Nanotechnology and the Environment, Applications and Implications*. ACS symposium series 890 (2005)

Chapter 14

Practicals

14.1 Introduction

In this chapter we shall briefly discuss some easy ways to perform experiments which will enable the beginners to get some experience in the synthesis of nanomaterials and characterization. We shall mention the equipment required and outline the procedure along with some results. These experiments were actually performed in the author's laboratory and are possible to perform in any moderately equipped laboratory without heavy investment.

For chemical synthesis discussed here, simple and inexpensive set of equipment is required. As shown in Fig. 14.1, a multiple neck, round bottom, glass flask can be used. Chemical synthesis reaction is generally carried out under some inert gas (N_2 or Ar) atmosphere to avoid uncontrolled or undesirable oxidation of the nanoparticles by constant flow of the gas in the round bottom flask. Dropwise addition of the reacting solution is carried out. A centrally inserted refluxer in the flask also helps in some cases. One can use other necks for introducing temperature measuring devices (e.g. a thermometer), pH electrode, gas insertion device etc.

After preparing the appropriate precursor solutions by proper weighing/measuring and using appropriate volumes of the chemicals, reactions can be carried out following the steps given in the flow charts.

The nanoparticles are usually separated out from the liquids as a precipitate by using a centrifuge (approximately 2,000–3,000 r.p.m. speed). Precipitate should be washed using appropriate solvents like water, acetone, ethanol or methanol to remove unreacted chemicals or byproducts. Precipitate can be then dried to get sample in the form of powder.

Characterization of samples can be done using some of the techniques given below. The basic principles of characterization techniques were discussed in Chap. 7. All these techniques may not be available/required for high quality research but sometimes observation of colour change and techniques like optical absorption spectroscopy and X-ray diffraction are good enough to indicate the formation

Fig. 14.1 A typical set-up for synthesizing nanoparticles by chemical method



of nanomaterials. Transmission Electron Microscope (TEM), Scanning Electron Microscope (SEM), and photoluminescence are further useful techniques available at many places.

Usually for optical absorption investigations, nanoparticle sample dispersed in a liquid is used. For transmission or scanning electron microscopy (TEM or SEM) or other microscopies a drop of liquid is placed on suitable grid or substrate and liquid is allowed to evaporate before the measurements are carried out. In some other analysis like photoluminescence, X-ray diffraction (XRD) powder samples are preferred.

It may be noted that the following procedures are only illustrative and only to give one a flavour of nanotechnology experiments and are not necessarily the best possible procedures for obtaining particular nanoparticles. One can follow the literature and obtain better materials following some more stringent conditions or procedures.

14.2 Synthesis of Gold/Silver Nanoparticles

Aim: To synthesize metal nanoparticles of gold and silver.

14.2.1 Chemicals

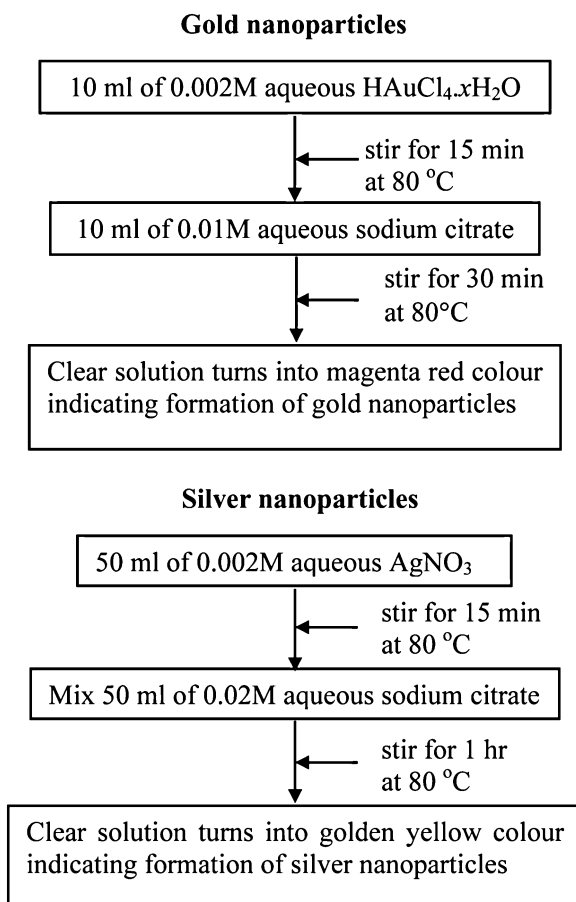
1. Chloro auric acid (HAuCl_4) for gold particles
2. Silver nitrate (AgNO_3) for silver particles
3. Trisodium citrate ($\text{C}_6\text{H}_5\text{O}_7\text{Na}_3$)
4. Double distilled water

14.2.2 Equipments

1. Round bottom flask
2. Magnetic stirrer cum heater
3. Optical absorption spectrometer ($\sim 250\text{--}700\text{ nm}$)

14.2.3 Synthesis Procedure

Procedures of gold and silver nanoparticles synthesis are given in the flow chart form. Synthesis can be carried out using the glass apparatus or set up as shown in Fig. 14.1.



14.2.4 Results

The magenta red and yellow colours for gold and silver solutions respectively indicate the formation of nanoparticles. Changing the concentrations, reaction time, temperature etc. one can obtain different shapes/sizes of the particles. This changes the solution colour or shades. There is large literature on these aspects. Typical photograph of gold and silver particles obtained using above procedure are shown in Fig. 14.2a.

Optical absorption spectra can be recorded using a simple absorption spectrometer discussed in Chap. 7. Figure 14.2b illustrates typical spectra obtained for the synthesis described here. It can be seen that peak for silver appears at approximately 396 nm and that for gold at about 530 nm. These are Surface Plasmon Resonance (SPR) peaks discussed earlier.

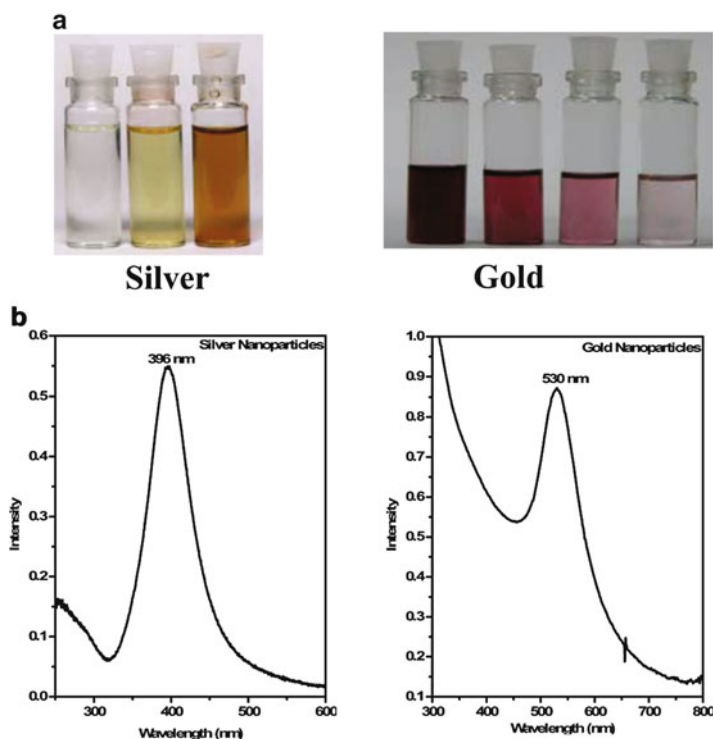


Fig. 14.2 (a) Gold and silver nanoparticles in solution. (b) Typical absorption spectra of silver and gold nanoparticles

14.3 Synthesis of CdS Nanoparticles

Aim: To synthesize CdS nanoparticles and cap them using thioglycerol.

14.3.1 Chemicals

1. Cadmium acetate ($(\text{CH}_3\text{COO})_2\text{Cd} \cdot 2\text{H}_2\text{O}$)
2. Thioglycerol ($\text{HSCH}_2\text{CH}_2\text{OH}$)
3. Sodium sulphide (Na_2S)
4. Ethanol

14.3.2 Equipments

1. Round bottom flask
2. Magnetic stirrer cum heater
3. Optical absorption spectrometer

14.3.3 Synthesis Procedure

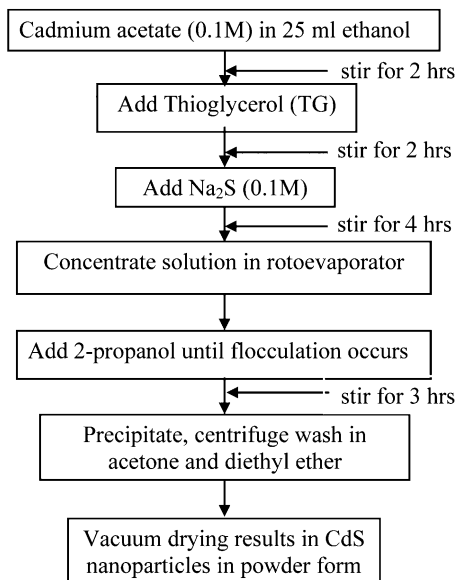
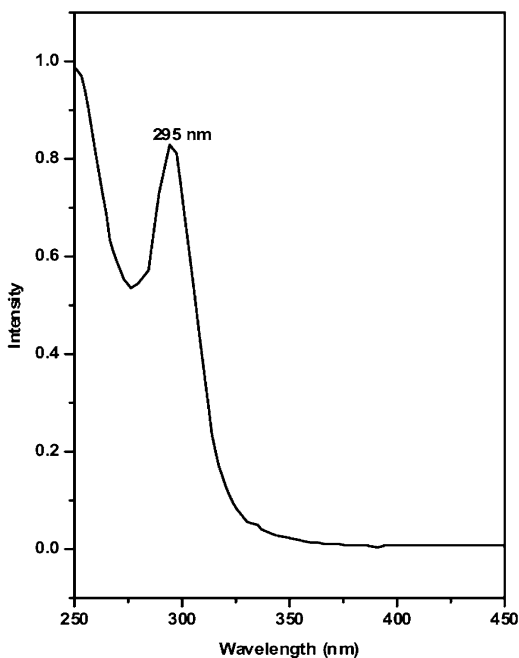


Fig. 14.3 UV-Vis absorption spectrum for CdS nanoparticles



14.3.4 Results

As can be seen from Fig. 14.3 an excitonic peak appears at a position much shifted compared to onset of bulk absorption at about 530 nm in case of CdS. Size determination of particles can be carried out using the procedure discussed below for ZnO or TiO₂ nanoparticles. In case of ZnO, procedure using optical absorption is discussed whereas for TiO₂ the XRD analysis is described. One can also obtain ZnS, PbS etc. nanoparticles by similar procedure using different precursors and analyze similarly.

14.4 Synthesis of ZnO Nanoparticles

Aim: To synthesize zinc oxide nanoparticles using a chemical route.

To cap the particles with thioglycerol.

To calculate the size of the synthesized particles using UV-Vis absorption spectrum.

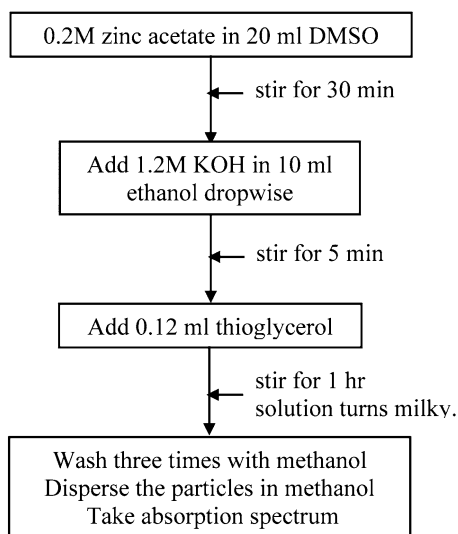
14.4.1 Chemicals

1. Zinc acetate dihydrate ($(\text{CH}_3\text{COO})_2\text{Zn}\cdot 2\text{H}_2\text{O}$)
2. Potassium hydroxide (KOH)
3. Thioglycerol (TG)
4. Di-methylene sulphoxide (DMSO)
5. Ethanol

14.4.2 Equipment

1. Round bottom flask
2. Magnetic stirrer cum heater
3. Optical absorption spectrometer

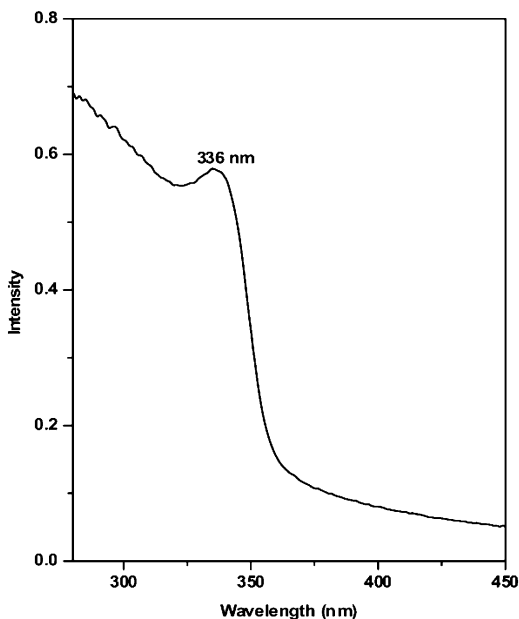
14.4.3 Synthesis Procedure



14.4.4 Results

An absorption spectrum as shown in Fig. 14.4 can be obtained. The absorption peak appears at 336 nm.

Fig. 14.4 UV-V is absorption spectrum for ZnO nanoparticles



Using the exciton peak position (336 nm) one can determine the energy gap ($E = hc/\lambda$). In this case it is 3.7 eV. This can be used to determine the average particle size using effective mass approximation (or tight banding calculations). Effective mass approximation formula is as follows

$$E = E_g + \frac{\hbar^2 \pi^2}{2R^2} \left\{ \frac{1}{m_e} + \frac{1}{m_h} \right\} - \frac{1.8e^2}{4\pi\epsilon_0 R} + \text{smaller term} \quad (14.1)$$

where, E is band gap of the synthesized particle, E_g – bulk band gap of ZnO (3.3 eV), R – radius of the particle, m_e – effective mass of electron (for ZnO it is $0.28 m_0$), m_h – effective mass of hole (for ZnO it is $0.49 m_0$), ϵ – dielectric constant of material (for ZnO it is 9.1), ϵ_0 – permittivity of free space, and h is Planck's constant $h = h/2\pi$.

The last small term in the Eq. (14.1) can be neglected. Substituting above values in equation, we obtain particle size as 4.6 nm.

14.5 Synthesis of TiO₂ Nanoparticles

Aim: To synthesize titanium dioxide (TiO₂) nanoparticles using a chemical route.

To determine the phase and calculate the size of the synthesized particles using X-ray diffraction (using Scherrer formula).

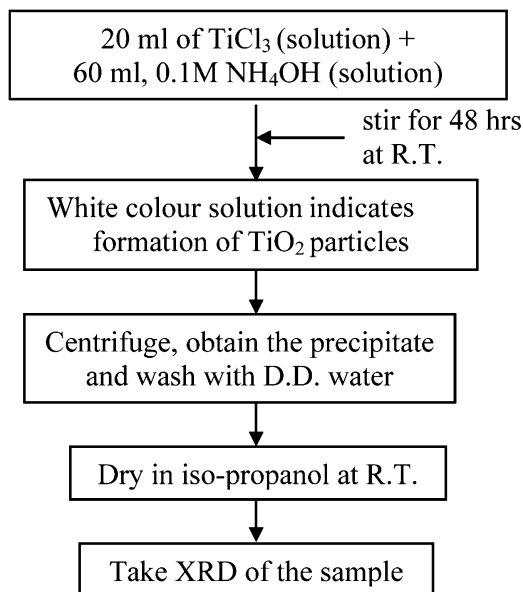
14.5.1 Chemicals

1. Titanium trichloride (TiCl_3)
2. Ammonium hydroxide (NH_4OH)
3. 2-propanol ($(\text{CH}_3)_2\text{CHOH}$)
4. Double distilled (D.D.) water

14.5.2 Equipment

1. Round bottom flask
2. Magnetic stirrer cum heater
3. Optical absorption spectrometer

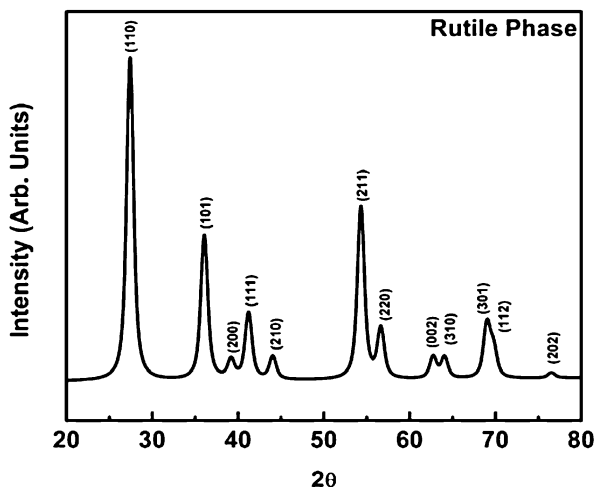
14.5.3 Synthesis Procedure



14.5.4 Results

A typical diffraction pattern of dry powder sample obtained using $\text{CuK}\alpha$ radiation ($\lambda = 0.154 \text{ nm}$) appears as shown in Fig. 14.5. Comparing with JCPD card no

Fig. 14.5 The X-ray diffraction pattern of TiO_2



21-1276, it was found that the particles are in rutile phase. The various peaks are identified with different planes. One can determine the particle size using Scherrer formula as

$$d = \frac{0.9\lambda}{\beta \cos \theta_B} \quad (14.2)$$

where, d is particle diameter, λ – wavelength of X-rays, β – full width at half maximum (FWHM) of diffraction peak, and θ_B is Bragg diffraction peak angle.

The size of the particles as calculated using Scherrer formula turns out to be ~ 7 nm in this case.

14.6 Synthesis of Fe_2O_3 Nanoparticles

Aim: To synthesize iron oxide particles of different shapes.

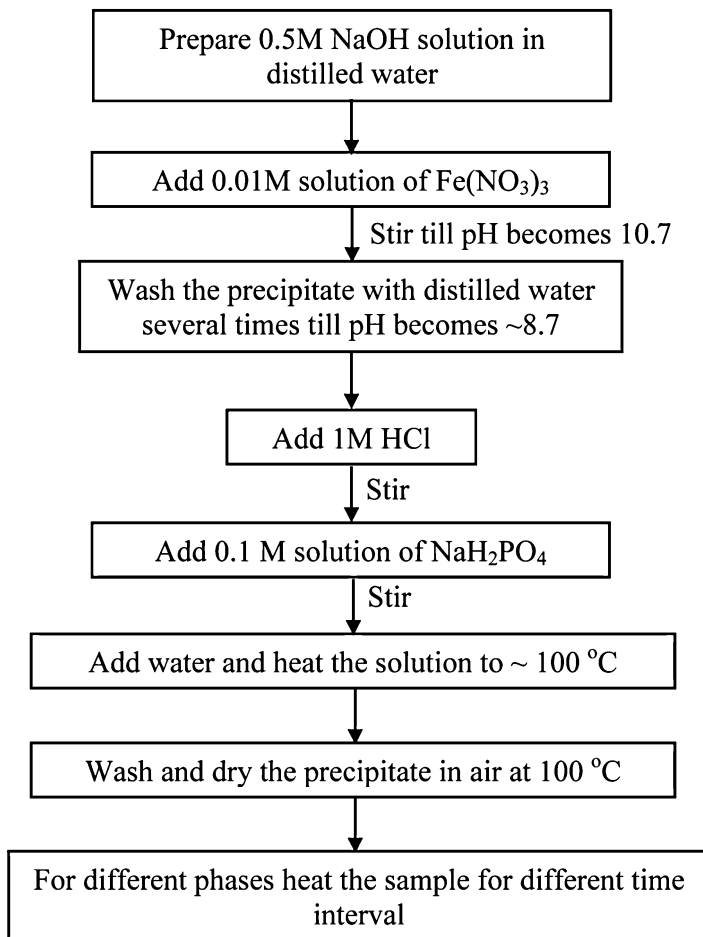
14.6.1 Chemicals

1. Sodium hydroxide (NaOH)
2. Iron chloride (FeCl_3)
3. Sodium hexametaphosphate (NaH_2PO_4)
4. Double distilled water

14.6.2 Equipment

1. Round bottom flask
2. Magnetic stirrer cum heater

14.6.3 Synthesis Procedure



14.6.4 Results

In this case different shapes of Fe_2O_3 particles are obtained (Fig. 14.6). X-ray diffraction analysis can be carried out to determine phase. It will be found that these

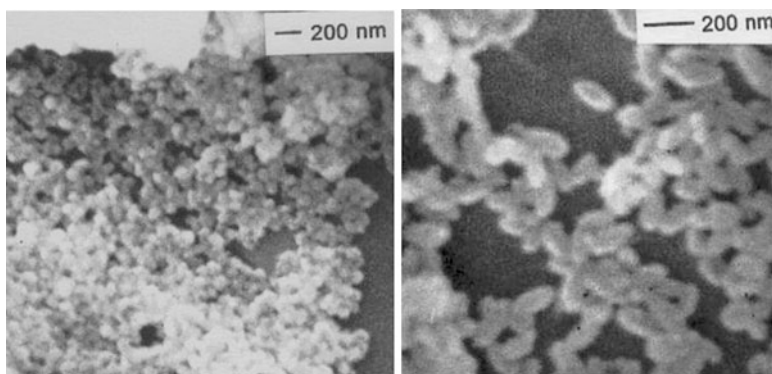


Fig. 14.6 SEM images of Fe_2O_3 nanostructures

particles can be converted into magnetic particles by annealing them in H_2 gas at $350\text{--}370^\circ\text{C}$ for 2 h. The $\alpha\text{-Fe}_2\text{O}_3$ particles convert into $\gamma\text{-Fe}_2\text{O}_3$ without substantial change of particle size. Note that here particles have ~ 200 nm size. However by changing the reaction conditions, smaller particles can be synthesized.

14.7 Synthesis of Porous Silicon

Aim: To obtain porous silicon and study its photoluminescence.

14.7.1 Materials

1. Silicon substrate (p-type) as anode
2. Graphite or platinum electrode as cathode
3. Hydrofluoric acid (48 %)
4. Ethanol, acetone and distilled water for cleaning substrates

14.7.2 Equipment

1. Teflon holder for electrochemical etching of substrate.
2. Voltmeter and ammeter for adjusting voltage and current.
3. Power supply (24 V, 30 mA)
4. Electrochemical cell (teflon or plastic) (Fig. 14.7)

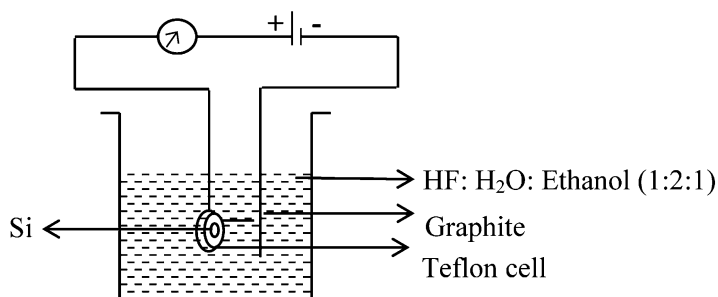


Fig. 14.7 Experimental set-up for porous silicon formation

14.7.3 Experimental Procedure

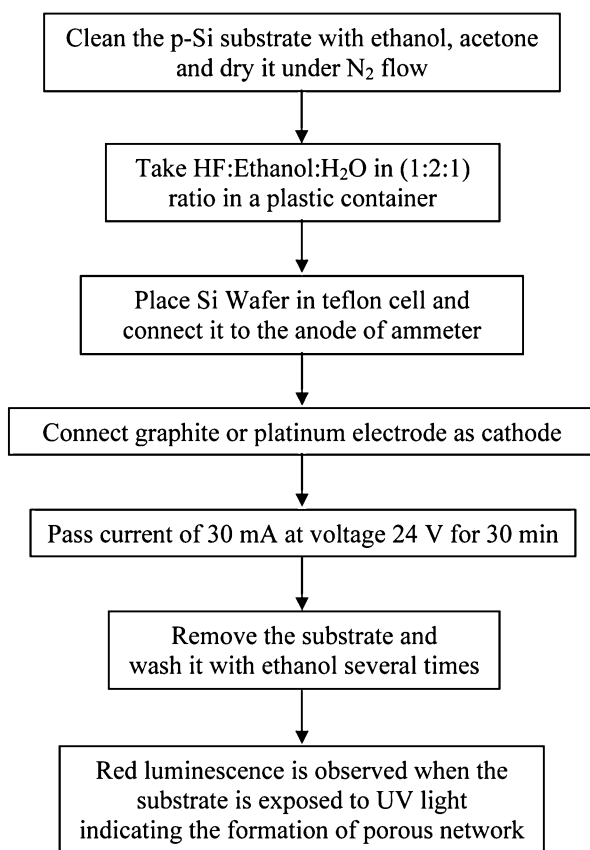
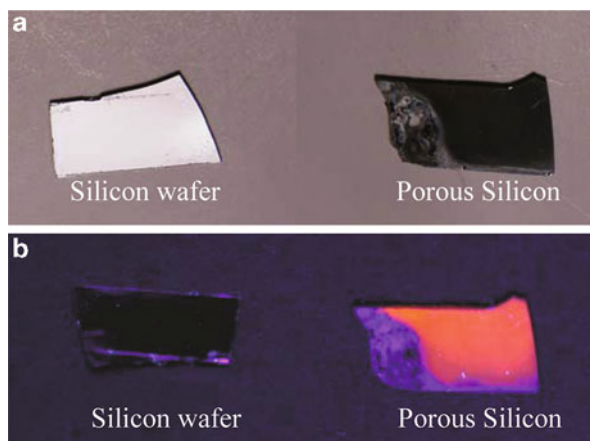


Fig. 14.8 (a) Silicon wafer and porous silicon in visible light and (b) under UV light



14.7.4 Results

A simple method to check whether silicon has become porous or not, is to just expose the dried sample to a UV radiation. A silicon piece would not show any luminescence but porous silicon glows in the visible range from blue to red colour depending upon the silicon substrate and conditions of the preparation. Figure 14.8 illustrates the photographs of silicon substrate and porous silicon sample prepared under the conditions mentioned above and under UV radiation.

One can also record a luminescence spectrum as illustrated in Fig. 14.9 which clearly shows that there is a broad luminescence band located at about 700 nm in the present case.

If one can inspect the sample under a scanning electron microscope, it will appear porous. However depending upon the substrate and processing of silicon substrate one may get different morphologies of porous silicon. Figure 14.10 illustrates the SEM image of the p-Si substrate used in these experiments.

14.8 Introductory Photolithography

14.8.1 Background

Photo-litho-graphy means light-stone-writing in *latin*.

Photolithography is the photographic process to transfer circuit patterns on a semiconductor wafer. This process is useful to fabricate variety of micro components and systems such as electronic circuits, air bag accelerometers, inertial guidance sensors, pressure and flow sensors. Using photolithography a pattern upto

Fig. 14.9 Excitation (---) and emission (—) spectra of porous silicon

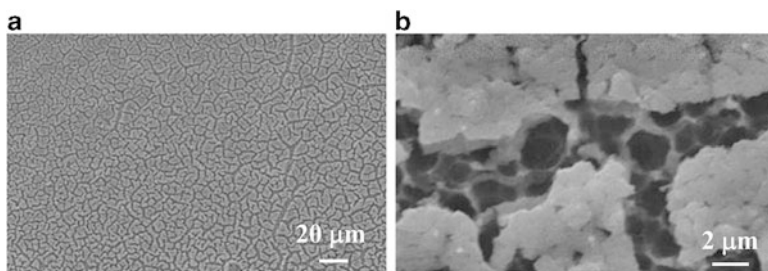
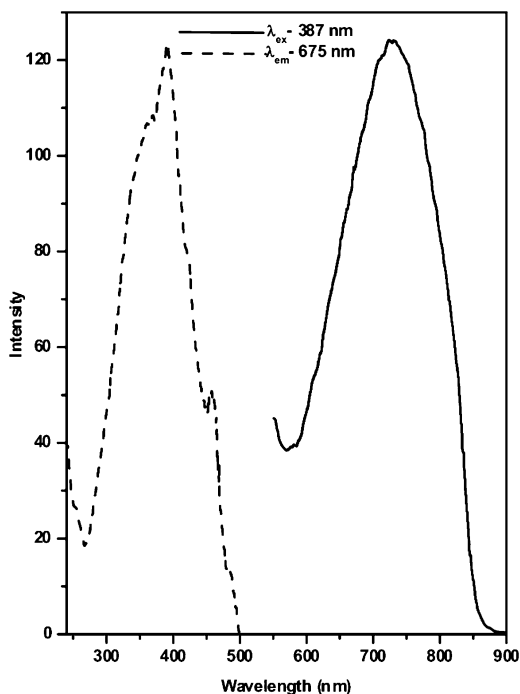


Fig. 14.10 (a) SEM image of porous silicon and (b) SEM magnified view porous silicon sample in (a)

hundreds of nanometer can be prepared on semiconductor surfaces. Size of the pattern is usually in the range of 10–100 μm . Pattern is designed using software such as Autocad. Usually, the designed pattern is bigger and hence is scaled down 100 times. The patterned design of circuit is then printed on a mask (also called as master) in a photography laboratory. Using a single mask, large number of patterns can be generated on semiconductor chips. This is done in various steps such as wafer cleaning, metallization, photosensitive material coating, soft baking, mask

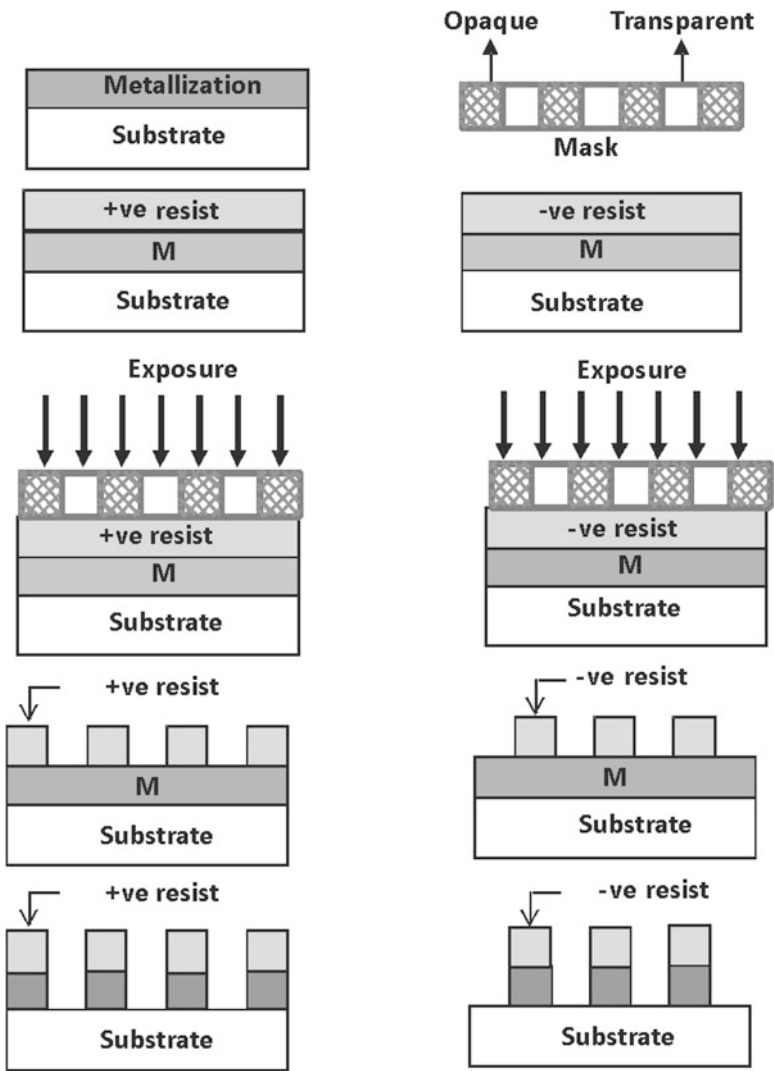


Fig. 14.11 Schematic diagram of photolithography process

alignment, UV exposure, developing, hard baking and etching. Photolithography process for transferring the pattern is shown in Fig. 14.11.

In the first step, the wafer is cleaned chemically to remove organic, ionic and metallic impurities. After cleaning, silicon dioxide which serves as a barrier layer is deposited on the surface of the wafer. Then photoresist is applied to the surface

of the wafer. The photoresist material is of two types viz. positive and negative. In positive resists, when UV light is exposed, chemical structure of exposed part of resist changes and becomes soluble in the developer. The exposed resist is then washed away by the developer solution, leaving behind windows of the bare underlying material. Negative resists behave in just the opposite manner. When UV light is exposed to resist, exposed part is polymerized and is difficult to dissolve in developer. Therefore, the negative resist remains on the surface wherever it is exposed, and the developer solution removes only the unexposed portions. These photoresists are generally organic polymers which are sensitive to light. Soft-baking plays critical role because photoresist coatings become photosensitive, or imageable, only after softbaking. Image is transferred from master to substrate by exposing it to UV radiation. Mask consists of clear and opaque areas, which allows exposure of resist material through clear part. Exposure alters the photoresist material chemically and alters its solubility, which is dissolved in suitable solvent called developer. Finally wafer is washed with etching solution to remove the surface not protected by photoresist.

Aim: To transfer the design pattern on given substrate.

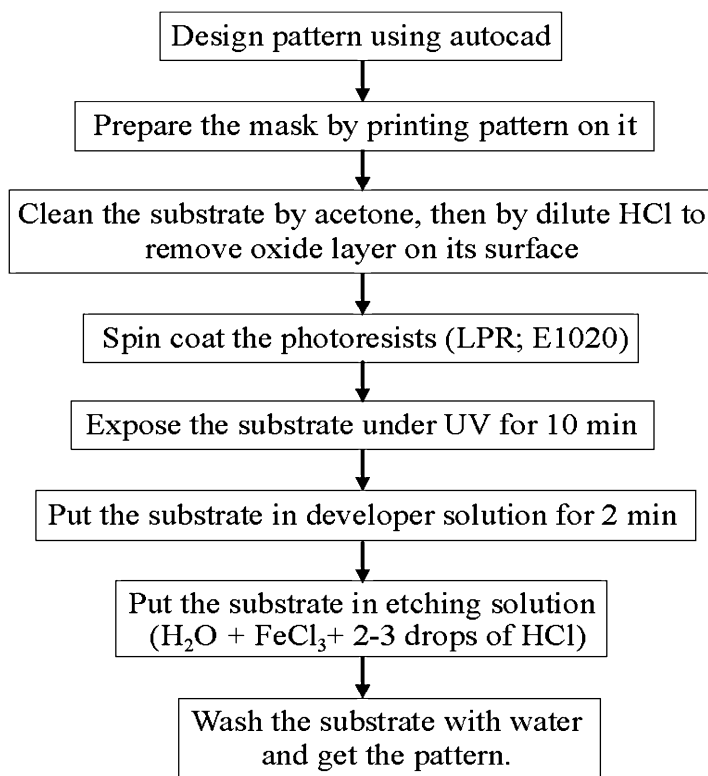
14.8.2 Chemicals

1. Silicon substrate
2. Acetone
3. Negative photoresist (LPR: E1020 of Kodak)
4. Developer (of Kodak)
5. Hydrogen peroxide (H_2O_2)
6. Ferrous chloride (FeCl_3)
7. Hydrochloric acid (HCl)

14.8.3 Equipment

1. Spin coating machine
2. UV lamp
3. Autocad software
4. Mask

14.8.4 Experimental Procedure



14.9 Introductory Nano (Soft) Lithography Using PDMS

A number of lithography techniques like electron beam, ion beam, X-ray etc. are possible to use in order to pattern nanometric features. However they are too expensive for large scale production. Recently some new techniques have been developed known as *Soft Lithography*. Soft-litho means use of 'soft materials' i.e. polymers.

An elastomer of PDMS (poly dimethyl siloxane) of Sylgard 184, Dow Corning is used to prepare microstructures. There are different techniques used in soft lithography: Microcontact printing (μ CP), Replica molding (REM), Microtransfer molding (μ TM), Micromolding in capillaries (MIMIC), and Solvent assisted micromolding (SAMIM).

This experiment involves synthesis of some micron size particles and then patterning them into small channels using micromolding in capillaries (MIMIC) technique. First we will discuss the synthesis of silica particles.

Aim: To synthesize silica nanoparticles using a chemical route.

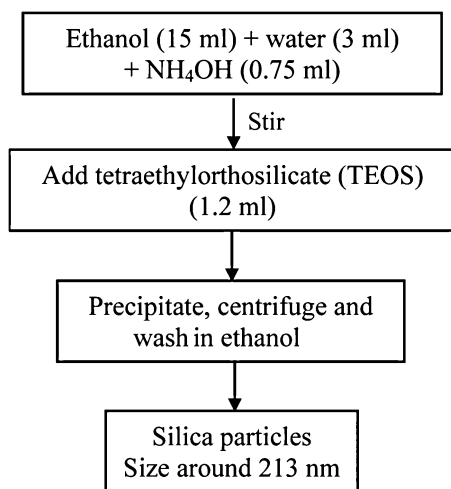
14.9.1 Chemicals

1. Tetraethyl orthosilicate (TEOS) $[\text{Si}(\text{OC}_2\text{H}_5)_4]$
2. Ammonium hydroxide (NH_4OH)
3. Distilled water (H_2O)
4. Ethanol ($\text{C}_2\text{H}_5\text{OH}$)

14.9.2 Equipment

1. Round bottom flask
2. Magnetic stirrer

14.9.3 Synthesis Procedure



14.9.4 Results

One may record SEM or TEM of the samples just by putting a drop of sample on suitable metal grid (in case of TEM) or substrate for SEM analysis. Here some TEM analysis is shown (Fig. 14.12).

Now we will pattern these particles using *Micromolding in Capillaries* technique. Various steps for the patterning are shown in Fig. 14.13. In this technique, first a PDMS mold is placed on the cleaned substrate (see Chap. 9 for details about

Fig. 14.12 TEM image of silica nanoparticles

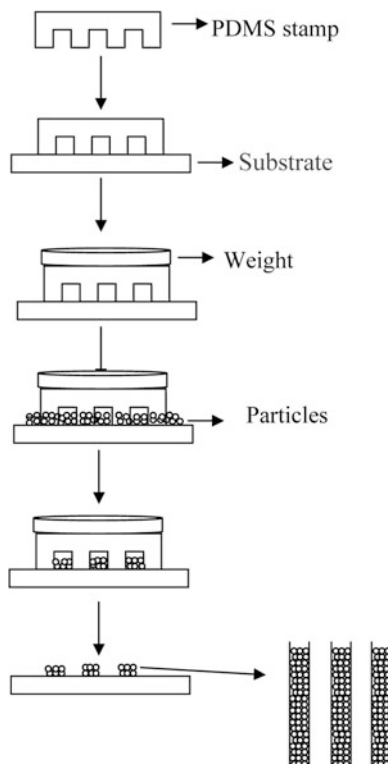
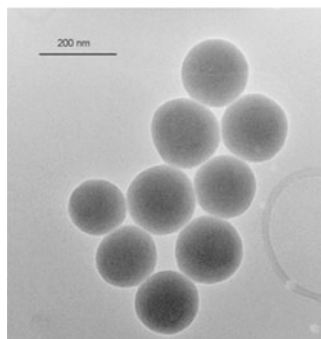


Fig. 14.13 Schematic diagram of micromolding in capillaries technique

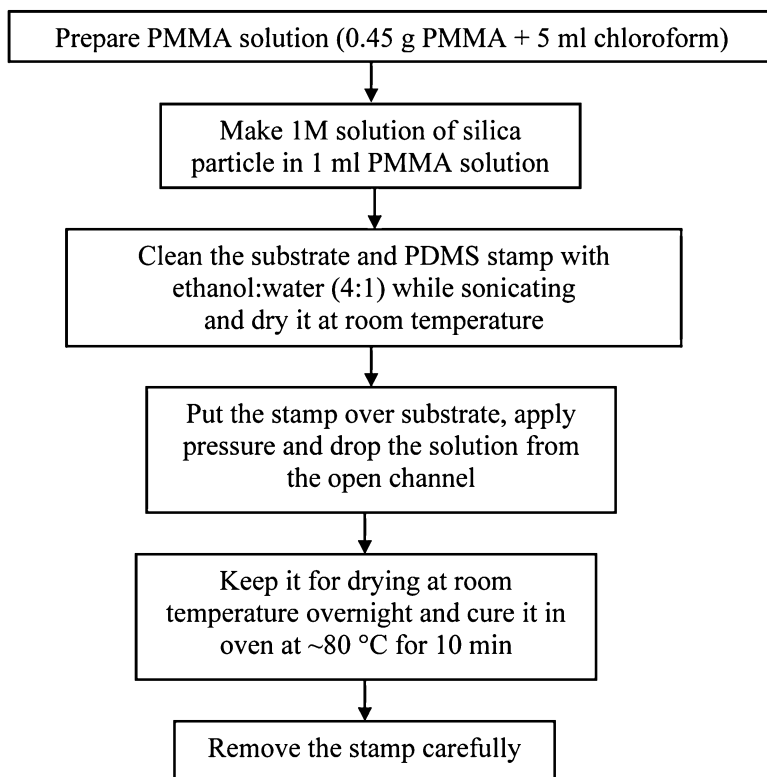
PDMS stamp). Then liquid drop is placed at the open ends of the channel which fills the channels by capillary action. Cure the sample in oven at $\sim 80^\circ\text{C}$ and remove the stamp carefully. Resulting structures are usually thinner than the height of the groove.

Aim: To pattern silica particles by MIMIC (micromolding in capillary) technique.

14.9.5 Materials and Equipments

1. Silica particles
2. Poly methyl metha acrylate (PMMA)
3. Chloroform (CHCl_3)
4. Ethanol ($\text{C}_2\text{H}_5\text{OH}$)
5. Poly dimethyl siloxene (PDMS) stamp
6. Substrate (silicon wafer or glass)

14.9.6 Experimental Procedure



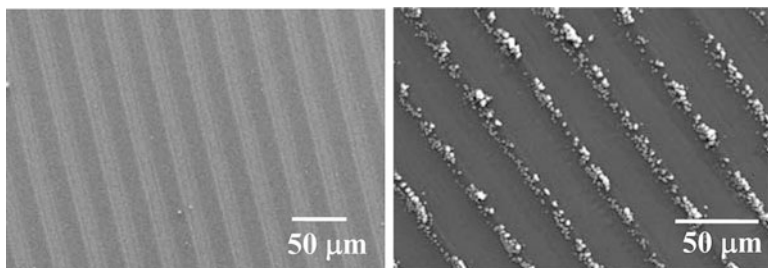


Fig. 14.14 SEM images of (a) pattern and (b) patterned silica particles

14.9.7 Experimental Set-Up

14.9.8 Results

SEM image (Fig. 14.14) shows (a) PMMA is filled in the grooves of pattern. The groove width is 15 μm and the spacing between the grooves is 10 μm . (b) Pattern of silica particles.

14.10 Fabrication of Porous Alumina or Anodized Alumina (AAO) Template

Aim: To fabricate uniform and ordered pores in alumina using electrochemical method.

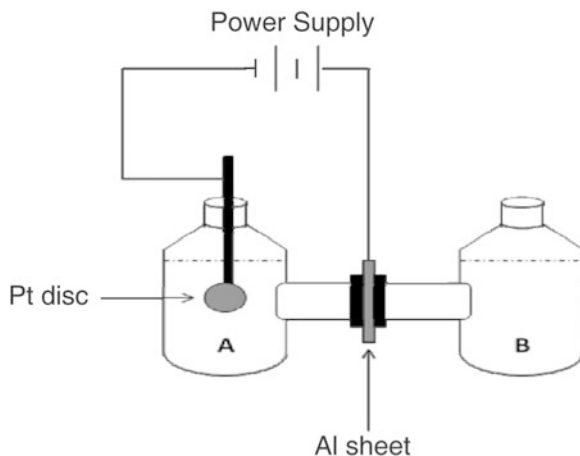
14.10.1 Chemicals

1. Oxalic acid
2. Perchloric acid
3. Ethanol
4. Copper chloride
5. Phosphoric acid
6. Methylene orange

14.10.2 Equipment

Electrochemical set up with power supply (see Fig. 14.15)

Fig. 14.15 Apparatus for preparation of AAO templates



14.10.3 Fabrication Procedure

Highly ordered porous alumina templates are prepared by a modified two-step anodizing process. Ultrapure aluminum foil (99.99 %) is used.

1. Cut the aluminum foil into $7.5\text{ cm} \times 4.5\text{ cm}$ sheets.
2. Etch the aluminum sheets in alkaline solution for 2 min and then rinse with ethanol, dry and anneal at 450°C for 4 h.
3. *Electropolishing*: Aluminum sheet is used as anode and platinum plate as cathode. 1:4 mixture of perchloric acid (HClO_4) and ethanol is used as the electrolyte. This solution is kept in a beaker. The electrolysis is done for one minute at the operating voltage of 18 V.
4. *First oxidation process*: Parameters for first anodization process are as follows:

Anode: Aluminum sheet

Cathode: Pt

Electrolyte: Oxalic acid in bottle A and Millipore water in the bottle B. Pt disc is placed in the oxalic acid solution (bottle A) and Al sheet is fixed between the two bottles as shown in Fig. 14.15.

Voltage: 50 V

Current: 0.03 and 0.07 Amp

Time: 1.5 h

5. *Reduction process*: Remove the platinum electrode from oxalic acid. Remove oxalic acid from the bottle and rinse the bottles with Millipore water. Let the Al sheet remain fixed between the two bottles. Place a mixture of chromic acid (1.8 wt%) and phosphoric acid (6 wt%) in bottle A and Millipore water in bottle B. This process is called as reduction process and is used to remove the oxide layer from the surface so that the inside part of alumina sheet can be made porous. The bottles are kept in water bath at 60°C for 2 h.

6. *Second oxidation step:* Repeat the oxidation process with the same parameter as in Step 4.
7. *Removal of alumina from the other side of the sheet:* Put saturated solution of copper chloride in the bottle B and Millipore water in bottle A. Wait for 1 min. The surface of the template sheet should be carefully watched through bottle A. Once the 75 % of the exposed portion of the template turns black, remove the solutions immediately from the bottles. This step is very critical and should be done very carefully. If the copper chloride solution is allowed to remain in the solution for long time, it will lead to breakage of template. This procedure removes residual alumina from the other side of the AAO sheet. After slowly rinsing the template sheet and bottles with Millipore water, it is observed that the sheet has become translucent.
8. *Pore widening at the bottom of the AAO template:* Put 5 % H_3PO_4 solution in bottle B to open the bottom pore of anodic alumina membrane and dilute aqueous methylene orange solution in bottle A. The bottles are kept until half of the solution in bottle A turns red. Then the template is carefully removed and rinsed with Millipore water. Let the template remain in a beaker containing Millipore water overnight and remove it from water next day.

14.10.4 Results

Typical pore structure obtained using above procedure is depicted in Fig. 14.16. The images are obtained using a scanning electron microscope. Optimization of the parameters and solution concentrations can be made in order to vary the pore size and control the ordered structure. Such pores can be used to deposit different materials in them and obtain 1-D nanomaterials.

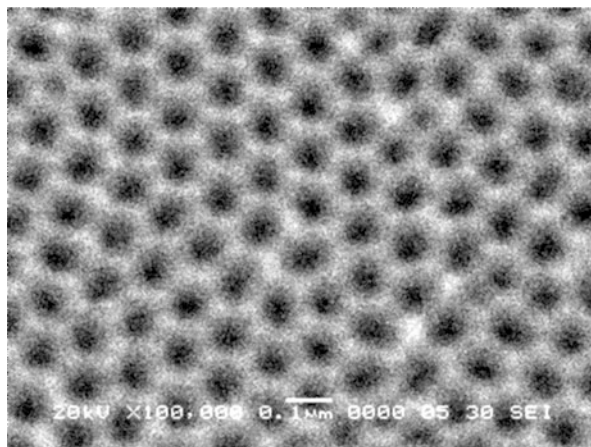


Fig. 14.16 SEM image of AAO template

Further Reading

- A.G. Cullis, L.T. Canham, P.D.J. Calcott, The structural and luminescence properties of porous silicon. *J. App. Phys.* **82**(3), 909 (1997)
- B.D. Gates, Q. Xu, J.C. Love, D.B. Wolfe, G.M. Whitesides, Unconventional nanofabrication. *Ann. Rev. Mater. Res.* **34**, 339 (2004)
- S. Link, M.A. El-Sayed, Optical properties and ultrafast dynamics of metallic nanocrystals. *Ann. Rev. Phys. Chem.* **54**, 331 (2003)
- O. Masala, R. Sheshadri, Synthesis routes for large volumes of nanoparticles. *Ann. Rev. Mater. Sci.* **34**, 41 (2004)
- C.B. Murray, C.R. Kagan, M.G. Bawendi, Synthesis and characterization of monodisperse nanocrystals and close-packed nanocrystal assemblies. *Ann. Rev. Mat. Sci.* **30**, 545 (2000)

Appendix I

PERIODIC TABLE SYMBOLS OF ELEMENTS AND THEIR ATOMIC NUMBERS

PERIODIC TABLE																	
SYMBOLS OF ELEMENTS AND THEIR ATOMIC NUMBERS																	
IA		VIII A										IB		IIB		Nobel Gases	
1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16	17	18
H	He	Li	Be	B	C	N	O	F	Ne	Na	Mg	Al	Si	P	S	Cl	Ar
19	20	21	22	23	24	25	26	27	28	29	30	31	32	33	34	35	36
K	Ca	Sc	Ti	V	Cr	Mn	Fe	Co	Ni	Cu	Zn	Ga	Ge	As	Se	Br	Kr
37	38	39	40	41	42	43	44	45	46	47	48	49	50	51	52	53	54
Rb	Sr	Y	Zr	Nb	Mo	Tc	Ru	Rh	Pd	Ag	Cd	In	Sn	Sb	Te	I	Xe
55	56	57	72	73	74	75	76	77	78	79	80	81	82	83	84	85	86
Cs	Ba	La	Hf	Ta	W	Re	Os	Ir	Pt	Au	Hg	Tl	Pb	Bi	Po	At	Rn
87	88	89	104	105	106	107	108	109	110	111	112	(113)	(114)	(115)	(116)	(117)	(118)
Fr	Ra	Ac	Rf	Db	Sg	Bh	Hs	Mt	Ds	Rg	Cn	(113)	(114)	(115)	(116)	(117)	(118)
(119)	(120)	(121)	(154)	(155)	(156)	(157)	(158)	(159)	(160)	(161)	(162)	(163)	(164)	(165)	(166)	(167)	(168)

LANTHANIDES

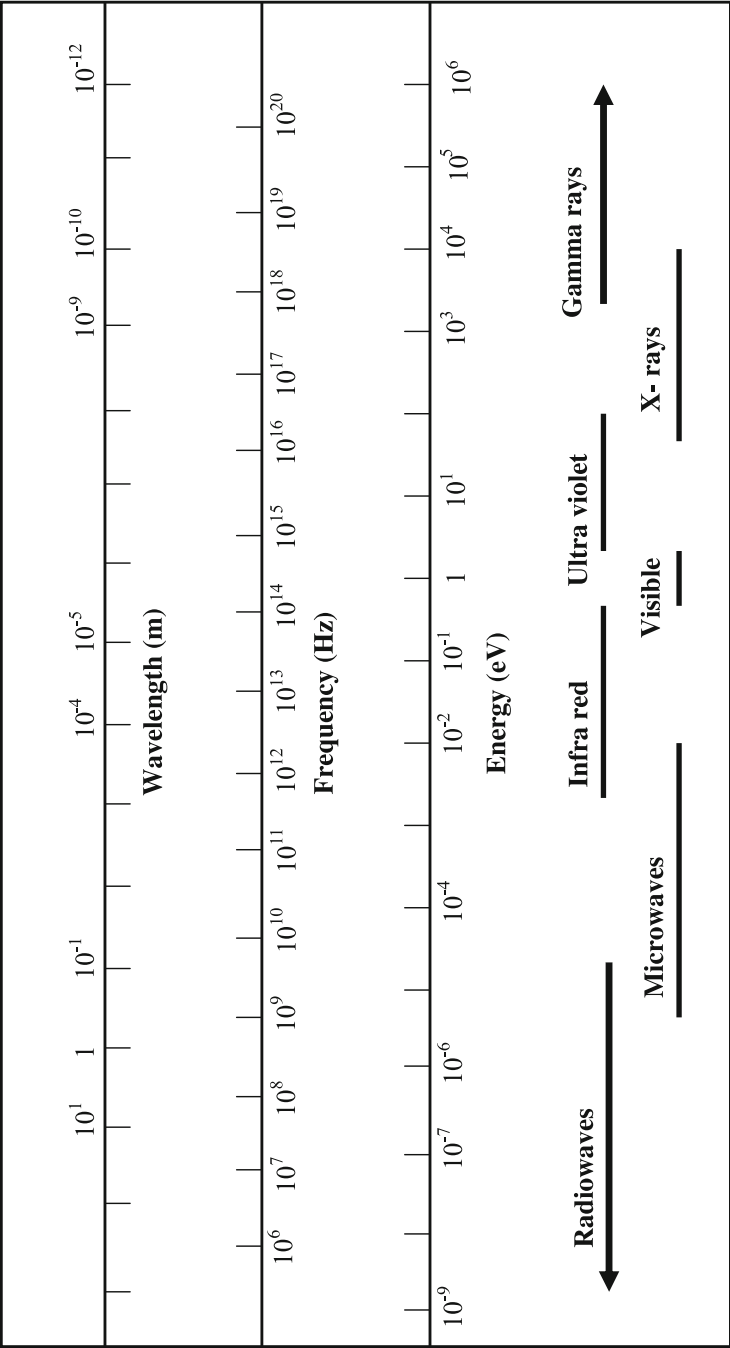
58	Ce	59	Pr	60	Nd	61	Pm	62	Sm	63	Eu	64	Gd	65	Tb	66	Dy	67	Ho	68	Er	69	Tm	70	Yb	71	Lu
----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----

ACTINIDES

90	Th	91	Pa	92	U	93	Np	94	Pu	95	Am	96	Cm	97	Bk	98	Cf	99	Es	100	Fm	101	Md	102	No	103	Lr
----	----	----	----	----	---	----	----	----	----	----	----	----	----	----	----	----	----	----	----	-----	----	-----	----	-----	----	-----	----

Appendix II

ELECTROMAGNETIC SPECTRUM



Appendix III

List of Fundamental Constants

Physical quantity	Symbol	Value	SI	CGS
Speed of light in vacuum	c	2.9979	10^8 ms^{-1}	$10^{10} \text{ cm. s}^{-1}$
Electronic charge	e	1.6021	10^{-19} C	–
		4.8032	–	10^{-10} esu
Boltzmann constant	k_B	1.3806	10^{-23} JK^{-1}	$10^{-16} \text{ erg. K}^{-1}$
Planck's constant	h	6.6262	10^{-34} J s	10^{-27} erg. s
Electron rest mass	m_e	9.1095	10^{-31} kg	10^{-28} g
Proton rest mass	m_p	1.6726	10^{-27} kg	10^{-24} g
One electron volt	eV	1.6021	10^{-19} J	10^{-12} erg
Avogadro's number	N_A	$6.0221 \times 10^{23} \text{ mol}^{-1}$	–	–
Permittivity of free space	ϵ_0	–	$10^7/4\pi c^2 \text{ Fm}^{-1}$	1
Permeability of free space	μ_0	–	$4\pi \times 10^{-7} \text{ Na}^{-2}$	1
Bohr magneton	μ_B	9.2741	$10^{-24} \text{ J T}^{-1}$	$10^{-21} \text{ erg G}^{-1}$
Bohr radius (radius of hydrogen atom)	r_0	5.2917	10^{-11} m	10^{-9} cm
Atomic mass constant	m_u	1.6605	10^{-27} kg	10^{-24} g

Appendix IV

Vacuum Techniques

Vacuum is the space devoid of any gas, liquid, solid or any particles. Vacuum is needed for various purposes such as packaging of food, materials deposition and processing, analysis equipment and many other applications. Extent of vacuum is measured with reference to normal atmospheric pressure at sea level i.e. one atmosphere. This in turn is measured as the pressure exerted by air on unit area of a flat surface. Pressure is defined as force per unit area. The SI unit for pressure is Pascal (Pa) defined as N/m^2 . At sea level mercury column height is 760 mm. This is usually used to measure pressure. Traditionally, in different parts of the world different groups have been using different units of vacuum. However, the reference is always atmospheric pressure at sea level and vacuum is the pressure smaller than atmospheric pressure. Some common units used to measure pressure are torr, Pa and mbar. They are inter-related as follows:

$$1 \text{ atmosphere} = 760 \text{ torr} = 1.013 \times 10^5 \text{ Pa}$$

$$1 \text{ torr} = 1 \text{ mm of Hg} = 10^{-3} \text{ mbar}$$

$$1 \text{ bar} = 750 \text{ torr}$$

$$1 \text{ Pa} = 7.5 \text{ mtorr}$$

For the sake of convenience the vacuum ranges are roughly identified as follows (note that the demarcation lines between different ranges are rather blurred).

Low vacuum	10^5 Pa (750 torr) to 3.3×10^3 Pa (25 torr)
Medium vacuum	3.3×10^3 Pa (25 torr) to $\sim 10^{-1}$ Pa (7.5×10^{-4} torr)
High vacuum	10^{-1} Pa (7.5×10^{-4} torr) to 10^{-4} Pa (7.5×10^{-7} torr)
Very high vacuum	10^{-4} Pa (7.5×10^{-7} torr) to 10^{-6} Pa (7.5×10^{-9} torr)
Ultra high vacuum	10^{-6} Pa (7.5×10^{-9} torr) to 10^{-10} Pa (7.5×10^{-13} torr) or below

From the kinetic theory of gases, mean free path of gas molecules λ is

$$\lambda = \frac{1}{\sqrt{2}\pi d^2 n}$$

where d is diameter of a molecule and n is number density of molecules.

Flux of gas molecules striking a surface, per unit time and on unit area are given as

$$\phi = 3.513 \times 10^{22} \frac{P}{(MT)^{1/2}} \text{ molecules/cm}^2.\text{s}$$

where P is pressure, M is weight of molecule and T is temperature.

Some molecules may stick to the surface (adsorbed) and some may get released back or simply bounce back to the ambient. In the vacuum technology terminology, when a surface releases gas molecules in the vacuum system it is known as ‘outgassing’.

Vacuum System

In order to obtain a vacuum in some chamber (usually made of glass, steel or some other metal or alloy) it is necessary to connect it through various pipe lines, vacuum valves, traps (to avoid gases or pump fluids entering the vacuum chamber) to different pumping devices called ‘vacuum pumps’. It may be necessary to use one or more pumps to achieve an adequate vacuum which is measured by using ‘vacuum gauges’. In order to achieve desired vacuum it is necessary to pay enough attention to the materials used not only for making the vacuum chambers but also to various sealings, valves, traps, vacuum pumps, vacuum gauges, fluids etc. which may be used to obtain, retain, or measure vacuum. The commonly used chamber materials are glass, stainless steel and in some cases aluminum alloys. Some materials used to connect ports with some modules or close some chamber parts are known as flanges. They are sealed using rubber gaskets (viton), gold, tin or indium wires. Often there are glass to metal connectors or ceramics to metal connectors needed to make electrical connections for performing different types of operations in a vacuum chamber. The detailed discussion on materials is beyond the scope of this book and can be found in the references given at the end of the appendix.

Vacuum Pumps

A variety of vacuum pumps have been designed with varying pumping speeds and are in use depending upon the vacuum requirement i.e. whether low vacuum, high vacuum or some other range is required. Before we proceed to the vacuum pumps let us define the pumping speed.

Pumping speed (S) of a vacuum system is defined as the volume (V) or amount of a gas removed from the pump in time t .

$$S = V/t$$

Throughput (Q) is defined as the quantity of a gas in units of pressure and amount of gas passing from some point in certain time. Throughput is also known as 'gas load'.

$$\begin{aligned} Q &= \text{pressure} \times \text{volume/time} \\ &= \text{torr} \times \text{litres/sec or Pa} \times \text{m}^3/\text{sec} \end{aligned}$$

Throughput is related to the pumping speed as

$$Q = P.V/t = P.S$$

The vacuum pumps can be divided into two types: (1) Gas transfer type and (2) Entrapment.

Major gas transfer type of pumps are rotary pump, root's pump, diffusion pump and turbo-molecular pump. Major entrapment type of pumps are sorption pump, ion pump, sublimation pump and cryo pump.

In the following we shall outline the principles of some of these pumps in brief. More details should be found in the references given at the end of the appendix.

Rotary Vane Pump

This type of pump is used to obtain the vacuum from atmospheric pressure to ~ 0.13 Pa. The pump body is usually made of steel. The pump consists mainly of a disk type rotor (see Fig. A.1) which rotates off-centric with respect to the chamber centre. Two vanes attached to the centre of the rotor make tight seal between the rotor and the chamber (stator) during rotation, thereby sweeping a gas from the inlet side to the outlet side without much of back-streaming of the gas. Due to compression of the gas on the outlet side, pressure increases which is sufficient to open a valve and let the gas be thrown out. The pump body is immersed in oil.

The rotors rotate at $\sim 2,000$ to $5,000$ rpm and small clearance needs to be kept between rotor and the stator which limits the vacuum that can be achieved. Most of the rotary pumps are also provided with 'gas ballast' arrangement. There are other versions of rotary pump like rotary plunger pump and trochoid pump. One may also use 'double stage or two stage' rotary pump.

Fig. A.1 Schematic diagram of a rotary pump

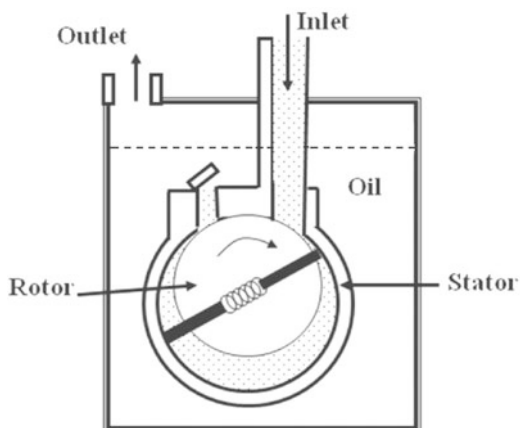
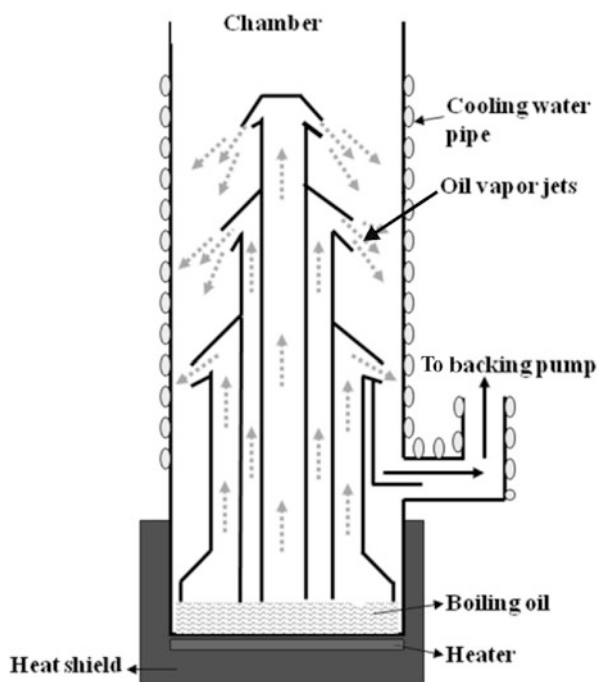


Fig. A.2 Schematic diagram of a diffusion pump



Diffusion Pump

The usual operating range of a diffusion pump is $\sim 10^{-1}$ to $\sim 10^{-3}$ Pa. However with the use of proper oil and good design it can be extended even to $\sim 10^{-7}$ Pa. However, like a rotary pump, diffusion pump cannot operate just by itself. It needs a backing pump like a rotary pump to pump away the gas at high pressure. A diffusion pump consists (see Fig. A.2) of an outer cylinder in which an assembly of a chimney-like

structure is kept. There is oil at the bottom of the pump body which can be heated by an external heater. There is also a metallic water tubing which surrounds the body of the diffusion pump. The chamber to be evacuated to high vacuum is first pumped down to $\sim 10^{-1}$ Pa using a rotary or other suitable pump and then the valve of diffusion pump is opened to the chamber to be evacuated. The heater when started evaporates the oil which passes at very high velocity through the nozzles in the downward direction. The gas molecules from the chamber diffuse in the oil jets and are forced downwards. The oil is cooled by the water flowing in the tubing surrounding the pump body and gas molecules rush in the outlet region where they are pumped by a pump like rotary pump. The process can keep on repeating to gain and continue the desired vacuum state inside the vacuum chamber.

Sorption Pump

It is one of the simplest pumps with the capability of achieving vacuum, starting from the atmospheric pressure upto $\sim 10^{-3}$ Pa. Pumping speed is initially upto $\sim 10^{-1}$ Pa very high and then slows down. The pump works on the principle of adsorption of gases on cooled surfaces. In a cylindrical steel chamber highly porous material like molecular sieves, zeolite or charcoal is filled. The pump body is cooled by liquid nitrogen by filling it in the outer jacket of the pump. The gases in the experimental chamber are simply sucked by the cool surface of the absorbant and retained as long as the surface is cold enough. After achieving the desired vacuum in the main chamber the intake port is isolated and liquid from the outer jacket is removed. The pressure due to gases can be released by simply venting through the degassing port. Often heating can be used to accelerate the degassing and making the sorbent surface fresh. The sorption pump can be reused. Oil-free pumping without any moving parts (no vibrations) makes this pump suitable for many experiments (Fig. A.3).

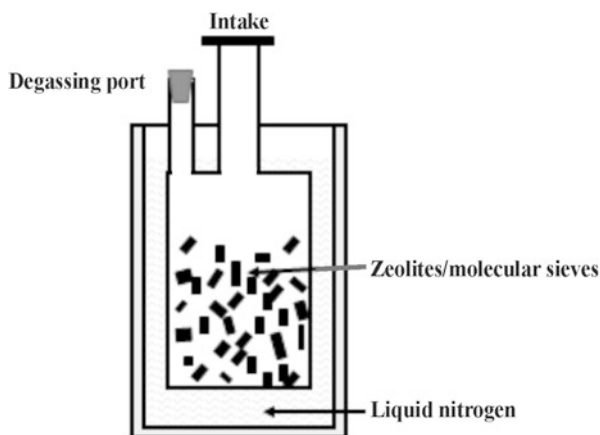


Fig. A.3 Schematic diagram of a sorption pump

Ion Pump

This pump operates in the range of $\sim 10^{-1}$ to 10^{-6} Pa. It also needs pre-pumping by another pump like a diffusion pump. However, unlike diffusion pump which constantly needs a backing pump to operate along with it, the ion pump can be left to itself once it starts. The pumping action simply requires that the gas be ionized. With the cathode-anode assembly (upto 4–5 KV) with parallel magnetic field (3–4 KGauss), initial ionization is triggered by stray electrons or radiation in the chamber. The electrons are accelerated in path of spiral towards anode and ionize on the way the neutral gas molecules inside the chamber. The gas atoms/molecules with positive charge are then accelerated towards the cathode where they can react and get adsorbed as well as sputter cathode material viz. titanium. It gets deposited in other parts like anode where too pumping of gas can occur. Usually diode and triode pump geometries are available (Fig. A.4).

Triode ion pump is just a modified version of the diode pump in order to avoid the argon instability. The insertion of a slotted type cathode assembly enables the argon ion to get embedded without resputtering. Different gases get pumped at different speeds depending upon their reactivity with the cathode material viz. titanium (Fig. A.5).

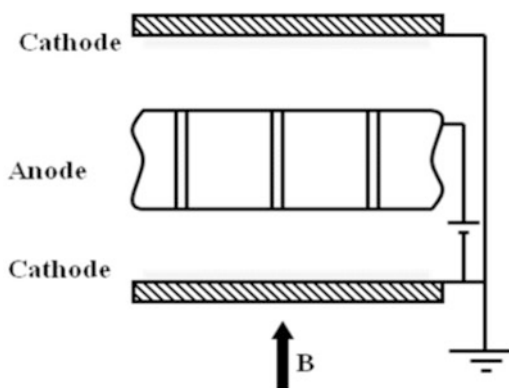


Fig. A.4 Schematic diagram of an ion pump with diode configuration

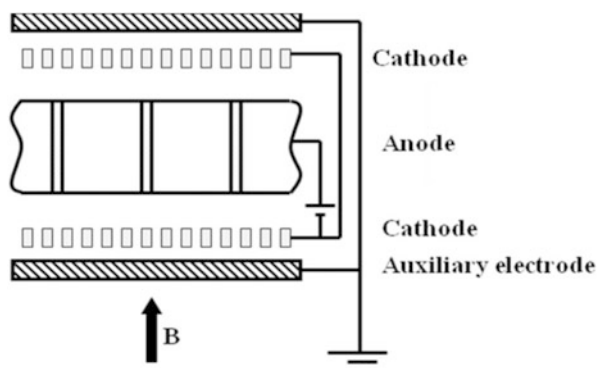


Fig. A.5 Schematic of a triode ion pump

Vacuum Gauges

There are three types of vacuum measurement devices called ‘vacuum gauges’.

1. Mechanical gauges – like U-tube manometer, dial strain gauge, capacitance gauge and McLeod gauge. These gauges measure the actual force exerted by the gas or air.
2. Transport gauges – like Pirani and thermocouple gauges. They work on the principle of conductivity of a gas or air.
3. Ionization Gauges – depend upon the ionization of the gas in vacuum system. Cold cathode ionization gauge, hot filament ionization gauge and Bayard-Alpert gauge are the examples of ionization gauges.

We shall discuss in brief few gauges below.

U-tube Manometer

As shown in Fig. A.6 one end of a glass tube is open to the atmosphere (pressure P_1) and other end is connected to the vacuum system of which pressure is to be measured. The tube is partially filled with mercury or some other liquid with low vapour pressure. Initially, when the vacuum system is not pumped height of the liquid column in both the arms would be equal. When there is a pressure difference due to vacuum system evacuation, the liquid level in the arm connected to the vacuum system would rise and a height difference ‘ h ’ would be observed in the two arms. If P_2 is the pressure in the vacuum system,

$$P_1 - P_2 = h\rho g$$

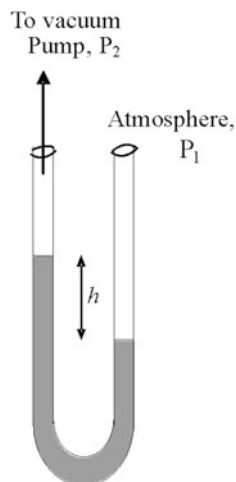


Fig. A.6 U-tube manometer

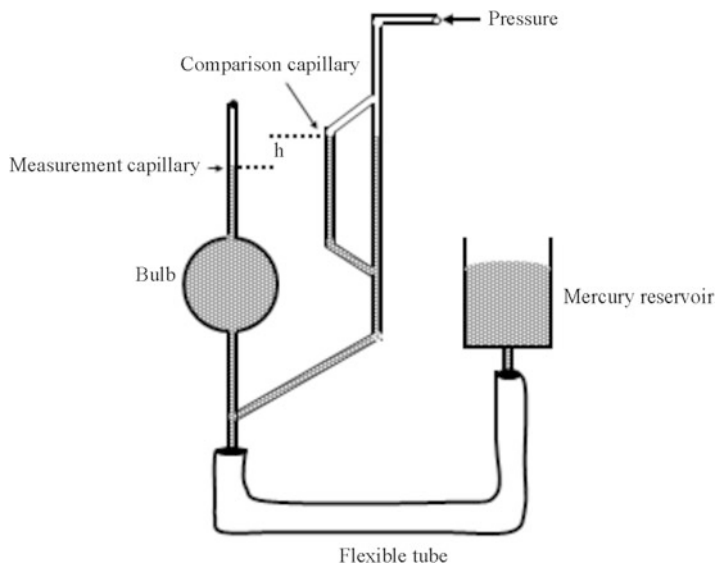


Fig. A.7 A McLeod gauge

where ρ is the density of the liquid and g is the gravitational constant. Knowing all the quantities except pressure (P_2), it can be estimated from above equation.

As illustrated in Fig. A.6 one end can be sealed to know vacuum before pumping mercury, so that changes in the atmospheric pressure need not be known.

McLeod Gauge

The construction of this gauge is as shown in Fig. A.7. Initially the mercury level is brought to point 1. Gas is at pressure P and bulb is also at same pressure. Volume of the bulb is V , hence gas volume in bulb is V . Reservoir of mercury is lifted up so that the mercury reaches top of capillary. Gas in the capillary of diameter d and bulb is compressed. It can be shown that pressure is given as reservoir is lifted up so that the level of mercury is at the level of capillary.

Pirani Gauge

A filament is made the part of Wheatstone bridge. Initially when the system is at atmospheric pressure the resistance bridge is balanced. When the system starts getting evacuated and the heated filament starts losing its heat, its resistance changes disturbing the balance. Consequently the current flowing in the ammeter in the

bridge serves as a pressure reading. The filament resistance is a function of the heat lost by it to the colliding gas molecules.

Thermocouple Gauge

Thermocouple gauge also works on the same principle as Pirani gauge i.e. thermal conductivity of gases. If large number of gas molecules are present in the system more heat will be taken off the heated filament or wire than otherwise. Heat conducted is proportional to the gas density and becomes a measure of vacuum. Only difference in Pirani and thermocouple gauge is that instead of resistivity change as in Pirani gauge, temperature of a heated wire is directly measured by a thermocouple (thermocouple is a temperature measuring device which works on the principle that if two dissimilar metal junctions are made and one is hot and the other is cold then there is a current flow which depends on the temperature difference between the two junctions). The temperature of the heated wire is in turn a measure of vacuum.

Cold Cathode Gauge (Penning Gauge)

It is similar in principle to an ion pump. However, the dimensions of the electrodes are small and it is not designed to pump large volume of gases in the vacuum system but derive only a small amount for the measurement of a gas.

An anode maintained at high voltage (~ 2 kV D.C.) is placed between two flat cathode plates. A magnetic field is applied parallel to the electric field. Gas is ionized under the influence of applied electric and magnetic fields. The ion current is a measure of amount of pressure of gas inside the system as amount of gas being ionized is a function of amount of gas molecules present inside (Fig. A.8).

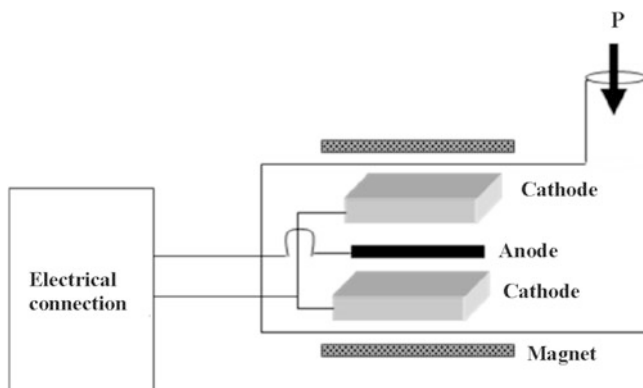


Fig. A.8 A penning gauge

Hot Cathode Ion Gauge

In this gauge a hot filament is surrounded by a metal grid and both are placed inside a cylindrical metal collector. The grid is held at ~ 200 V positive potential with respect to the filament. Electrons from the filament are accelerated towards the grid and can shoot up slightly beyond the grid but are repelled by the negatively held collector and again attracted by the grid. They may oscillate to some extent around the grid and cause ionization of the gas inside the gauge. The ionized gas molecules or atoms when positively charged get accelerated towards the collector causing a current which is proportional to the density of gas molecules or pressure inside the vacuum system (Fig. A.9).

Major problem with this kind of ion gauge is that at low pressure, when ions are accelerated towards the collector, without any collisions with other molecules, they are able to produce X-rays or UV radiation while they strike the collector. This radiation in turn can produce additional ionization which results into additional ion current. Thus even with lower vacuum one observes larger ion current or higher pressure than actually exists in the system. Thus the gauge becomes inaccurate in the ultra high vacuum regime.

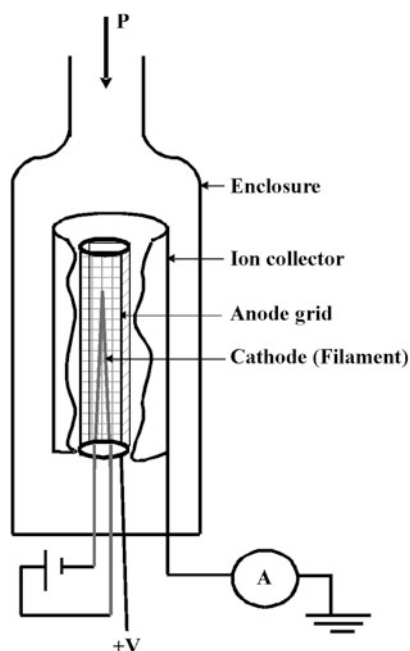
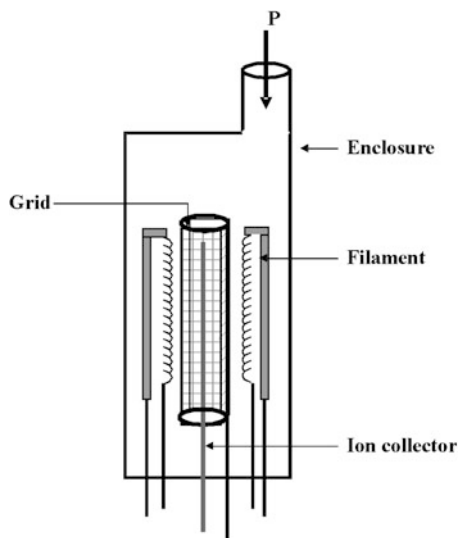


Fig. A.9 Schematic of an ion gauge

Fig. A.10 Schematic of a B-A gauge



Bayerd-Alpert (B-A) Gauge

An improvement in the design of ionization gauge is made in the Bayerd-Alpert gauge so as to increase the accuracy of the gauge in the ultra high vacuum range. In order to overcome the limitation due to the so-called 'X-ray limit' of an ionization gauge discussed above, the large area collector is replaced by a thin collector, a metal wire, placed at the centre of the gauge head. The collector is surrounded by the metal grid and the filament is placed outside the grid. This enables the pressure measurement upto $\sim 10^{-8}$ Pa (Fig. A.10).

Below $\sim 10^{-8}$ Pa, pressure/vacuum measurement even with a B-A gauge becomes unreliable. Often partial pressure gauges like mass spectrometers are employed to understand the composition of residual gases inside the vacuum systems at such a high vacuum.

Further Reading

- A. Chambers, R.K. Fitch, B.S. Halliday, *Basic Vacuum Technology* (IOP Publishing Ltd., Bristol, 1989).
- L.I. Maissel, R. Glang, *Handbook of Thin Film Technology* (McGraw Hill Company, New York, 1970)
- A. Roth, *Vacuum Technology*, 2nd edn (North-Holland, Amsterdam, 1982)

Appendix V

Properties of Some Semiconductors

Semi-conductor	Structure(s)	Lattice constant(s)/Å	Bulk band-gap at 300 KeV	^a Electron effective mass/m ₀	^a Hole effective mass/m ₀	Static dielectric constant
Si	Cubic, Diamond	5.431	1.12 (i)	0.98; 0.19	0.16;0.49	11.4
GaAs	Zinc blende	5.653	1.42 (d)	0.067	0.082	13.1
GaSb	Zinc blende	6.098	0.72 (d)	0.042	0.4	15.7
Ge	Cubic, Diamond	5.646	0.66 (i)	1.64; 0.082	0.04;0.28	16
ZnO	Wurtzite	$a = 3.25, c = 5.2; a = 4.58$	3.35 (d)	0.29	1.21	9
	Rock salt					
ZnS	Zinc blende	5.42	3.68 (d)	0.4	0.61	5.2
ZnSe	Wurtzite; Zinc blende	$a = 3.82, c = 6.626; a = 5.668$	2.7 (d)	0.17	1.44	9.12
InP	Zinc blende	5.869	1.35 (d)	0.077	0.64	12.4
InAs	Zinc blende	6.058	0.36 (d)	0.023	0.4	14.6
InSb	Zinc blende	6.479	0.17 (d)	0.0145	0.4	17.7
CdS	Wurtzite	$a = 4.16, c = 6.756$	2.42 (d)	0.21	0.8	5.4
CdSe	Zinc blende	6.05	1.7 (d)	0.13	0.45	10
CdTe	Zinc blende	6.482	1.56 (d)	0.11	0.4	10.2
PbS	Rock salt	5.9362	0.41 (d)	0.25	0.25	17
PbSe	Cubic	6.117	0.27 (d)	0.047	0.041	280
PbTe	Rock salt	6.462	0.31 (d)	0.17	0.2	30

^a 'm₀' stands for the standard electron mass = 9.1×10^{-31} kg. 'i' and 'd' are for indirect and direct band gaps respectively

Appendix VI

Kronig Penney Model (1-D)

Free electron theory cannot explain the various observed properties of solids. The band theory could explain the various properties like conductivity, Hall effect etc. correctly. The origin of bands in a solid can be understood by considering the motion of an electron in a one dimensional periodic lattice. The model was first proposed by Kronig and Penney way back in 1930 and bears their name. Kronig and Penney assumed that the electron experiences, as illustrated in Fig. A.11 a periodic potential with a period $(a + b)$.

Potential energy is zero between 0 and a and maximum (V_0) at 0 and a . Atomic nuclei are separated by a period of $(a + b)$. Thus

$$0 < x < a, \quad V = 0 \quad (\text{A.1})$$

$$-b < x < 0, \quad V = V_0 \quad (\text{A.2})$$

Schrödinger equations for the two regions are

$$\frac{d^2\psi}{dx^2} + \frac{2m}{\hbar^2} E \psi = 0 \quad (0 < x < a) \quad (\text{A.3})$$

and

$$\frac{d^2\psi}{dx^2} + \frac{2m}{\hbar^2} (E - V_0) \psi = 0 \quad (-b < x < 0) \quad (\text{A.4})$$

We assume that, for the electron, $E < V_0$.

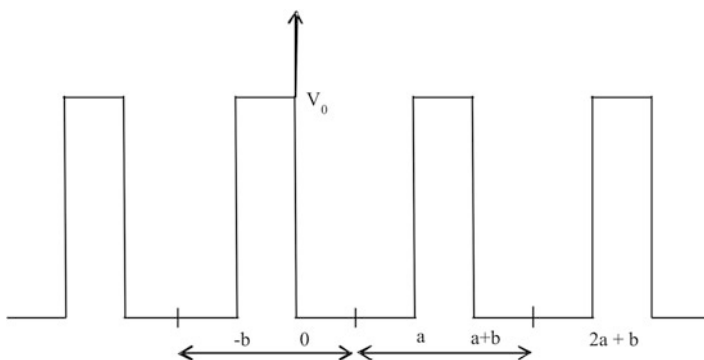


Fig. A.11 Periodic potential in a 1-D lattice

Put

$$\frac{2mE}{\hbar^2} = \alpha^2 \quad (\text{A.5})$$

and

$$\frac{2m}{\hbar^2} (V_0 - E) = \beta^2 \quad (\text{as } E < V_0) \quad (\text{A.6})$$

The Bloch theorem for the motion of an electron in the periodic potential states that if $u(x) = u(x + a)$ where 'a' is the lattice spacing then

$$\psi(x) = e^{\pm i k x} u(x) \quad (\text{A.7})$$

is the plane wave solution of the Schrödinger equation modulated by the periodicity of the lattice.

Using (A.7) and substituting in Eqs. A.3 and A.4 we get

$$\frac{d^2 u}{dx^2} + 2ik \frac{du}{dx} - (a^2 + k^2) u = 0 \quad (\text{for } 0 < x < a) \quad (\text{A.8})$$

and

$$\frac{d^2 u}{dx^2} + 2ik \frac{du}{dx} - (\beta^2 + k^2) u = 0 \quad (\text{for } -b < x < 0) \quad (\text{A.9})$$

Consider the following solutions for these equations as

$$u_1 = Ae^{i(a-k)x} + Be^{-i(a+k)x} \quad (\text{for } 0 < x < a) \quad (\text{A.10})$$

and

$$u_2 = Ce^{(\beta-ik)x} + De^{-(\beta+ik)x} \quad (\text{for } -b < x < 0) \quad (\text{A.11})$$

where, A, B, C and D are constants. These constants have to be chosen such that they satisfy the boundary conditions

For continuity:

$$u_1(0) = u_2(0) \text{ and } \left(\frac{du_1}{dx} \right)_{x=0} = \left(\frac{du_2}{dx} \right)_{x=0} \quad (\text{A.12})$$

For periodicity:

$$u_1(a) = u_2(-b) \text{ and } \left(\frac{du_1}{dx} \right)_{x=a} = \left(\frac{du_2}{dx} \right)_{x=-b} \quad (\text{A.13})$$

These lead to four linear equations involving constants A, B, C and D .

The four equations have a solution if the determinant of coefficients of A, B, C , and D vanishes.

This leads to the equation

$$\frac{\beta^2 - \alpha^2}{2\alpha\beta} \sinh(\beta b) \sin(\alpha a) + \cosh(\beta b) \cos(\alpha a) = \cos k(a + b) \quad (\text{A.14})$$

If $V_0 \rightarrow \infty$ and $b \rightarrow 0$ but $V_0 b$ is finite then

$$\frac{mV_0 b a}{\hbar^2 \alpha} \sin(\alpha a) + \cos(\alpha a) = \cos(ka) \quad (\text{A.15})$$

By writing

$$p = \frac{mV_0 b a}{\hbar^2} \quad (\text{A.16})$$

We get simplified form of Eq. A.14 as

$$p \frac{\sin(\alpha a)}{\alpha a} + \cos(\alpha a) = \cos(ka) \quad (\text{A.17})$$

R.H.S can take values between $+1$ and -1 . This implies only allowed values of αa for which L.H.S. lies between ± 1 . Thus one can see that

1. Electron energy spectrum has values separated by forbidden regions.
2. With increasing αa band width increases.
3. Width of a band decreases as P increases (Fig. A.12).

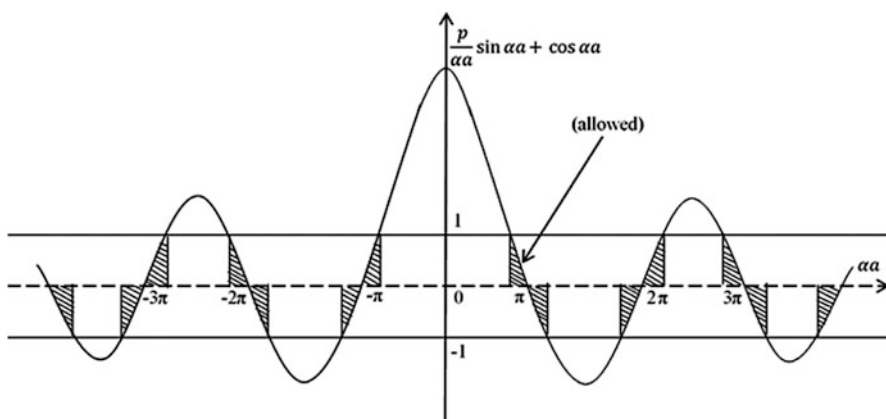


Fig. A.12 Graphical representation of Eq. A.17

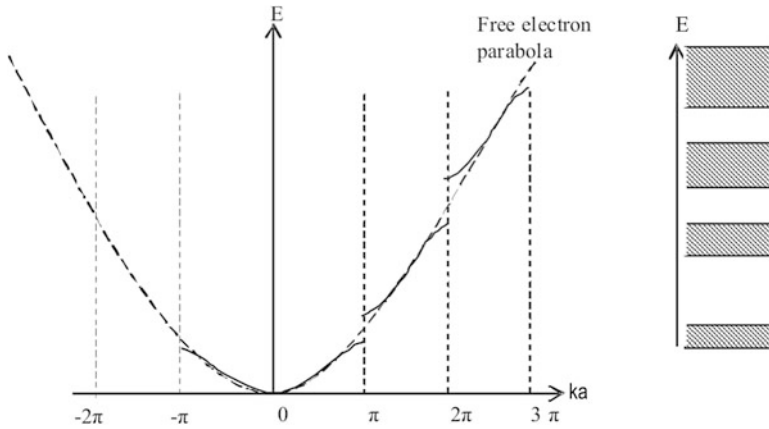


Fig. A.13 On the left hand side a plot of E against ka for an electron in a periodic lattice, as predicted by Kronig Penney model is shown. The *dotted* free electron parabola shows the effect of neglecting the periodic potential in a solid and simply assuming the electron gas. Illustration on the right hand side is schematics showing how electron energy band separation decreases as we go on increasing the energy

When $P \rightarrow \infty$ allowed region narrows down to a line spectrum.

If, $P \rightarrow \infty$,

$$\sin \alpha a = 0 \text{ i.e. } \alpha a = \pm n\pi, \text{ where } n = 1, 2, 3, \dots$$

$$\alpha^2 = \frac{2mE}{\hbar^2} = (n\pi)^2 \quad (\text{A.18})$$

Therefore

$$E_n = \frac{n^2 \hbar^2 \pi^2}{2m a^2}$$

When E is plotted against ka , we get a plot as illustrated in Fig. A.13.

Further Reading

A.J. Dekker, *Solid State Physics* (Macmillan, London, 1952)

C. Kittel, *Introduction to Solid State Physics*, 5th edn (Wiley Eastern Ltd., New Delhi, 1994)

Index

A

AAO, 376–378
Aerogel, 78, 103, 104, 273, 296–303, 319, 334, 347
Agriculture, 317, 345–346, 350, 351
Amorphous, 32, 33, 56–58, 63, 106, 112, 115, 165, 167, 169, 184, 282, 283, 288, 294, 306, 320
Amphiphilic, 95, 97, 99, 306
Apoferritin, 121, 122
Atomic scattering factor, 161–162, 165, 166
Autoclave, 105, 207, 300, 304, 308
Automobile, 302, 318, 328, 336–337, 350

B

Ball milling, 55–57, 308
Biocomposites, 111
Biological labeling, 340
Biomaterials, 113, 114
Biominerals, 111, 112
Black body, 1–3, 5, 6, 11
Bonding, 31–53, 64, 123, 126–128, 132, 136, 199–201, 203, 235, 244, 277, 285, 287, 301, 303, 310
 covalent, 45–49, 126, 201, 235, 310
 ionic, 45–49
 metallic, 45, 47, 48
 secondary, 45, 48–49
Bottom up, 88, 241, 313
Bragg condition (law), 162–165, 167
Bravais lattice, 36, 37, 40, 159
Brownian motion, 31, 79, 80, 84, 85, 171

C

Cancer, 113, 214, 339, 340, 342–344, 346, 353
Cantilever, 148, 153–155
Carbon nanotubes (CNT), 73, 74, 184, 259, 267, 273–285, 301, 318, 329, 330, 332, 335, 337, 338, 345, 351–353
 multiwall, 281
 single wall, 65
Cloaking, 214, 311, 313, 348
Clusters, 50, 55, 59, 61–66, 90, 199–203, 207, 238, 261–263, 273–276
Colloids, 55, 78–87, 93, 103, 122, 131, 135, 176, 248, 308, 349
Compton effect, 2, 9, 14
Core-shell particles, 273, 306, 308–311, 340, 341, 343, 344
Cosmetics, 339, 353, 354
Coulomb
 blockade, 260–263
 potential, 25–27
Counter ions, 84
Crystal structure factor, 166–167
Curie law, 229
Curie–Weiss law, 229

D

de-Broglie wave, 11, 19
Defense, 347–348
Density of states, 22–25, 188
Deposition
 chemical vapour (CVD), 71–73, 103, 188, 275, 281–283, 285
 ECR, 71
 electric arc, 73–74, 283

Deposition (*cont.*)

- ion, 66
- ionized cluster beam, 63–64
- laser ablation, 275, 281, 283
- laser pyrolysis, 65
- magnetron sputtering, 67, 69–70
- molecular beam epitaxy (MBE), 75
- physical vapour (PVD), 62, 248
- plasma, 70–71
- RF, 67
- sputter, 61, 65–71, 75
- Diamond, 33, 39, 45, 47, 148, 183, 196, 197, 236, 249, 273, 285, 302, 394
- Diatom, 113, 115, 116
- Dielectric confinement, 215, 220
- Diffraction
 - electron, 13, 136, 239
 - neutron, 136, 159, 169
 - X-ray, 42, 136, 159, 160, 162, 169, 237–239, 355, 356, 360, 362, 364
- DLS. *See* Dynamic light scattering (DLS)
- DLVO, 83
- DNA, 113–115, 121–123, 126, 128, 130–132, 353
- Drug delivery, 98, 113, 276, 286, 307, 308, 339, 341–343
- Dye, 299, 309, 311, 319, 321–326, 338–340, 344, 347, 350, 351, 353
- Dynamic light scattering (DLS), 136, 171–173

E

- Effective mass approximation, 205–208, 295, 362
- Elay–Riedel, 73
- Electrochemical, 92, 287–289, 292, 293, 333–334, 366, 376
- Electron confinement, 19–27, 214, 215
- Emulsion, 98, 309
- Energy, 2, 44, 55, 81, 114, 125, 136, 202, 247, 261, 274, 317, 350, 362
- Energy gap, 52, 136, 173–175, 204, 205, 207–209, 212, 263, 295, 323, 362
- Environment, 115, 117, 147, 148, 188, 294, 328, 334, 338, 347, 349–354
- Enzymes, 113, 114, 119, 342
- Eukaryotes, 113
- Exciton
 - biexciton, 204
 - bright, 204
 - dark, 204
 - Frenkel, 205
 - Mott–Wannier, 204–206
- Exclusion principle, 45, 52, 153, 203, 227

F

- Far field, 156
- Fermi level, 52, 149, 187, 188, 263, 264
- Ferritin, 121, 122, 131
- Field emission, 212, 213
- Fluorescence, 185, 210, 340, 341, 352
- Food, 107, 116, 128, 345–346, 349, 353, 383
- Fuel cells, 317, 321, 327–335, 337, 350, 352
- Fullerene, 73, 74, 201, 273–278, 281–283, 285, 326, 335, 345, 352–354
- Fundamental constants, 6, 383

G

- Gecko effect, 315
- Grain, 32, 56, 79, 140, 165, 235, 237
- Graphene, 259, 273, 285, 301, 321, 329, 335, 336, 351, 352
- Graphite, 33, 39, 74, 273, 275–278, 281, 283, 289, 302, 366
- Grätzel cell, 321–325
- Green synthesis, 111

H

- Hamaker constant, 84
- Horse spleen, 122
- Hund's rules, 227
- Hydrogen storage, 318, 329, 334–335, 337
- Hydrophilic, 95, 96, 99, 299, 313, 314
- Hydrophobic, 95, 96, 99, 299, 313, 314, 343, 351
- Hydrothermal synthesis, 105

I

- Imaging, 142, 155, 157, 249, 339–341, 343, 353
- Inverse micelles, 99, 100, 102
- Ionization energy, 47, 202, 261

L

- Lab-on-chip, 107–109, 244
- LaMer diagram, 88, 89
- Langmuir–Hinschelwood, 72, 73
- Lattice, 40, 131, 132, 146, 167, 204, 205, 217, 220, 234, 237–239, 272, 277, 394–396, 398
 - body centered (bcc), 37
- Bravais, 36, 37, 40, 159

- hexagonal, 130, 169
 - oblique, 130
 - primitive, 33, 34, 37, 38
 - reciprocal, 38–41
 - square, 130
 - Laue method, 161, 162
 - L-B films, 95–99
 - Light confinement, 272
 - Liquid crystals, 43–44, 126, 308
 - Lithography, 107, 126, 127, 148, 241, 243–256, 259, 260, 372–376
 - dip-pen, 126, 127, 249–251
 - electrical SPM, 252
 - electron beam, 243, 247–248, 253
 - ion beam, 248
 - laser, 246
 - nanosphere, 248–249
 - neutral beam, 248
 - optical SPM, 251
 - photon, 252
 - scanning probe, 252
 - soft, 252–256
 - thermo-mechanical, 251–252
 - UV beam, 246
 - X-ray, 243, 246
 - Localized surface plasmon, 215–222
 - Lotus effect, 313–314
 - Luminescence, 136, 184–186, 209–213, 287, 288, 296, 311, 368
 - cathode, 203, 213
 - electro, 126, 185, 209, 211–213
 - high field, 212, 213
 - injection, 211–213
 - photo, 185–186, 209–211, 288, 293, 295–296, 311, 356, 366
 - thermo, 203, 213–214
- M**
- Macromolecules, 114
 - Magnetic domains, 130, 232
 - Magnetic susceptibility, 228, 229
 - Magnetic Tunnel Junction (MTJ), 242, 267, 271
 - Magnetism
 - antiferromagnetism, 231–232
 - diamagnetism, 228
 - ferrimagnetism, 227, 231
 - ferromagnetism, 231
 - nanomagnetism, 225
 - paramagnetism, 227
 - superparamagnetism, 233, 344
 - Magnetoresistance
 - collosal (CMR), 234
 - giant (GMR), 234, 268
 - Melt mixing, 57
 - Membranes, 113, 121, 329
 - S-layers, 121, 130
 - Metal organic framework (MOF), 307–308, 334
 - Metamaterials, 311–313, 348
 - Micelles, 99–102, 126, 306–307, 341, 346
 - Microemulsion, 98–102, 309
 - Microorganism, 111, 113, 116–120
 - Microreactor, 108, 109
 - Microscopes
 - atomic force (AFM), 136, 148, 152
 - confocal, 136, 140
 - field emission (FESEM), 143
 - optical, 135, 136
 - scanning electron (SEM), 136, 142, 143, 355
 - scanning near-field optical (SNOM), 148, 155
 - scanning tunnelling (STM), 142, 148, 244
 - transmission electron (TEM), 136, 142, 146, 356
 - Microwave, 70, 107, 158, 159, 313
 - Mie theory, 215–217
 - Moor's law, 242, 260
 - Multilayers, 59, 96, 135, 169, 234, 241, 268–270
- N**
- Nanocomposites, 334, 335, 346
 - Nanoelectronics, 126, 259–272
 - Nanoindentor, 197
 - Nanolithography, 241–256
 - Nanomagnetism, 130, 225–235
 - Nanophotonics, 222, 224, 272
 - Near-field, 135, 136, 155–159, 251
 - Néel temperature, 232
 - Negative refractive index, 311, 312
 - Neutron diffraction, 136, 159, 169
- O**
- Ostwald ripening, 89
- P**
- Penrose tiling, 41, 42
 - Phosphorescence, 185, 210

Photoelectric effect, 2, 6, 7, 9, 10, 186
 Photon tunnelling, 157–158
 Photoresist, 244, 245, 247, 251, 370, 371
 Plasma, 14, 67, 68, 70–71, 218, 243, 306
 Plasmonic, 188, 214, 313, 340
 Pollution, 98, 103, 109, 320, 328, 334, 337, 338, 349–352, 354
 Polycrystalline, 32, 33, 165, 167, 170, 235–237
 Poly dimethyl sulphoxide (PDMS), 107, 253–256, 372–376
 Polymer, 86, 128, 132, 147, 244, 248, 253–255, 273, 310, 326–332, 334, 336, 338, 340, 342, 343, 347, 348
 Porous silicon, 185, 286–289, 291–296, 303, 366–369
 Primitive cell, 33, 34, 37
 Printing
 micro contact (μ CP), 253, 372
 micromoulding in capillaries (MIMIC), 372
 micro transfer molding (μ TM), 255, 372
 replica molding (REM), 255, 372
 soft lithography, 252, 372
 solvent assisted micromolding (SAMIM), 255, 372
 Properties of nanomaterials
 magnetic, 225
 mechanical, 235
 melting, 238
 optical, 208
 structural, 237
 Protein, 114, 117, 119, 121–123, 130, 131, 272, 342, 353

Q

Quantum dot, 20, 75, 132, 135, 242, 262, 263, 272, 324–326, 340, 350, 352
 Quantum well, 20, 29, 75, 135
 Quantum wire, 75, 135

R

Raman scattering, 181

S

Scherrer formula, 169, 362, 364
 Schottky barrier, 212
 Schrödinger equation, 2, 15–20, 25–27, 50, 207, 395, 396
 Self assembly
 co-assembly, 128
 dynamic, 127

 hierarchical, 128
 static, 127
 Self cleaning, 313–314, 329, 337, 338
 Single crystal, 32, 33, 132, 161, 165, 167, 181, 235, 277, 286
 Single electron transistor (SET), 242, 263–267
 S-layer, 121–122, 128, 130, 131
 Solar cell
 dye-sensitized, 321, 353
 hybrid, 321, 335
 organic, 320, 326
 Sol-gel, 103–104, 132, 296, 304, 309
 Sonochemical, 106–107
 Spectroscopies
 electron spectroscopy (XPS, UPS, ESCA, Auger), 136, 186
 Fourier transform infrared (FTIR), 136, 177, 179, 180
 infra-red (IR), 176–181, 287, 320
 luminescence, 136
 Raman, 136, 181–184, 352
 UV-Vis-NIR, 175–176
 Spin field effect transistor (Spin FET), 271
 Spintronics, 65, 267–271, 285
 Spin valve, 242, 267, 270–271
 Sports, 338
 Supramolecules, 126
 Surface plasmon polariton, 215, 222–225
 Surface plasmon resonance (SPR), 175, 215–222, 311, 343, 358
 Surface tension, 81, 82, 89, 99, 101, 296–298
 Surfactant, 95, 97, 99–102, 310

T

Textile, 338
 Tissue, 111, 339, 342, 344–345, 353
 Top down, 241
 Toys, 338
 Tunnelling, 27–29, 135, 136, 148–152, 157, 158, 249, 262, 263, 271, 290

U

Uncertainty principle, 2, 12–15
 Unit cell, 33, 34, 38, 41, 166, 167, 199, 201, 303

V

Vacuum
 gauges, 384, 389
 pumps, 298, 384, 385
 units, 351, 383

Van der Waals, 48, 49, 81, 97, 133, 315
Vibrating sample magnetometry (VSM), 136,
192–194

W

Winsor phase diagram, 102
Work function, 8, 149, 187, 188, 203, 212

X

X-ray diffraction, 42, 136, 159, 160, 162, 169,
197, 237–239, 355, 356, 362, 364

Z

Zeolites, 103, 104, 303–308, 352