

CHAPTER 2

IMAGE FORMATION

"Plenty and variety of images exist in life"

There are plane images

..... grey images

..... colored images

..... three dimensional images

"However, humans are always images of each other"

A . OPTICAL IMAGES

A.1. Image formation by the eye

In principle, the eye is similar to a camera. The eye-lens forms a real image on the retina which is one part of the four essential parts of the human eye shown in Fig. 2.1.

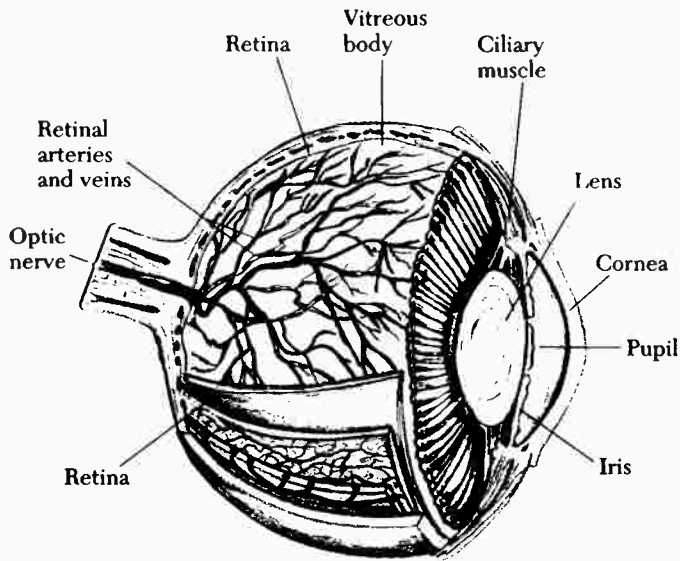


Fig. 2.1. The human eye.

The retina is a delicate membrane packed with light sensitive cells which send nerve impulses to the brain. The space between the cornea and the crystalline lens is filled with a transparent material, aqueous humor. The main body of the eye is filled with vitreous humor, a jelly-like transparent material having nearly the same refractive index as the aqueous humor, 1.34. The cornea and the aqueous humor act as a lens providing most of the light rays entering the eye. The crystalline lens merely provides the fine adjustment of the focal length required to bring the image of the object exactly on

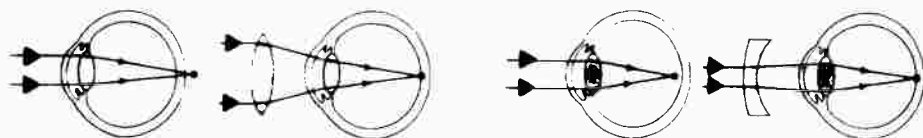
the retina, regardless of the object distance. The power of the lens could be adjusted by means of ciliary muscles. The shortest distance at which an object can be placed from the eye and be seen sharply is called the near point (about 25 cm) for a normal young person. This distance changes by age due to the lack of accommodation of the ciliary muscles. For instance, at an age about 60 years, the near point becomes at about two meters.

When the image of an object is formed behind the retina, as in Fig. 2.2, the object is seen blurred. This condition can be corrected with a converging lens that helps in bringing the image exactly on the retina, and is known as "hyperopia".

On the other hand, if the eye is longer than normal or due to some defect in the eye lens, the image of a distant object is thus not seen clearly and this known as "myopia". A diverging lens is needed to correct the eye defect.

A person may also have an eye defect known as "astigmatism", in which a point source produces a line image on the retina. This defect arises when the eye-ball is not perfectly spherical. Astigmatism can be corrected with lenses having different curvature called cylindrical lenses. The different curvatures of the cylindrical lens compensates for the curvature of the cornea in the particular astigmatic plane.

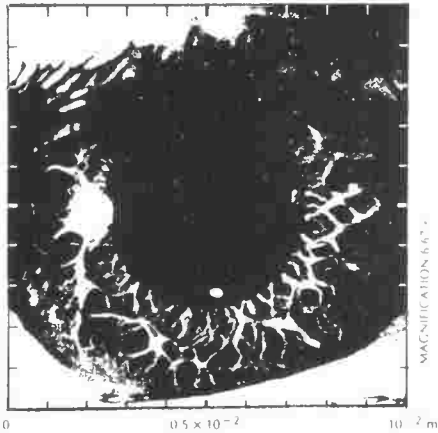
Image formation by the eye could also be affected by other diseases such as "cataract" where the lens becomes partially or totally opaque. This could be remedied by surgical removal of the lens. Another disease is called "glaucoma" arising from an abnormal increase in fluid pressure inside the eye-ball. This disease might cause total blindness. It could be treated by medicine, surgery, and most recently by the help of lasers.



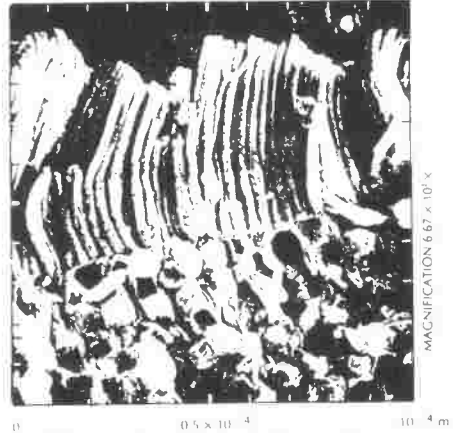
A hyperopic eye could be corrected by a converging lens

A myopic eye could be corrected by a diverging lens

Fig. 2.2. Eye glasses to correct defects of sight.



The human eye



The rod cells of the retina

A.2. Microscope images

Optical images are formed when spherical waves fall on plane and spherical surfaces. Mirrors and lenses are devices that work on the basis of image formation by reflection and refraction. These devices are used in optical instruments that help us in magnifying very small objects or seeing very far objects. Because physics encompasses the smallest particles, such as electrons and quarks, and it also encompasses the largest bodies, such as galaxies, thus looking at these bodies by some means is of vital importance to our knowledge of the physical world around us. The smallest particles and the largest bodies differ in size by a factor of more than 10^{40} . Therefore, to have a look at each of these bodies one would choose the suitable means for that.

The optical microscope and the telescope were the first instruments to explore the unknowns in our physical world. These instruments employ usually a system of two lenses: the objective and the eyepiece.

A.3. The microscope

The microscope consists of two lenses of very short focal lengths. The objective is placed near the object, and it forms a real, magnified

image of the object. This image serves as object for the eyepiece, which acts as a magnifier and forms a virtual image at infinity (see Fig. 2.3). Thus, both the objective and the eyepiece contribute to the magnification of the microscope. Good microscopes operate at magnifications of up to 1400. Higher magnifications could be obtained using the oil-immersion microscope.

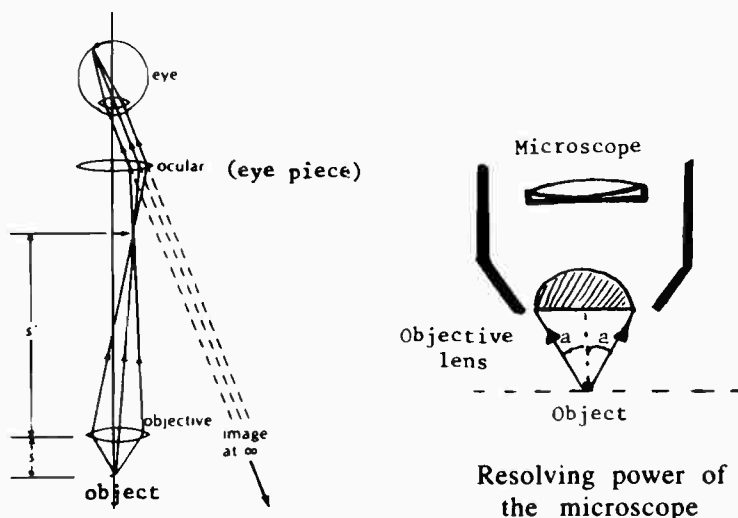


Fig. 2.3. Arrangement of lenses in a microscope. The object to be magnified is below the objective and the eye is just above the eyepiece. The eyepiece forms a virtual image at infinity, and the lens of the eye focuses on the retina the parallel rays emerging from the eyepiece.

The resolving power of the microscope, d , is the limit of resolution for the object and is given by:

$$d = \lambda / 2 \sin a$$

where λ is the wavelength of the light used, and a is the half angle subtended by the objective at the object O.

The eye can resolve about 0.01 cm. The largest useful magnifying power of a microscope is one which magnifies the limit of resolution of the objective to that of the eye. This is about 1000 with glass lenses and visible light. Higher resolving power may be obtained using light of shorter wavelengths such as ultraviolet. An electron microscope (see later) contains electron lenses and utilizes electrons in place of light, it has a limit of resolution less than 10^{-9} m owing to the much shorter wavelength of moving electrons compared with that of light.

A.4. The telescope

The simple astronomical telescope consists of an objective of very long focal length and an eyepiece of short focal length. These two lenses are separated by a distance nearly equal to the sum of their individual focal lengths, so that their focal points coincide. The objective forms a real image of the distant object. This image serves as object for the eyepiece, which forms a magnified virtual image at infinity (see Fig. 2.4).

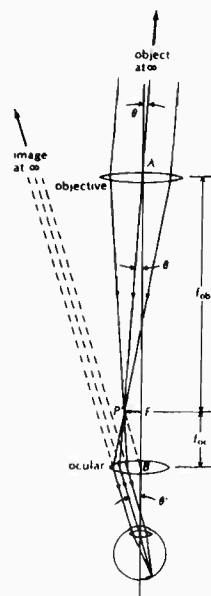


Fig. 2.4. An astronomical telescope. The object is at a large distance above. The observer's eye is below the eye-piece which forms a virtual image at infinity, and the lens of the eye focuses on the retina the parallel rays emerging from the eyepiece.

Many astronomical telescopes are reflecting telescopes in which a concave mirror plays the role of the objective. The mirror forms a real image that serves as object for the eyepiece. The eyepiece is placed in front of the mirror and blocks out some of the light. Because large mirrors of good quality, free of aberrations, are easier to manufacture than large lenses of good quality, the largest astronomical telescopes all use mirrors (about 5 meters in diameter and a focal length of about 17 meters).

B. IMAGE FORMATION BY ELECTRONIC DEVICES

B.1. Introduction

The understanding of how images formed by electronic devices requires the study of that branch of physics which is concerned with electric and magnetic phenomena. The laws of electricity and magnetism play a central role in the operation of such devices as televisions, electron microscopes, computers, electronic devices used in medicine, such as sonography, magnetic resonance imaging, etc.

It is well known that interatomic and intermolecular forces are responsible for the formation of matter in its three forms, solid, liquid and gas. These forces are electric in origin. Elastic forces arise from electric forces at the atomic level.

In the following non-traditional way of teaching physics, namely, under a story cover, we are going to give in sequential form, the physical principles underlying each part of the unit we wish to show how it works. This new method will be stimulating for the student to learn physics since it deals directly with devices and new technologies that every human presently meets in his every day life.

As an example of the new method here adopted for learning physics, let us go through how a television picture is formed on the TV screen. We have always seen TV pictures, and in colors, and we were always eager to know how a movie picture is transmitted through air to be seen every where.

In order to give information about how this happens, one should know some background and fundamental knowledge about the following topics.

1. The electrons coming out of the filament of the TV, and where they come from, how they are ejected. This requires knowledge of the electron structure of the filament, the electron gas theory, thermionic emission from hot bodies, etc.

2. The acceleration of electrons in an electric field, deflection of electrons in electric and magnetic fields, electronic circuits that make the electron beam to scan the TV screen to form the image.
3. How electrons are made to appear on the screen, so we should know something about fluorescence and phosphorescence, color, complementary colors, etc.

This is just an example of how the physics will be presented in a story-like way, and not in the traditional way we used to teach physics, namely, properties of matter, optics, electricity and magnetism, etc.

B.2. Electron microscopy and TV-images

How electron microscope and TV images are formed

In order to know how TV images are formed we have to start with the laws of electricity and magnetism which play a central role in the operation of various devices not only television but also radios, electric motors, high energy accelerators and most of the electronic devices used in medicine. It is well known that the formation of solids and liquids are electric in origin. Electromagnetic interactions hold electrons and nuclei to form atoms, atoms together to form molecules and molecules together to form macroscopic objects. Atoms are hollow structures. The nucleus is a point mass formed of protons and neutrons. Its diameter is of the order of 10^{-12} cm, while the atomic diameter is of the order of 10^{-8} cm.

Magnetism was discovered by the Chinese as early as 2000 B.C. while electrification was discovered by ancient Greeks at about 700 B.C. The word electron came from the Greek word for amber, which was the first material showing the phenomenon of electrification. Two kinds of electrification were discovered, namely, positive and negative electricity.

Towards the end of the nineteenth century Sir J.J. Thomson discovered the existence of the electron, which was the lightest particle known at that time. Its mass was about $1/1840$ of the mass

of the hydrogen atom. Experiments showed that the electron carries a tiny quantity of negative electricity, namely, 1.6×10^{-19} Coulomb.

In 1909, Millikan discovered that electric charge occurs always as some integral multiple of unit charge, e . In modern terms, any charge, q , is said to be quantized. That is, electric charge exists as discrete packets, so we can write, $q = N e$, where N is some integer. Other experiments showed that the proton has an equal and opposite charge $+e$. The neutral atoms, thus, must contain as many protons as electrons.

B.3. Conductors and insulators

Conductors are materials in which electric charges move quite freely whereas insulators are materials that do not readily transport charge. There is a third class of material, called semiconductors which have electrical properties somewhere between those of conductors and insulators. Copper, silver and aluminium are good conductors. Rubber and glass are insulators. Silicon and germanium are semiconductors. The electrical conductivity changes by several orders of magnitude from one class of material to the other. It is of the order of $10^5 \text{ Ohm}^{-1}\text{cm}^{-1}$ for good conductors while it is about $10^{-10} \text{ Ohm}^{-1}\text{cm}^{-1}$ for insulators.

Earthing an apparatus means that we connected it to the ground. The Earth can be considered an infinite sink to which electrons can easily migrate.

B.4. Particle nature of electricity

The chemical effect of the electric current was first studied by Faraday who found that when a current is passed through copper sulphate solution with copper electrodes, copper is deposited on the cathode and is lost from the anode. He, therefore, put forward his first law of electrolysis:

"The mass of any substance liberated in electrolysis is proportional to the quantity of electric charge that liberated it".

Faraday's second law of electrolysis concerns the masses of different substances liberated by the same quantity of charge. His second law is"

"The masses of different substances, liberated in electrolysis by the same quantity of electric charge, are proportional to the ratio of the relative atomic mass to the valency".

This law implies that the same quantity of charge is required to liberate one mole divided by the valency of any substance. This quantity known as the Faraday and is equivalent to 96500 Coulombs.

Faraday's laws indicate that the charge carried by each ion is proportional to its valency. The charge on a monovalent ion could thus be found using the following argument:

Avogadro's constant, about 6.02×10^{23} , is the number of molecules in one mole. In electrolysis, 96500 Coulombs is required to deposit one mole of a monovalent element. When the element is monoatomic, the number of ions of one kind which carry this charge is equal to the number of molecules. Thus the charge on each ion is given by $96500/6.02 \times 10^{23}$ or 1.6×10^{-19} Coulombs (C). If 1.6×10^{-19} C is denoted by the symbol e , the charge on any ion is then e , $2e$, $3e$, etc., depending on its valency. Thus:

" e is basic unit of charge"

All charges, produced by any means, are multiples of the basic unit, e . This was evidenced by the famous Millikan oil-drop method for the determination of the charge e on the electron.

B.5. Millikan's oil-drop experiment (Charge quantization)

Millikan measured the terminal velocity of an oil-drop falling through air. He then charged the oil-drop by using some ionizing radiation, e.g. X-rays. He then applied an electric field to oppose the gravity. The drop now moved with a different terminal velocity, which was again measured (see Fig. 2.5).

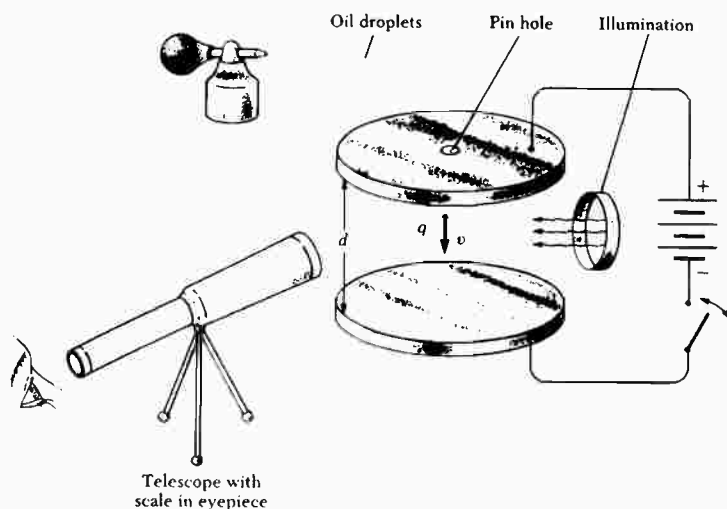


Fig. 2.5. A view of the Millikan oil-drop apparatus.

Suppose the radius of the oil-drop is a , the densities of oil and air are ρ and σ respectively, and the viscosity of air is η . When the drop without a charge falls steadily under gravity it has a terminal velocity v_1 .

upthrust force + viscous force = weight of drop

$$\frac{4}{3} \pi a^3 \sigma g + 6 \pi \eta a v_1 \text{ (Stoke's law) } = \frac{4}{3} \pi a^3 \rho g$$

$$6 \pi \eta a v_1 = \frac{4}{3} \pi a^3 (\rho - \sigma) g$$

Therefore, $a = (9 \eta v_1 / 2(\rho - \sigma) g)^{1/2}$

Suppose the drop now acquires a negative charge e' and an electric field intensity, E , is applied to oppose gravity, so that the drop now has a terminal velocity, v_2 . The force due to E on the drop is Ee' . Therefore,

$$\frac{4}{3} \pi a^3 \sigma g + 6 \pi \eta a v_2 = \frac{4}{3} \pi a^3 \rho g - Ee'$$

$$E e' = \frac{4}{3} \pi a^3 (\rho - \sigma) g - 6 \pi \eta a v_2$$

Hence,

$$E e' = 6 \pi \eta a v_1 - 6 \pi \eta a v_2 = 6 \pi \eta a (v_1 - v_2)$$

Thus:

$$e' = \frac{6 \pi \eta}{E} (9 \eta v_1 / 2(\rho - \sigma) g)^{1/2} (v_1 - v_2)$$

Knowing experimentally the values of v_1 , v_2 , ... etc., we get the value of the charge e' on the drop. Repeating the experiment several times, Millikan found that the charge on the drop was always a simple multiple of the basic unit, $e = 1.6 \times 10^{-19}$ C.

Example 2.1:

Calculate the radius of the oil-drop in Millikan's experiment from the following information:

Density of oil = 900 kg/m^3 ; terminal velocity = $2.9 \times 10^{-2} \text{ cm/s}$;
viscosity of air = $1.8 \times 10^{-5} \text{ N/m}^2$; (ignore density of air).

If the charge on the drop is $-3e$, what potential difference must be applied between the two plates 5 cm apart in order to keep the drop stationary between them? ($e = 1.6 \times 10^{-19} \text{ C}$).

Solution:

At the terminal velocity, force due to viscous drag = weight of spherical drop.

$$a = (9 \eta v / 2 \rho g)^{1/2} = 1.6 \times 10^{-4} \text{ cm}$$

Assuming V to be the potential difference between the two plates to keep the drop stationary, the electric field is (V/d) .

$$\text{The upward force} = E \times 3e = \frac{4}{3} \pi a^3 \rho$$

$$\text{Therefore,} \quad E = V/d = 4 \pi a^3 \rho / 9 e$$

$$V = 1600 \text{ volts.}$$

B.6. The classical free electron theory

From where electrons come?

Drude and Lorentz pictured a metallic conductor as an array of positive ions permeated by a gas of free electrons in which the energy distribution was assumed to be Maxwellian. The electrons were supposed to behave towards each other as uncharged particles, in spite of their electrostatic repulsion. The electrons were thus strictly free and not bound to their parent atoms.

In spite of all these simplifying assumptions the theory succeeded in the following respects:

1. It allowed the derivation of the functional form of Ohm's law connecting electric current with electric field.
2. It proved the validity of the empirical Wiedmann-Franz law stating that the ratio of electrical to thermal conductivities at the same temperature is the same for all metals.
3. The optical properties of metals concerning opacity and luster were also explained as follows. The free electrons in the metal could oscillate in an alternating electromagnetic field forming an incident beam of light. Absorption of the incident photon energy at all wavelengths might thus take place and hence the metal appears opaque. The metallic luster comes from the emission of the excited electrons at the surface to radiation when they fall back to lower energy levels.

B.7. The electron gas in a metal

Metals usually have a small number of electrons in their outermost shells. Metallic atoms give up their outer electrons so that no particular electron is assigned to a particular atom. These electrons are responsible for the binding between the atoms, also they are responsible for the observed high electrical and thermal conductivities of metals.

A metal crystal might be pictured as formed of an array of positive ions embedded in a uniform sea of electrons, or an electron cloud. The flexibility of the metals is attributed to the metallic bond arising from the attraction of the ions to the electron cloud. The ions are thus not tied to each other directly as in other types of bonding, e.g. ionic and covalent. This flexibility gives rise to ductile and plastic properties of metals. The absorption of the electron cloud to all wavelengths of light incident on its surface, and the subsequent emission of all these wavelengths give the metals the well known luster and glitter as in the case of gold, silver, platinum and copper.

B.8. Microscopic nature of electric current

Drude and Lorentz pictured a metallic conductor as formed of an array of positive ions permeated by a gas of free electrons. The electrons were supposed to behave towards each other as uncharged particles. The periodic lattice field of the positive ions was supposed to be smoothed out into a uniform potential and the electrons were thus considered to be strictly free and are not bound to their parent atoms.

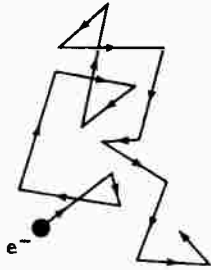
An experimental evidence for the existence of the free electron gas was given by Tolman. A piece of metal wire was accelerated sharply. Electrons were thrown due to their inertia to the back end of the metal thus producing a detectable current that could be measured by a sensitive galvanometer. It was found that the charge-mass ratio, e/m , for the particles producing the Tolman effect was the same as that of electrons.

B.9. Electrical conductivity and Ohm's law

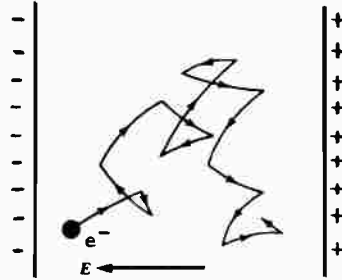
(Microscopic nature of conductivity)

In spite of the Drude and Lorentz simplifying assumptions of the free electron theory of metals, yet it succeeded in the derivation of the functional form of the empirically well known Ohm's law connecting electric current with electric field.

Consider the effect of an electric field X on a classical free electron gas of density n electrons per unit volume. The electrons move around in a random way with zero drift velocity since there are just as many electrons moving in one direction as in the opposite direction (Fig. 2.6).



Random motion of electrons.



Drifting motion in electric field.

Fig. 2.6.

We introduce the concept of relaxation time, τ , usually defined as the average time taken by the electron to describe a free path, λ . If c is the average molecular velocity of the electron, then

$$\tau = \lambda / c$$

The relaxation time governs the establishment of equilibrium through collisions, from an initial disturbed state in which the drift velocity is not zero, and it depends on the scattering factors of the conduction electrons.

In the absence of an external field, no forces act on the electron and its motion is free and satisfies the equation:

$$(dv/dt) + v/\tau = 0$$

where v is the drift velocity.

If v_0 is the steady drift velocity when an electric field is applied, and if v_t is the drift velocity after a time t from the moment of removing the field, then:

$$v_t = v_0 \exp(-t/\tau)$$

The equation of motion for the electron drift in the presence of field is given by:

$$m (dv/dt + v/\tau) = F$$

where $F = X e$, is the average external force acting on an electron, and m is the electron mass. When the drift velocity does not change with time, the term dv/dt equals zero. A particular solution of the equation is thus given by:

$$v = X e \tau / m$$

The drift mobility is μ_D and is defined as the drift velocity per unit electric field, i.e.

$$\mu_D = v / X = e \tau / m$$

The electric current density, I , is defined as the electric charge passing perpendicular through unit area in unit time:

$$I = n e v$$

Thus in the steady state we have:

$$I = (n e^2 / m) X$$

showing a direct proportionality between electric current and field, which is Ohm's law.

The electrical conductivity, σ , is given by:

$$\sigma = I / X = n e^2 \tau / m = n e^2 \lambda / m c$$

But from the classical theory of a perfect gas, the kinetic energy, E , of the electron is given by:

$$E = \frac{1}{2} m c^2 = \frac{3}{2} k T$$

Therefore,

$$\sigma = \lambda n e^2 c / 3 k T$$

This equation gives the variation of electrical conductivity, σ , with temperature, T . According to this equation, in a perfect lattice and at absolute zero of temperature, the lattice conductivity should be infinite and the resistivity zero. But, this was not found to be the case, since experimentally, it was found that the resistivity approached a value ρ_0 at 0 K. This residual resistivity was attributed to the scattering of electrons on impurities and lattice imperfections (see Fig. 2.7).

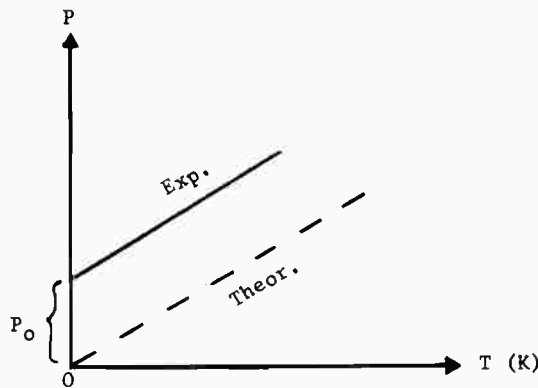


Fig. 2.7. Temperature dependence of a normal conductor. (Notice the residual resistivity at 0 K).

B.10. The electron gun in the electron microscope or in the tube of the TV set

All objects in the Earth's gravity could be considered as dropped in a potential well. In order to send an object, e.g. out of the Earth's gravity this object has to be supplied with enough energy to

overcome the potential barrier, afterwhich it becomes free from the grasp of the Earth.

In a similar manner, the electrons of the free electron gas in a metallic conductor are dropping in a potential well of an energy barrier called the work function. If the electron is given by some means energy equivalent to this work function, it will be freed out of the metal. The thermoionic emission of electrons from bodies is the process of releasing electrons by heat. This forms the principal practical source of electrons in commercial electron tubes, fluorescent lamps, and in all electronic apparatuses that need electrons.

The supply of electrons usually consists of a fine tungsten wire, which is heated to a high temperature when a low voltage source of 4-6 volts is connected to it. When the temperature of the filament is high enough, the thermal velocities of the electrons inside the metal will be increased. The chance of electrons escaping from the attraction of the positive ions, fixed in the lattice, will then be raised. Thus electrons will be boiled off and become released from the filament, forming a space charge around it. When a positively charged anode is introduced near the filament, it will attract the electrons thus forming an electric current, like that happens in a diode valve. The relation between the anode current and anode potential is called the diode's characteristic curve, shown in Fig. 2.8. The current increases with the positive anode potential until an anode potential that collects all the electrons emitted by the filament, and the current is said to be saturated.

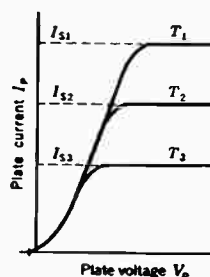


Fig. 2.8. Characteristic curves for a diode valve.

B.11. Photons and electrons

Plank's quantum energy hypothesis did not provide a new view of light. The oscillators discussed by Plank were thought to be atomic oscillators. In 1905, Einstein extended the idea of quantization to light itself, thus giving an answer to the puzzling question of the photoelectric effect.

In the photoelectric effect, electrons are emitted from a material when light is incident on its surface. The light must supply the electron with sufficient energy to escape from the surface. The minimum energy required to release the electron is called the **work function**, ϕ . It depends on the material.

Einstein proposed that electromagnetic radiation consists of particle-like packets of energy. He regarded a light wave of some given frequency, f , as a stream of energy packets, he called **photons**, each with one quantum of energy, hf , where h is Plank's constant. Each photon travels at a speed of light, c . When a photon interacts with an electron at the surface of a metal, the electron acquires all the energy of the photon, which exists no longer. If the photon energy, hf , given to the electron is just equal to the work function, ϕ , the electron will be just released from the solid. If the photon energy is greater than ϕ the electron will be released and the excess energy will appear as kinetic energy of motion for the electron. The principle of conservation of energy requires that the following equation be valid:

$$hf = \phi + \text{kinetic energy of electron}$$

The kinetic energy of the ejected photoelectrons can be determined using the apparatus shown in Fig. 2.9. An evacuated glass bulb contains a metal plate, C, connected to the negative terminal of a battery. Another metal plate, A, is maintained at a positive potential by the battery. When the bulb is kept in the dark, the ammeter reads zero, indicating that there is no current in the circuit. When monochromatic light of frequency, f , is incident on the electrode C, electrons are ejected from it and travel to the collecting electrode A. The ammeter will detect the flow of a current.

If we reverse the polarity of the cathode C and the electrode A, by putting negative voltage on A and a positive voltage on C, then the photoelectrons are repelled by the negative plate A. At a certain negative potential, V_s , called the stopping potential, the electrons will stop from reaching the plate A and no current will thus be detected. The stopping potential is independent of the radiation intensity. The measured value of this stopping potential gives us the kinetic energy of the ejected electrons:

$$\text{Kinetic energy} = e V_s$$

The kinetic energy increases with the increase of the frequency of incident light. A linear relationship is observed as shown in Fig. 2.10.

Fig. 2.9. Circuit diagram for observing the photoelectric effect. When light strikes the electrode C, photoelectrons are ejected from the plate. The flow of electrons to plate A constitutes a current in the circuit. The polarity of C and A is reversed when we need to measure the stopping potential, V_s .

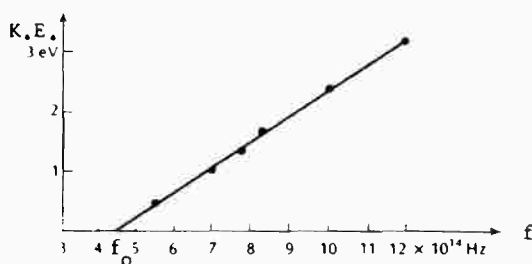
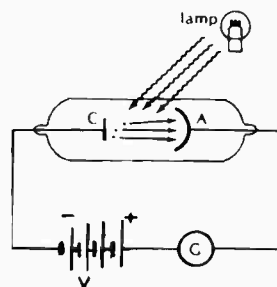


Fig. 2.10. A sketch of the maximum kinetic energy versus frequency of incident light for photoelectrons in a typical photoelectric effect experiment. f_0 is the threshold frequency below which no electrons could be ejected from the metal. The work function, $\phi = h f_0$.

B.12. The Compton effect

Very clear experimental evidence of the particle-like behavior of photons was given by Compton who was investigating the scattering of X-rays by a target of graphite. He found that the wavelength of the scattered X-rays was larger than that of the original X-rays. This was attributed to the collisions of photons with the electrons of the target. The electrons pick up some of the photon energy of the incident X-rays. The deflected photon is left with reduced energy implying an increase in wavelength.

For a quantitative discussion of the photon-electron collision, we need an expression for the momentum of a photon. Einstein concluded that a photon of energy E traveling in a certain direction will carry a momentum equal to E/c or hf/c . Compton used Einstein's idea for photon momentum in the scattering of X-ray photons from electrons, and conserved the energy and momentum of the photon-electron pair in a collision. Fig. 2.11 shows the quantum picture of the transfer of momentum and energy between an individual X-ray photon and an electron.

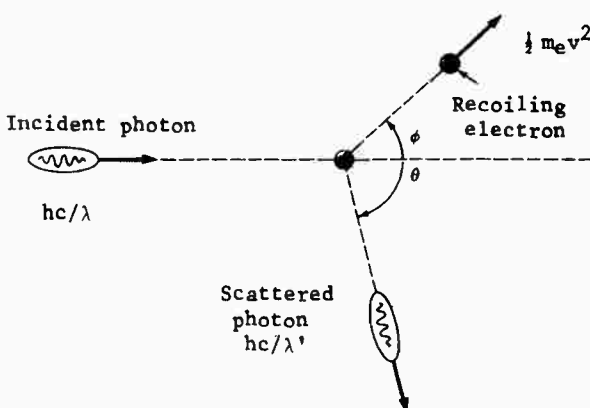


Fig. 2.11. Diagram representing Compton scattering of a photon by an electron. The scattered photon has less energy (or longer wavelength) than the incident photon.

Applying the conservation of energy to this process gives:

$$h f = h f' + K_e \quad \text{or} \quad hc/\lambda_0 = hc/\lambda' + K_e$$

where K_e is the kinetic energy of the recoil electron and λ' is the wavelength of the scattered photon.

Before collision the electron is at rest and the photon has a momentum hf/c . After collision, the electron has a recoil momentum $m_e v$ and the photon has a reduced momentum hf'/c . Conservation of momentum demands that the sum of the two final momentum vectors equal to the initial momentum vector. From these two conservation laws Compton arrived at his formula:

$$\Delta\lambda = (h/m_e c) (1 - \cos \theta)$$

where θ is the scattering angle and $\Delta\lambda$ is the change in wavelength effected by collision with the electron. For the forward incident direction, i.e. $\theta = 0$, the factor $(1 - \cos \theta) = 0$ and there is no change in wavelength.

Compton's measurements were in excellent agreement with the predictions of the above equation. It could be stated that these experimental results of Compton effect form the first proof of the validity of the quantum theory.

B.13. The electric field

We defined in chapter one, the gravitational field, g , as the gravitational force, F on an object of unit mass. The definition of electric field is similar to that of the gravitational field. Suppose a unit charge, q , is placed at a point P near a group of charged particles. The electric field, E , at P due to the group of charges is defined as the electric force, F , exerted by the charges on the unit charge. If a charge q_0 is placed at the point P then the force on it will be:

$$F = q_0 \cdot E$$

The electric field due to a point charge q at a point P distant r from the charge is given by Coulomb's law as:

$$E = \frac{1}{4\pi\epsilon_0} \frac{q}{r^2}$$

where ϵ_0 is the permittivity of space and has the value $8.854 \times 10^{-12} \text{ C}^2/\text{N.m}^2$.

B.13.1. Comparison between electric and gravitational fields

We consider the case of the hydrogen atom which is formed of a proton and an electron separated by a distance of about $5.3 \times 10^{-11} \text{ m}$.

From Coulomb's law, we find that the attractive electrical force has the magnitude: ($e = 1.6 \times 10^{-19} \text{ C}$)

$$F_e = \frac{1}{4\pi\epsilon_0} \frac{e^2}{r^2} = 8.2 \times 10^{-8} \text{ N}$$

Using Newton's universal law of gravity we find the gravitational force between the electron and proton:

$$F_g = G \frac{m_e m_p}{r^2} = 3.6 \times 10^{-47} \text{ N}$$

Thus the ratio: $F_e/F_g = 3 \times 10^{39}$

It could be seen that the gravitational force between charged atomic particles is negligibly small compared with the electrical force.

B.13.2. Motion of charged particles in a uniform electric field

The motion of a charged particle in a uniform electric field is equivalent to that of a projectile moving in a uniform gravitational field. A particle of mass m and charge q placed in an electric field E will experience an electric force qE . If this is the only force acting on the charge, then using Newton's second law of motion we get:

$$F = qE = ma$$

where the acceleration of the particle, a , will be given by:

$$a = q E / m$$

If E is uniform, the acceleration is constant and will be in the direction of the electric field. For a negative charge, the acceleration will be in the direction opposite to the electric field.

Suppose an electron of charge $-e$ is ejected horizontally with velocity v_0 into a uniform field (see Fig. 2.12) of intensity E , produced by two charged plates. The electron undergoes a downward acceleration (opposite E). The path of the electron will be parabolic.

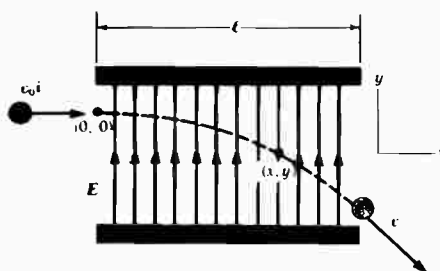


Fig. 2.12. An electron is projected horizontally into a uniform electric field produced by two charged plates.

The vertical velocity of the electron after it has been in the electric field a time t is given by:

$$v_y = a t = - (e E / m) t$$

The coordinates of the electron after a time t in the electric field are given by:

$$x = v_0 t \quad ; \quad y = \frac{1}{2} a t^2 = -\frac{1}{2} (e E / m) t^2$$

It could be seen by eliminating t from the two equations, that the trajectory is a parabola since y is proportional to x^2 .

B.14. The electron microscope

The electron microscope relies mainly on the wave-characteristics of electrons. It is similar in many respects to an ordinary compound microscope. The important difference between the two is that the electron microscope has a much greater resolving power because electrons can be accelerated to very high kinetic energies, giving them very short wavelengths. Any microscope is capable of detecting details that are comparable in size to the wavelength of the radiation used to illuminate the object. The wavelength of electrons might be 100 times shorter than those of the visible light used in optical microscopes. Thus, electron microscopes are able to distinguish details about 100 times smaller.

B.14.1. Wavelength associated with a moving electron

De Broglie, a french physicist, thought that since a beam of electrons passing through a narrow slit experiences diffraction phenomenon, like light, and since light photons behave like particles, then electrons could be considered as a new kind of object having both particle and wave properties. De Broglie introduced the idea that matter might have a dual nature. He started with Einstein's relation for the energy E of a photon of frequency f :

$$E = h f \quad \quad h \text{ being Plank's constant}$$

The relation between energy and momentum for photons is given by

$$P = E / c \quad \quad c \text{ being the velocity of light}$$

The relation between wavelength λ and frequency of light is

$$c = f \lambda$$

An expression relating momentum and wavelength could thus be obtained

$$P = E / c = h f / c = h / \lambda$$

Thus: $\lambda = h / P$

is called the De Broglie wavelength.

An electron dropping in a potential V volts will acquire an energy equal to (eV) which appears as kinetic energy $\frac{1}{2} mv^2$, where v is the electron's terminal velocity. The De Broglie wavelength associated with the moving electron will be:

$$\lambda_e = h / m v$$

It could thus be seen that increasing the electron velocity by using an electron gun, would decrease the associated De Broglie wavelength. An increased resolution could thus be obtained when such electrons are used in an electron microscope to identify small objects.

B.14.2. Operation of the electron microscope

The electron microscope consists mainly of an electron gun, magnetic lenses and a fluorescent screen. Electrons are first accelerated by the electron gun. The beam of electrons then falls on a thin slice of the material to be examined. The examined material should be very thin, typically a few hundred angstroms, in order to minimize the scattering and the absorption of electrons. The electron beam is controlled by electrostatic or magnetic deflection. The magnetic lines of force of the magnetic lens act on the charged electrons deflecting them to a focus forming an image to the body under test. Another magnetic lens forms an image on a fluorescent screen. This is done in the same way as an eyepiece in an ordinary microscope examining the image formed by the objective. The fluorescent screen is necessary because the image produced would not otherwise be visible (see Fig. 2.13).

The electron microscope is used to show strands of DNA, or deoxyribonucleic acid at very high magnification. DNA is found in the nuclei of cells. It is a long molecule made by stringing together a large number of nitrogenous base molecules on a backbone of sugar and phosphate molecules. The base molecules are of four kinds, the same in all living organisms. But the sequence in which they are

strung together varies from one organism to another. This sequence spells out a message - the base molecules are the letters in the words of this message. The message contains all genetic instructions governing the metabolism, growth, and reproduction of the cell.

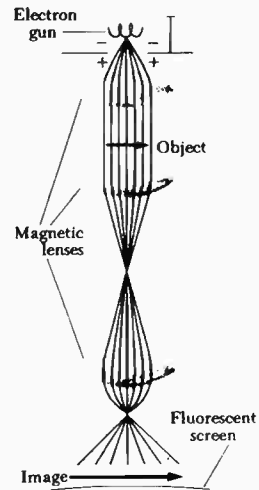


Fig. 2.13. Schematic diagram of an electron microscope. The lenses that control the electron beam are magnetic deflection coils.

The strands of DNA shown in Fig. 2.14 are encrusted with a variety of small protein molecules. At intervals, the strands of DNA are wrapped around larger protein molecules that form lumps looking like the beads of a necklace.

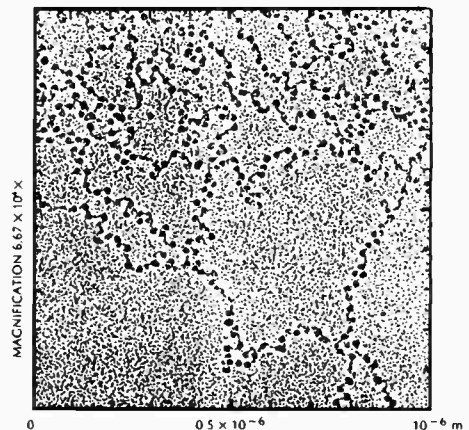


Fig. 2.14. Electron microscope image for strands of DNA. Magnification about 100,000.

B.14.3. The field-ion microscope

It is a known phenomenon that charges accumulate on sharp and pointed objects. Accordingly, the electric field intensity can be very high in the vicinity of pointed charged conductors. A device that makes use of this intense field is the **field-ion microscope**, which was invented by Muller in 1956. This microscope provides a magnification of 1,000,000 which allows seeing individual atoms on surface layers.

The basic construction of the field-ion microscope is shown in Fig. 2.15. The specimen to be studied is made in the form of a fine wire with a very sharp tip, which could be done by etching the wire in an acid. The diameter of the wire tip is about $0.1\ \mu\text{m}$. The specimen is

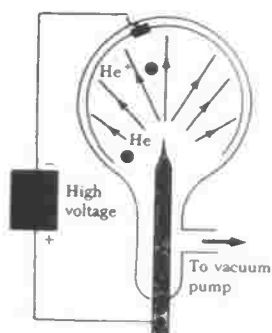


Fig. 2.15. Schematic diagram of a field-ion microscope.

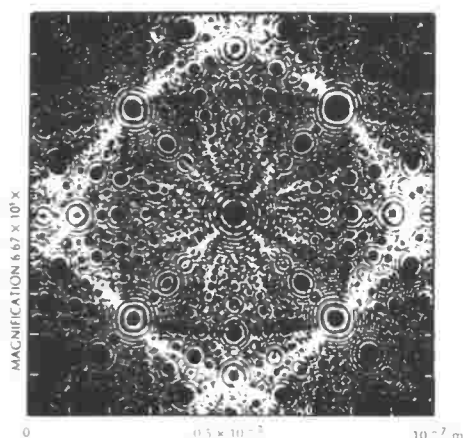


Fig. 2.16. Field-ion microscope picture of the surface of a platinum platinum crystal.

placed at the center of an evacuated glass bulb containing a fluorescent screen. A small amount of helium gas is introduced in the bulb. A very high potential difference is applied between the needle and the screen, thus producing a very intense electric field near the tip of the wire. It is important to cool the tip to liquid nitrogen temperatures in order to obtain stable pictures. The helium atoms in the vicinity of this high field region are ionized by loss of an electron

leaving helium positively charged. The positively charged He^+ ions accelerate to the negatively charged fluorescent screen. This results in a pattern on the fluorescent screen that represents an image of the tip of the specimen. Images of the individual atoms on the surface of the sample are visible, and the atomic arrangement on the surface can be studied. Fig. 2.16 represents a typical field-ion microscope pattern of a platinum crystal.

C. ATOMIC SPECTRA

C.1. Atomic structure and spectral lines

When light emitted by a light source is analyzed by a prism a spectrum is formed. The prism breaks the light into its component colors. All substances at some high temperature emit continuous distribution of wavelengths. The shape of the distribution depends on the temperature (see Wien's displacement law). In contrast to this continuous spectrum, we get discrete line spectrum emitted by excited atoms in an electric discharge tube containing gas at a very low pressure. The glow of light coming out of the discharge tube, then viewed through a spectrometer, shows a series of spectral lines of different wavelengths. The first series discovered was the Balmer series for hydrogen, Fig. 2.17. Balmer discovered a simple mathematical formula for the wavelengths of the visible lines in the hydrogen spectrum. His empirical equation was:

$$\frac{1}{\lambda} = R_H \left(\frac{1}{2^2} - \frac{1}{n^2} \right) \quad n = 3, 4, \dots$$

where $R_H = 1.097 \times 10^7 \text{ m}^{-1}$, is the Rydberg constant for hydrogen.

Other line spectra for hydrogen were found following Balmer's discovery. The Lyman series in the ultraviolet has wavelengths given by:

$$\frac{1}{\lambda} = R_H \left(\frac{1}{1^2} - \frac{1}{n^2} \right) \quad n = 2, 3, \dots$$

The Paschen and the brackett series in the infrared range, etc. All these series were combined into a single general formula:

$$\frac{1}{\lambda} = R_H \left(\frac{1}{n_2^2} - \frac{1}{n_1^2} \right) \quad n_1 > n_2$$

The important conclusion from the above equation for the spectral series is that the frequencies of the spectral lines of hydrogen are written as differences between two terms:

$$c R_H / n_2^2 \quad \text{and} \quad c R_H / n_1^2, \quad c \text{ being the velocity}$$

of light.

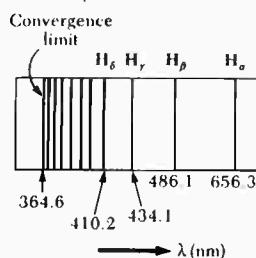


Fig. 2.17. The Balmer series of spectral lines for hydrogen (emission spectrum).

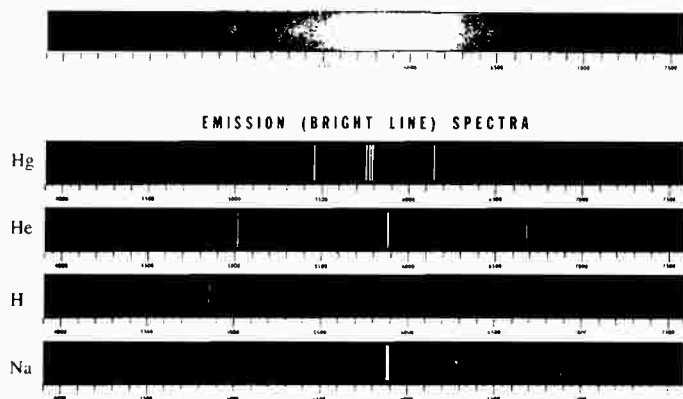


Fig. 2.18. Line spectra of some elements.

C.2. The Bohr quantum model of the hydrogen atom

Classical physics could not explain the characteristics of atomic spectra, and why hydrogen can emit only certain lines in the visible spectrum. Bohr presented a model for the hydrogen atom based on Plank's quantum theory, Einstein's photon theory of light, and

Rutherford's nuclear model of the atom. The Bohr postulates for the hydrogen atom are:

1. The electron moves in circular orbits about the proton under the influence of Coulomb force of attraction.
2. Only certain orbits are stable, in which the electron does not radiate. In classical electrodynamics, an electron moving in an orbit radiates energy. If it is so, then the electron would have to decay and spiral into the nucleus. That is why Bohr assumed the existence of stable orbits that do not radiate energy.
3. Radiation is emitted by the atom when the electron jumps from an outer orbit to an inner one. The frequency of the light emitted is related to the change in the atoms energy according to Plank-Einstein formula:

$$E_i - E_f = h f$$

where E_i and E_f are the initial and final energies.

4. Bohr quantized the orbital angular momentum of the electron. For an electron in a circular orbit of radius r and speed v , the Bohr model requires that the orbital angular momentum $L = m v r$, be a positive integer multiple of $h/2\pi$, i.e.

$$L = n h / 2\pi \quad n = 1, 2, 3, \dots$$

The radius of each orbit can be obtained by applying Coulomb's law and the quantization condition:

$$\frac{m v^2}{r} = \frac{1}{4\pi\epsilon_0} \frac{e^2}{r^2}$$

$$m v r = n h / 2\pi$$

$$\text{Thus,} \quad r_n = \epsilon_0 h^2 n^2 / \pi m e^2 \quad n = 1, 2, 3, \dots$$

The radius of the orbit depends on the integer n . For $n = 1$, we get Bohr's orbit of radius:

$$a_0 = \epsilon_0 h^2 / \pi m e^2$$

The energy E_n of an electron in the n^{th} orbit is thus:

$$E_n = - m e^4 / 8 \epsilon_0^2 h^2 n^2 \quad n = 1, 2, 3, \dots$$

The ground state for the hydrogen atom corresponding to $n = 1$, is given by:

$$E_1 = - 2.17 \times 10^{-18} \text{ J} = - 13.6 \text{ eV}.$$

C.3. The energy level diagram

The term orbit, in Bohr's theory, quickly fell out of use, because the orbits themselves could not be seen; nothing definite could be said about the paths taken by electrons. Physicists came to speak of "energy levels" instead, as most of the relevant data available at that time consisted of energies.

The energy ascribed to an orbit is negative because it is a potential energy. A low energy orbit has low n and a high energy orbit has high n . Because there is no upper limit to the value of the quantum number n , there is no upper limit to the value of the quantum number n , there is no limit to the upper value of E_n , however, as n becomes very large, E_∞ approaches zero. In Bohr's model this would correspond to an electron at a great distance from a nucleus experiencing hardly any attraction towards it at all. The electrons closer to the nucleus would feel an increasingly stronger attraction and so would have an increasingly negative potential energy. This is illustrated in the energy level diagram shown in Fig. 2.19. The energy levels in an atom can be compared with a set of steps having different depths. A ball placed on different steps represents an electron in different energy levels in the atom. Moving down from one step to another corresponds to the release of a precise amount of energy, responsible in the hydrogen atom for the lines of the Balmer and spectral series in the hydrogen spectrum. There are

no in-between lines because there are no in-between steps for the electron to rest.

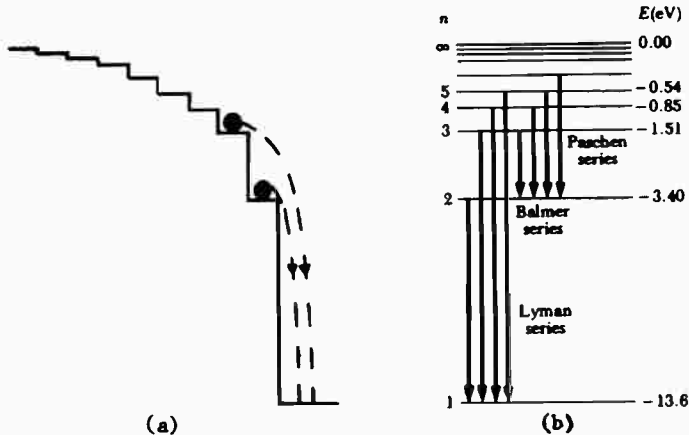


Fig. 2.19. a. Energy levels in a simple atom can be compared with a set of steps having different depths. b. Energy level diagram for hydrogen showing the transitions responsible for the Lyman and the Balmer series.

C.4. The hydrogen atom and the periodic table

We have seen that the allowed energies for the hydrogen atom are given by:

$$E_n = -\frac{m e^4}{8 \epsilon_0^2 h^2 n^2} = -\frac{13.6}{n^2} \text{ eV} \quad n = 1, 2, 3, \dots$$

An expression for the radius of the stationary orbits, r_n , is given by:

$$r_n = \frac{\epsilon_0^2 h^2 n^2}{\pi m e^2} \quad n = 1, 2, 3, \dots$$

The orbit for which $n = 1$ has the smallest radius; and is called the Bohr radius a_0 , and has the value:

$$a_0 = \frac{\epsilon_0^2 h^2}{\pi m e^2} = 0.529 \text{ Angstrom (\AA)}$$

The lowest stationary or non-radiating state is called the ground state for which $n = 1$ and has energy $E_1 = -13.6 \text{ eV}$. The next state is the first excited state and has $n = 2$ and an energy $E_2 = E_1/2^2 = -3.4 \text{ eV}$. An energy diagram showing the energies of these discrete states is shown in Fig. 2.19. Each state is characterized by a quantum number, n , called the **principal quantum number**. The allowed energies depend only on one quantum number, n , because we dealt with the problem as one-dimensional problem.

In the three dimensional problem of the hydrogen atom, three quantum numbers are needed for each stationary state, corresponding to three independent degrees of freedom for the electron. Shrodinger solved this problem completely and presented three quantum numbers, n , L , m_l characterizing every energy state. The quantum number n is called the **principal quantum number**, L is called the **orbital quantum number**, and m_l is called the **orbital magnetic quantum number**.

There are certain important relations between these quantum numbers, as well as certain restrictions on their values. These restrictions are:

The value of n can range from 1 to infinity.

The value of L can range from 0 to $(n-1)$.

The value of m_l can range from $-L$ to $+L$.

For example, if $n = 1$ only $L = 0$ and $m_l = 0$ are permitted. The states with the same principal quantum number, n , are said to form a "shell". These are identified by the letters K, L, M, ... which designate states for which $L = 0, 1, 2, 3, \dots$. Similarly, the states having the same values of n and L are said to form a **subshell**. The letters s, p, d, f, g, ... are used to designate the states in the subshells.

C.5. The electron spin

The idea of electron spin was first introduced by Goudsmidt and Uhlenbeck to account for the observed two D-lines in the atomic spectra of sodium. Close examination of the yellow line in the sodium spectrum showed that it was a doublet differing in wavelength by about six Angstroms. This was explained by the existence of two electron states in the same quantum state of energy. Accordingly, a fourth quantum number, called **the spin quantum number** was introduced.

In order to describe the spin quantum number, it is convenient to think of the electron as spinning about an axis while moving in its orbit in the atom. Two spins are possible, clockwise and anticlockwise, usually called spin up and spin down. The energy of the electron is slightly different for the spin directions. This difference in energy accounted for the observed splitting of the yellow D-lines of sodium. The quantum numbers associated with the spin of the electron are $m_s = \pm 1/2$ for the spin up and spin down respectively.

Electron spin explains the fine structure usually observed in the spectrum of different atoms. It was verified experimentally by Stern and Gerlach who directed a beam of silver atoms (neutral particles) through a non-uniform magnetic field. They found that the beam split into two components. Classically, the beam should be spread out continuously. But due to the existence of two electrons of opposite spin in each energy level, the effect of the magnetic field on the two electrons of opposite spin will be in opposite directions.

It is known that a silver atoms have one electron in the outermost shell. This electron might have positive spin or negative spin. Thus, the effect of the magnetic field on the neutral atoms with different spin will be different, and that is why the beam splits into two.

Goudsmidt and Uhlenbeck proposed that the electron has an intrinsic angular momentum apart from its orbital angular momentum. From a classical point of view, this intrinsic angular momentum is attributed to a charged electron spinning about its own axis. Accordingly, the total angular momentum of the electron in a

particular electronic state contains both an orbital contribution and a spin contribution.

C.6. Pauli exclusion principle

After it has been recognized that the electronic state of an atom could be described by the four quantum numbers, n , L , m_l , m_s , a question arose about how many electron can have a particular set of quantum numbers. Pauli answered this question by presenting his exclusion principle. Pauli principle stated that:

"No two electrons in an atom can ever have the same four quantum numbers, i.e. no two electrons can be in the same quantum state".

Pauli thought that if this principle is not valid, then all electrons would end up in the ground state of the atom, and there would be no change in the chemical behavior of the elements. The electronic structure of complex atoms should be a succession of filled energy levels, where the outermost electrons are responsible for the chemical properties of the element.

To fill an atom's subshell with electrons we start with the lowest subshell. Once it is filled the next electron goes to the next subshell that is lowest in energy, and so on. According to Pauli exclusion principle there can be only two electrons in any orbital ($m_s = \pm 1/2$). Since each orbital is limited to two electrons, the number of electrons that can occupy the various levels is also limited. In this way the atoms of all elements were characterized and their electronic configurations specified. The periodic table, first made by Mendeleev, arranged the atoms according to their atomic weights and chemical similarities. It contained several blank spaces which he considered undiscovered elements. At present the electronic configuration of all elements are arranged in the periodic table such that the elements in a vertical column have similar chemical properties.

Electronic configuration of some elements.

Z	Symbol	Electron configuration	Element	
1	H	1s ¹	Hydrogen	} -- K-shell (1,0,0,±1/2)
2	He	1s ²	Helium	
3	Li	1s ² 2s ¹	Lithium	} -- L-shell (2,0 or 1) m _l = 0 (1,0) m _s = ± 1/2
4	Be	1s ² 2s ²	Beryllium	
5	B	1s ² 2s ² 2p ¹	Boron	
6	C	1s ² 2s ² 2p ²	Carbon	
7	N	1s ² 2s ² 2p ³	Nitrogen	
8	O	1s ² 2s ² 2p ⁴	Oxygen	
9	F	1s ² 2s ² 2p ⁵	Fluorine	
10	Ne	1s ² 2s ² 2p ⁶	Neon	

C.7. Frequency of photon emitted in transition

Let us now calculate the frequency of light emitted in a quantum jump from some initial state *i* to a final state *j*. In this jump the electron releases an energy:

$$\Delta E = E_j - E_i$$

Therefore,
$$\Delta E = -\frac{m_e e^4}{2 (4\pi\epsilon_0)^2 (h/2\pi)^2} \left(\frac{1}{n_j^2} - \frac{1}{n_i^2} \right)$$

According to Bohr's postulate, this energy is radiated as a single photon of frequency $f = \Delta E/h$, i.e.

$$c/\lambda = f = (E_j - E_i)/h = \frac{m_e e^4}{4\pi (4\pi\epsilon_0)^2 (h/2\pi)^3} \left(\frac{1}{n_j^2} - \frac{1}{n_i^2} \right)$$

This equation looks like the general formula for the spectral series, giving a theoretical formula for the Rydberg constant:

$$R_H = \frac{m_e e^4}{8 \epsilon_0^2 h^3 c} = 109,737 \text{ cm}^{-1}$$

This theoretical value of R_H agrees quite well with the experimentally obtained value.

C.8. The wave nature of matter

Electrons were considered as particles in Bohr's theory. However, when a beam of electrons is made to pass through an extremely narrow slit, the electrons show diffraction phenomena just like light. This means that electrons are neither classical particles nor classical waves. Electrons, just like photons, are new kind of objects that have combination of particle and wave properties. They were termed at that time **wavicles**.

De Broglie thought that if light could behave like both a wave and a particle, matter might also have a dual nature. To describe the wave associated with matter, De Broglie started with Einstein's relation for photons: $E = h f$, together with the relation between energy and momentum for photons: $P = E/c$, and the relation between speed c , wavelength λ , and frequency f for light: $c = f \lambda$. An expression relating momentum and wavelength can be obtained:

$$P = E/c = h f/c = h/\lambda$$

Thus,

$$\lambda = h/p = h/m v$$

where m is the electron mass, and v is its velocity. This equation shows the dual nature of matter. It contains both particle concepts (mv and E) and wave concepts (λ and f).

C.9. Deflection of electrons by an electrostatic field (The TV tube)

In 1895, J.J. Thomson applied a large potential difference between a hot filament and a fluorescent screen. Electrons came out

of the filament and accelerated to the screen. A tiny flash of light was seen at the point of their arrival. These electrons were called by Thomson at that time **cathode rays**. When these cathode rays were passed between two oppositely-charged plates they were deflected. Presently, we know that an electron charge, e , passing through an electric field of intensity, E , experiences a force equal to the product (eE). This force is responsible for the deflection of the cathode rays discovered by Thomson. A schematic representation of Thomson's apparatus is shown in Fig. 2.20. The apertures at A and B are slits in two oppositely charged plates used to accelerate and focus the electron beam. The apparatus was evacuated in order to allow the cathode rays to reach the fluorescent screen without being scattered by air molecules.

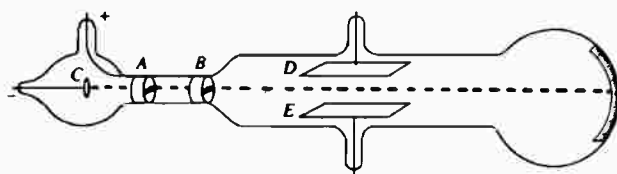


Fig. 2.20. Thomson's apparatus for cathode rays.

The modern application of Thomson's technique was the television picture tube. A hot filament emits electrons which are then accelerated through a vacuum to a fluorescent screen. In color television sets, the potential difference applied to cause this acceleration is about 30,000 volts. The beam of electrons is directed to different parts of the screen by the effect of magnetic and electric fields. By controlling these fields the beam of electrons is made to sweep rapidly the entire screen. By varying the intensity of the electron beam, the picture tube is made to glow with different intensities in different parts of the sweep, and a picture is formed by the light and dark areas thus produced. Each sweep of the screen is so rapid that someone watching the television is not aware of the

individual lines of the sweep, nor of the transition from one sweep to the next (see Fig. 2.21).

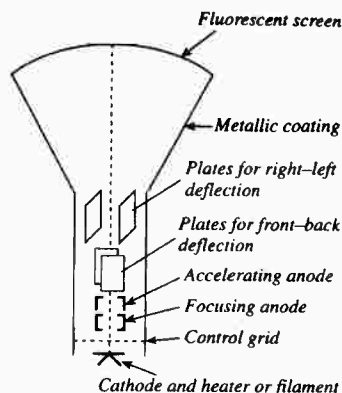


Fig. 2.21. The TV picture tube.

C.10. Motion of electrons and charges in electro-magnetic fields

In order to understand the operation of many modern devices and instruments, we must consider the motion of electrons and of charged particles in electric and magnetic fields. Electromagnetic forces dominate the motion of charged particles. If an electric field, E , and a magnetic field, B , exist in a region, then the combined force, F , acting on a particle of charge, e , and moving with a velocity, v , is given by:

$$F = eE + e v \times B$$

This force is called the **Lorentz force**.

Let us first consider the motion of an electron in a magnetic field with no electric field present. Suppose that the magnitude of the field is B and that the direction is out of the plane of the page, as shown in Fig. 2.22.

The magnetic force on the electron is perpendicular to B and to the velocity of the electron v . That is, the force and the velocity are perpendicular to each other and lie in a plane which is perpendicular

to B . If the magnetic force is the only force acting on the electron, then, from Newton's law, the acceleration of the electron is perpendicular to the velocity and also lies in the plane perpendicular to B . Since the acceleration is perpendicular to the velocity, only the direction of the velocity changes, and the path of the particle is in circles of radius r with constant speed v .

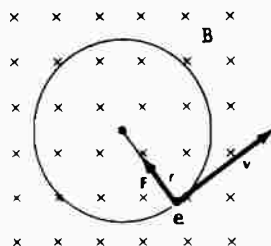


Fig. 2.22. The electron moves in a circular path perpendicular to a uniform magnetic field.

The centripetal acceleration 'a' is given by:

$$a = v^2 / r$$

The centripetal force provided by the magnetic force $= e v B$.
Newton's second law thus gives:

$$m v^2 / r = e v B \sin 90 = e v B$$

We put $\sin 90$ because the angle between v and B is 90° . In case there is an angle θ between v and B then the magnetic force will be ($e v B \sin \theta$). The above formula could be simplified thus the radius of the circular path is:

$$r = m v / e B$$

where m is the mass of the electron.

The deflection of electrons by magnetic fields forms the basis of the action of magnetic lenses. The lenses that control the electron beam in the electron microscope are magnetic deflection coils.

C.11. Complimentary colors and the colored TV

The electron beam scanning the television screen excites only three primary colors, namely, red, green and blue. Nearly all colors of the spectrum could be obtained by superposing these primary colors in varying proportions. Artists and painters know that when colored pigments are mixed we do not get the sum of their colors but only the color that is absorbed by none.

A blue body will absorb all wavelengths of the white light incident on it except the blue. A red body will reflect only the red light, and so on. The sensation of light by the eye depends on the wavelength of the photons incident, e.g. a photon of wavelength 6000 Angstrom will give the sensation of yellow light, etc. When two photons of different wavelengths fell on a human eye, it will not see a mixture of the two, but another color will be seen. Red and green photons will give the sensation of yellow, i.e. the eye sees the yellow if it is excited by red and green photons. The three primary colors could never be seen by mixing any of the other colors.

The result of mixing the three principal colors is shown in Fig. 2.23. Three circles representing the three primary colors are crossed together. Mixing of the three colors in the middle portion results in a white color.

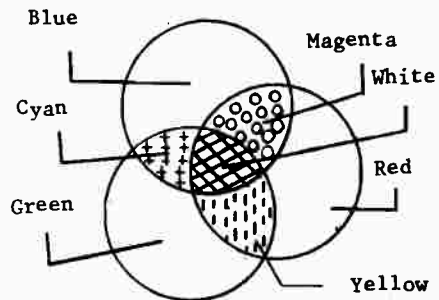


Fig. 2.23. Mixing of the primary colors.

The resulting colors of mixing each two of the primary colors are called the **complementary colors**, namely, yellow, cyan, magenta. On the other hand, the mixing of any two of the complementary colors will give a primary colors, as shown in Fig. 2.24. The mixing of any three complementary colors give the black color.

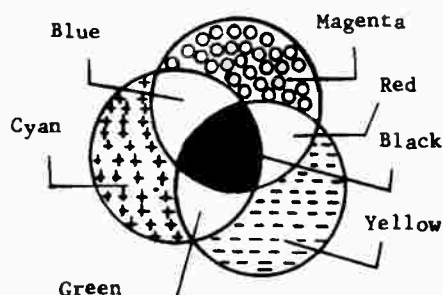


Fig. 2.24. Mixing of the complementary colors.

In color photography and color television, there is no need to reproduce every color in the exact wavelength in which it was taken. Several processes have been founded on the theory of primary colors and are in wide use. A combination of the three fundamental colors will produce the same color sensation in the eye. In color photography, they employ three color filters or screens, a red one, a green one, and a blue one. To reproduce the picture on the television screen we have to excite by the electron beam the three primary colors. This could be done by careful choice of the coating mixture of fluorescent materials fixed of the TV screen.

C.12. Fluorescence and phosphorescence

Fluorescence is the process of converting ultraviolet light to visible light. An atom absorbs a photon of energy hf_1 ends up in an excited state, E_3 . The atom can return to the ground state E_1 via some intermediate states, e.g. E_2 , as shown in Fig. 2.25. The photons emitted by the atom will have lower energy, hf_2 , when the electrons move from the intermediate state, E_2 , back to the ground state.

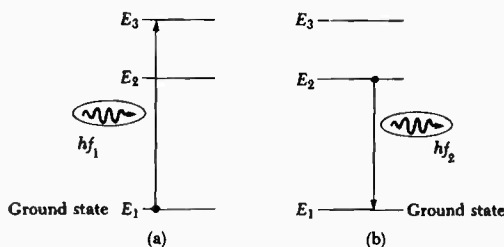


Fig. 2.25. The process of fluorescence. (a) An atom absorbs a photon of energy hf_1 and ends up in excited state, E_3 . (b) The atom emits a photon of energy hf_2 when the electron moves from an intermediate state, E_2 , back to the ground state.

In the television tube, the electrons are produced from the heated filament, then accelerated by an applied voltage on the electron gun. The accelerated electrons will cause the atoms of the fluorescent material coated on the inner surface of the tube, to get excited momentarily. Coming back to the ground state they emit visible light to form the observed picture on the TV screen.

Fluorescence analysis is sometimes used to identify compounds. Every compound has its own energy level diagram and so will fluoresce with a specific wavelength characteristic to it.

There is another class of materials that show illumination long after the exciting agent has been removed. These are called **phosphorescent** materials. The atoms of such materials remain in excited metastable state for periods ranging from a few seconds to several hours. Eventually, the excited atoms when dropping to their ground states will emit visible light long after being placed in the dark. Such materials are commonly used as paints to decorate hands of watches and clocks and to outline the exits in cinemas and large buildings to show the doors if there is a power failure.

D. HOLOGRAPHY

D.1. Formation of TV images and the holograms

Holography is a technique for producing three-dimensional images by using the interference of light waves. In order to make a hologram of an object we use a collimated beam of coherent monochromatic light from a laser source. In order to understand what is a laser beam, and how it is formed, we have to begin with the structure of the atom and its energy levels.

D.2. Laser light

Ordinary light is emitted by electrons in the atoms of light sources. The atoms emit their light spontaneously in random directions and at irregular times, over a broad spectrum, resulting in isotropic illumination of incoherent light.

Laser light originates from atoms, ions or molecules through a process of "stimulated emission" of radiation. The term **LASER** is a short hand writing coming from: **L**ight **A**mplification by **S**timulated **E**mission of **R**adiation. Laser light has three characteristics that makes it different from ordinary light:

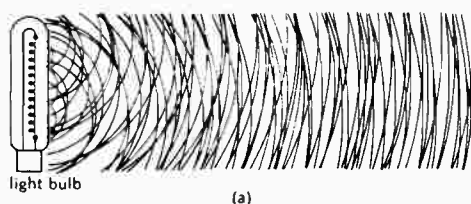
1. It is highly monochromatic.
2. It is very coherent.
3. It is well collimated, the angle of divergence is very small.

D.2.1. Stimulated emission

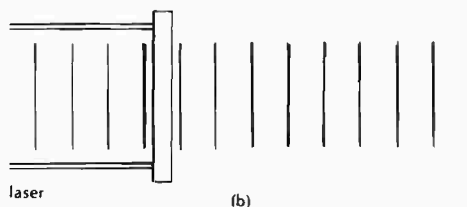
According to Bohr's model of the atom, energy must be provided to the atom in order to excite the electron to a higher energy state. When the electron returns from this excited state to the ground state, a photon of energy $hf = E_2 - E_1$, is **spontaneously emitted**. The time interval between the absorption of energy to form an excited state and the emission of the photon is unpredictable. Because different atoms in a light source emit photons at random, the light emerging from an ordinary source is a confused combination of many

different waves with no special directions of propagation and no special phase relationships (see Fig. 2.26).

In a laser, the atoms emit their light in unison. The electrons in the different atoms jump down at the same time emitting their light waves in a coherent combination in the same direction. By coherent we mean that all light waves from different atoms are in phase, i.e. the light wave contributed by each atom combines crest to crest with the light waves contributed by the other atoms. The waves combine constructively thus producing the laser beam.



Light waves emitted by ordinary high source (non coherent).



Light waves emitted by a laser (coherent).

Fig. 2.26.

Suppose a photon of energy, $hf = E_2 - E_1$ interacts with an atom that is already in the excited state E_2 . Einstein predicted that the incident photon may stimulate the atom to emit a photon very identical to the exciting one, i.e. both photons will have the same energy, phase, and direction of travel. If these two photons then interact with two more atoms in excited state, two more photons are produced, and so on. This stimulation process leads to photon amplification. In order to maintain this stimulation process (lasing action), we need to have more atoms in the excited state than in the ground state, a condition called **population inversion**. One method for achieving the required population inversion used an intense flash of light to lift electrons into excited orbits, this is called **optical pumping**.

D.3. The helium-neon gas laser

The laser tube (Fig. 2.27) contains a mixture of helium and neon at low pressure with mirrors at each end. A large electric field from a RF power supply is established in the tube by electrodes connected to the high voltage source. The electrons from ionized atoms are accelerated by the field and collide with atoms. Neon atoms are excited to state E_3 (metastable) through a process of electrical pumping as a result of collisions with excited helium atoms. Stimulated emission occurs as the neon atoms make a transition to state E_2 and the neighboring excited atoms are stimulated (see the energy diagram of Fig. 2.28).

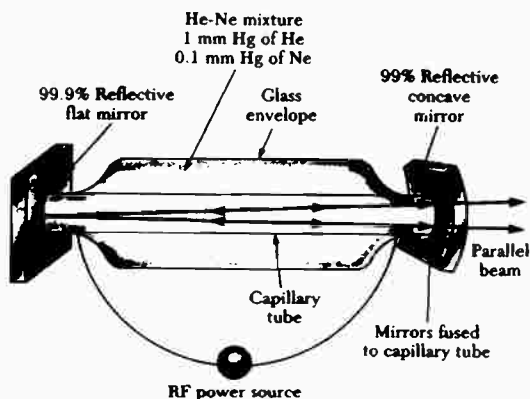
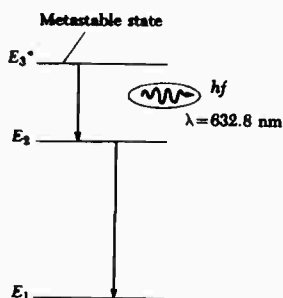


Fig. 2.27. Schematic diagram of a helium-neon laser.

Fig. 2.28. Energy diagram for the neon atom, which emits photons at a wavelength of 6328 Angstrom through stimulated emission. The photon at this wavelength arises from the transition $E_3 - E_2$.



The mirrors at each end of the tube encourage emissions along the tube axis (at the expense of emissions in other directions) by reflecting the light back and forth inside the tube. The tube is like a resonant cavity similar to an organ pipe. The atoms that are stimulated to emit are those that produce light which is in step with the light already propagating along the tube. One of the mirrors is partially reflecting and so is slightly leaky, transmitting about 1% of the incident light. This transmitted light forms the laser beam.

Lasers have presently too many applications, such as: surgical welding of detached retina; precision surveying and length measurement; source for nuclear fission reactions; precision cutting of metals and other materials; in telephone communication along optical fibers; and in the formation of three dimensional images, known as **holograms**.

D.4. Laser applications

Since the development of the first laser in 1960, there has been a tremendous growth of laser technology. Lasers are now available covering wavelengths in the infrared, visible, and ultraviolet regions. Applications include astronomical and geophysical purposes, to measure precisely the distance from various points on the surface of the Earth to a point on the moon's surface. It is also used to decode the digital information on the compact audio laser disc, the so-called **CD**. On the compact disc, the music has been digitized as pits and grooves embedded into a plastic-covered metal foiled. The fluctuation reflection of the weak laser spot from the foil surface is detected by a photocell and decoded by digital to analog circuits to reproduce music with extremely high fidelity, without noise. Besides, if a compact audio disc is used to store digital data, then the amount of information stored in this way is enormous, and is estimated to be about 1 gigabyte, i.e. 10^9 characters. The CD disc is now used to store encyclopedia or dictionary volumes.

Medical applications of lasers utilize the fact that different laser wavelengths can be absorbed in specific biological tissues. A well-known eye disease, glaucoma, is manifested by a high fluid pressure in the eye, which can lead to destruction of the optic nerve. A simple

laser operation called iridectomy, can burn open a tiny hole in the clogged membrane, thus relieving the destructive pressure.

In the case of diabetes, a serious side effect is the formation of weak blood vessels which leak blood into extremities. When this occurs in the eye, vision deteriorates (diabetic retinopathy) leading to blindness in diabetic patients. It is now possible to direct green light from argon ion laser through the clear eye lens and eye fluid, focus the laser light on the retina edges to cause photocoagulate the leaky vessels. Such operations have greatly reduced cases of blindness due to glaucoma and diabetes.

Laser surgery is now a practical reality. Infrared light from a carbon-dioxide laser can be used as a laser knife and welder for operations on blood-rich tissues, such as the liver. The cauterization of the tissues by the laser beam prevents excessive bleeding. In addition, this technique of using the laser knife, virtually eliminates cell migration, which is very important in cancer surgery and tumor removal.

In the field of energy production, powerful lasers are used to cause thermonuclear fusion of heavy hydrogen (deuterium and tritium), thus producing energy. At a temperature of 10^8 K, the violent thermal collisions between the nuclei of heavy hydrogen merge them together into helium. This reaction releases a large amount of heat. Obviously, in order to exploit this thermal energy we need to build accessories that convert the heat from fusion reaction into electric energy.

D.5. Lasers and fiber optics

Total internal reflection can occur when light comes from a medium of high index of refraction to one having a lower index of refraction. Consider a light beam traveling in medium 1 and is incident on the boundary of medium 2 (see Fig. 2.29), where n_1 is greater than n_2 . Various possible directions of the beam are indicated by rays 1 through 4. The refracted beams are bent away from the normal. At some angle θ_c , called the critical angle, the refracted light will move parallel to the boundary so that the angle of refraction is 90° . For incident angles greater than θ_c , total internal

reflection takes place, and the boundary layer will behave as a perfect reflecting mirror. According to Snell's law we have:

$$n_1 \sin \theta_c = n_2 \sin 90 = n_2$$

Thus: $\sin \theta_c = n_2 / n_1$

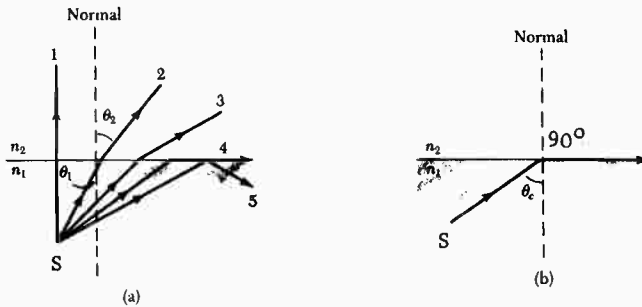


Fig. 2.29. Total internal reflection $n_1 > n_2$.

An important application of this phenomenon is in the use of glass or transparent plastic rods to "pipe" light from one place to another, as indicated in Fig. 2.30. Light is confined to traveling within the rods even around curves. Such a light pipe will be flexible if thin fibers are used. If a bundle of parallel fibers is used to construct an optical transmission line, images can be transferred from one point to another. This technique is known as **fiber optics**.



Fig. 2.30. Light travels in the fiber by multiple internal reflections.

Fiber optics is useful when we wish to view an image produced at inaccessible locations. For example, physicians often use it to examine internal organs of the body. A laser beam can be trapped in a fine glass fiber light guides (endoscope), and by means of total internal reflection one can get clear images of internal organs of the body. The endoscope can be introduced through natural orifices, conducted around internal organs and finally directed to specific interior body locations, thus eliminating the need of massive surgery. For example, bleeding in the gastrointestinal tract can be optically cauterized by fiber optic endoscope inserted through the mouth.

D.6. Formation of a hologram

A hologram is a three-dimensional image of an object. To produce a hologram, a film negative is first made by means of a laser beam and using the phenomenon of interference of light as shown in Fig. 2.31. Light from the laser is split into two parts by a half silvered mirror at B. One part of the beam reflects off the object to be photographed and then strikes an ordinary photographic film. The other hand of the beam is made to diverge by a lens L_2 , and is then reflected from the mirrors M_1 and M_2 , and finally strikes the film. The two beams overlap to form an extremely complicated interference pattern on the film. (In the next section we give some information about the interference of light and its diffraction through narrow slits and diffraction gratings). The interference pattern on the film is produced because the phase relationship of the two halves of the waves is maintained constant throughout the exposure because laser light is coherent and all photons of the beam have the same phase. The hologram records not only the intensity of light scattered by the object as in conventional photography, but also the phase difference between the reference beam and the beam scattered from the object. Because this phase difference, an interference pattern is formed on the film which will show like dark and light fringes on the photographic plate after its development. The light and dark fringes (see Fig. 2.32) on the developed film record the constructive and destructive interference of the reference and object beams.

When we shine the developed photographic plate with a laser beam and look at it from the far side, we will see an exact replica of the original object. The image is three-dimensional, by moving the eye up, down, or sideways we can bring the top, bottom or sides of the object into view. But, the image is virtual and if we try to touch the observed thing behind the hologram with a finger, we find that nothing is there.

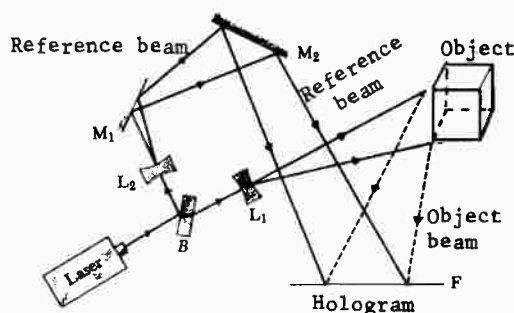
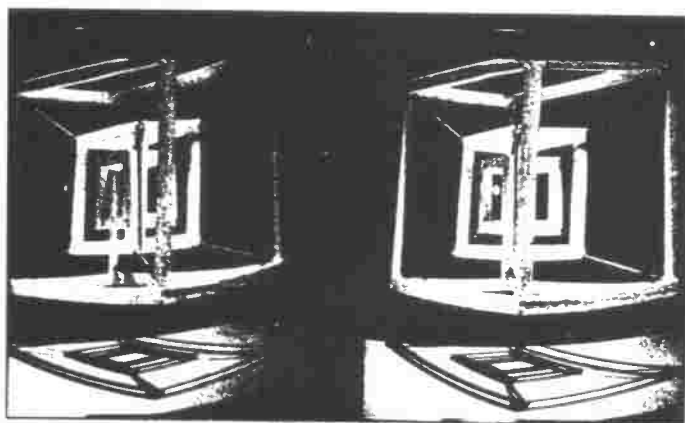
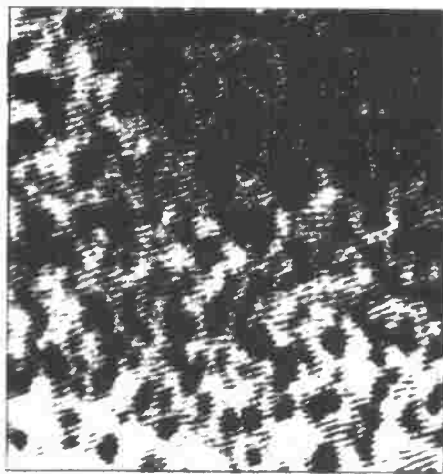


Fig. 2.31. Experimental arrangement for making a hologram.

D.7. Viewing of the hologram

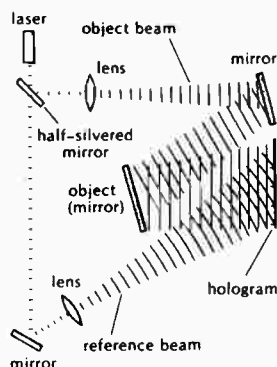
To understand how a three-dimensional image of the object is formed after the hologram has been prepared, suppose a plane mirror (Fig. 2.33) is placed instead of the object. The wavefronts emerging from the illuminated mirror are then simply plane waves and, at the photographic plate, these plane waves interfere with the plane waves of the reference beam. Since the angle of emergence of the diffracted beam, produced by the illuminated hologram, coincides with the angle of incidence of the original beam, produced by the illuminated mirror, then, the hologram reconstructs the wavefronts. It generates a light wave that has exactly the same characteristics as the light wave emitted by the object that was photographed. If we place our eyes beyond the hologram, we will see exactly what we would see if instead of the illuminated hologram we had the illuminated object in front of our eyes. We will believe that we see the object with all its three-dimensional features.

Fig. 2.32. A hologram. The dark and light fringes on this photographic plate record the constructive and destructive interference of the reference beam and the object beam.



The holographic image of a piece of modern sculpture has been photographed from two different directions, giving different points of view. (If you can de-couple your eyes, and view the left image with the left eye and the right image with the right eye, you will be able to perceive photos as three dimensional).

Fig. 2.33. Arrangement for viewing a hologram. The hologram will generate a wavefront having the same features of the wavefront emitted from the object.



D.8. Interference of light waves (Young's double slit experiment)

The wave theory of light was first proposed by Huygens. Young's double-slit experiment provides a simple demonstration of the wave nature of light. In this experiment, Young considered that light waves like other electromagnetic waves, obey the principle of linear superposition: if two waves meet at some points, the resultant electric or magnetic field is simply the vector sum of the individual fields. If two waves of equal amplitude meet crest to crest, they combine and produce a wave of doubled amplitude; if they meet crest to trough, they cancel and give a wave of zero amplitude. The former case is called **constructive interference** and the latter **destructive interference**.

A schematic diagram of Young's apparatus is shown in Fig. 2.34. Monochromatic light passes through a narrow slit and then through two parallel slits S_1 and S_2 before reaching a screen. The two slits serve a pair of coherent light sources because waves emerging from them originate from the same wavefront and therefore maintain a constant phase relationship. The light from the two slits produces a series of bright and dark parallel bands called **fringes**. When the light from slits S_1 and S_2 arrives at a point on the screen such that constructive interference occurs at that location, a bright line appears.

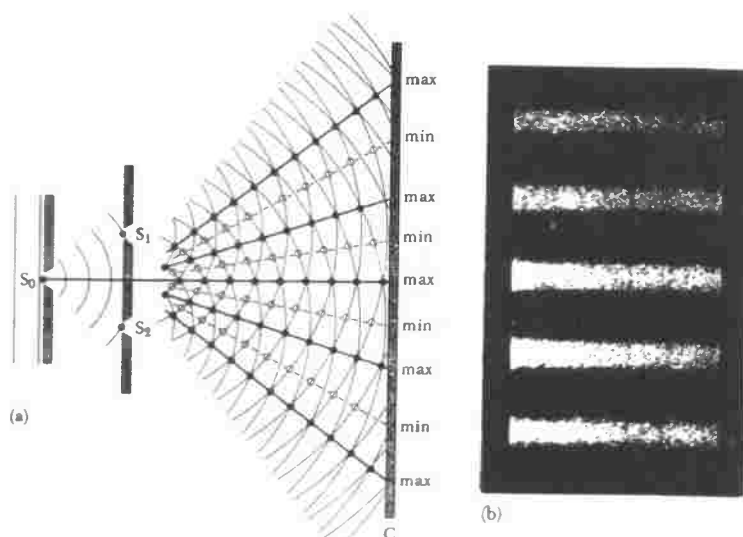


Fig. 2.34. Young's double-slit experiment. The narrow slits S_1 and S_2 behave as coherent sources which produce an interference pattern on the screen.

When the light of the two slits combines destructively at any location on the screen, a dark line results. Fig. 2.35 is a schematic diagram of how two waves combine constructively or destructively, according to the path difference of light, i.e. the difference in length traveled by the two waves to reach the screen. From the geometry of the figure the path difference is $(d \sin \theta)$. This path difference will determine whether or not the two waves are in phase when they arrive at the screen. If the path difference is zero or some integral multiple of the wavelength, the two waves are in phase and constructive interference results. Therefore, the condition for bright fringes, or constructive interference is given by:

$$d \sin \theta = n \lambda \quad n = 0, \pm 1, \pm 2, \dots$$

The central bright fringe is at $\theta = 0$, when $n = \pm 1$ we get the first order maximum, and so on, see Fig. 2.36.

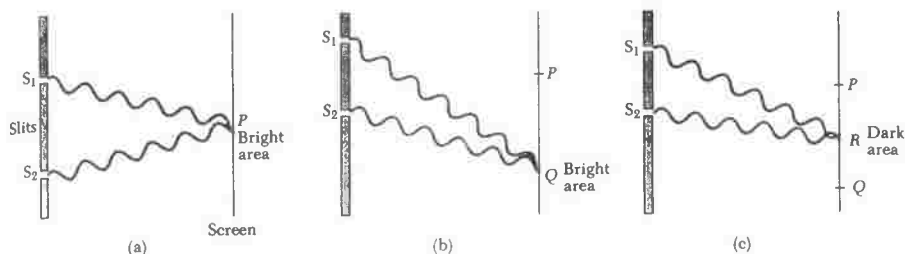
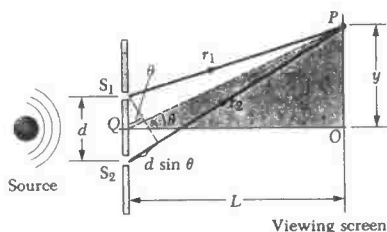


Fig. 2.35.

- a) Waves from the two slits arrive at the screen P_0 in phase, crest on crest and trough on trough so they interfere constructively producing the central bright fringe.
- b) The waves reach the screen out of phase, crest on trough and trough on crest, so they produce a dark fringe.
- c) Waves arriving at P_1 in phase so they produce the bright fringe of order 1, $n = 1$.

Fig. 2.36. Geometric construction for describing Young's double-slit experiment. Light waves from S_1 and S_2 travel different distances to reach P on the screen. The path difference between the two rays is $r_2 - r_1 = d \sin \theta$.



Similarly, for destructive interference, the path difference is an odd multiple of $\lambda/2$, the two waves arriving at the screen will be 180° out of phase, thus the destructive interference is given by:

$$d \sin \theta = \left(n + \frac{1}{2}\right) \lambda \quad n = 0, \pm 1, \pm 2, \dots$$

A demonstration of interference effects from two sources of water waves is shown in Fig. 2.37. The interference of these waves can be used to visualize the way the double slit interference pattern is formed.

Fig. 2.37. Water waves in a ripple tank give a two-dimensional analog of the double slit experiment. The water waves are set by two vibrating sources, going up and down together to create coherent waves of the same wavelength.



D.9. Magnetic resonance imaging (MRI)

(Its use in medicine)

A very important diagnostic tool in medicine is magnetic resonance imaging, which is widely used in different fields. The most successful and highly developed application in the area of medicine has been computerized X-ray tomograph imaging, which is a diagnostic technique that is noninvasive and has minimal impact on the subject being imaged.

The surface of an object as seen in an ordinary photograph is of limited value, particularly if the regions of interest are located within the interior of the object. In the medical field, we always need a tool for investigating regions of the body that lie beneath the skin, such as a broken bone, a foreign body, or a malformed internal organ. In many cases X-rays could provide the medically required information. Other modes of testing, such as ultrasound, gamma rays, might provide images of internal regions of a body. These techniques produce a single sheet of information coming from integrated radiation attenuation or the photon counts in the case of X-rays or gamma rays, or coming from reflected sound wave intensities in the case of the sonar images.

Computerized tomography (CT) is a procedure which can make use of the above sheet of information to provide, with the help of a computer, a slice or sheet image of the interior of the object. We are to apply this technique of CT to the new method of magnetic resonance imaging.

In order to get a CT scan of some cross-section of a body, assume that the body is located above the x-y plane and that X-rays uniformly illuminate the body from above. An X-ray detector with an adjustable iris is positioned beneath the body and the source as shown in Fig. 2.38. The detector is scanned over the surface of the body, thereby generating a sheet of information about the intensity of the transmitted X-ray beam from the different elements of the x-y plane. At each location the X-ray absorbance could be found and the information is fed into a computer and is subjected to an imaging algorithm which assigns to every element a brightness corresponding to the object's X-ray absorbance.

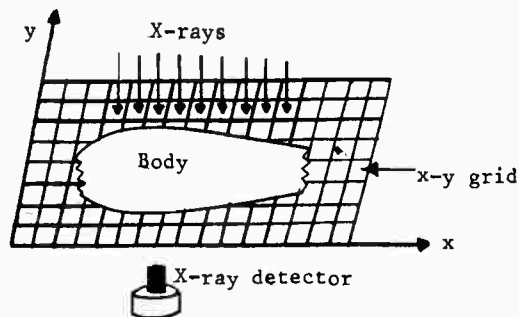


Fig. 2.38. Method of obtaining images by X-ray computerized tomography.

D.10. Magnetic resonance

To understand magnetic resonance imaging, it is necessary first to know something about magnetic resonance spectroscopy. We know from Plank's quantum theory that the energy of any system is quantized. For example, the electrons in an atom have discrete energy levels, and according to Pauli exclusion principle, an energy state cannot accommodate more than two electrons of opposite spin ($\pm 1/2$). When a magnetic field is applied to this system, each energy level will be split into two as shown in Fig. 2.39. Particles with spin

direction in the same direction as that of the magnetic field will have their energy increased, while those having their spin in the opposite direction to the field have their energies decreased.

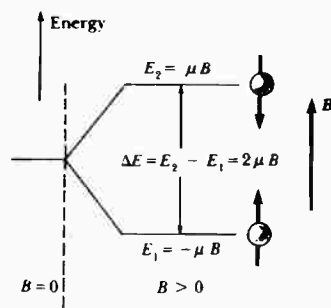


Fig. 2.39. A schematic representation of the energy dependence on magnetic field of a two-state electron quantum mechanical system.

The process here is a statistical one, and the spin might change from positive to negative and vice versa. These reversals cause the charged particles to change their energies while the magnetic field is on, from $E + \Delta E$ to $E - \Delta E$, and conversely. The population numbers of the particles will remain always the same on each energy level according to Boltzmann's distribution function, thus there will be a characteristic jump frequency from one energy level to the other.

The nuclei of atoms contain protons having spin moments. When placed in an external magnetic field a similar behavior occurs and the single energy level will be split into two energy levels due to the effect of magnetic field on the different spins $\pm 1/2$. Opposite to the case of electrons, a proton with spin $+1/2$ can occupy one of the two energy states, but the lower energy state corresponds to the spin aligned with the field and the higher energy state corresponds to the case where spin is opposite to the field.

It is possible to observe transitions between these two spin states using a technique known as nuclear magnetic resonance. A steady d.c. magnetic field is introduced, along with a second weak and oscillating magnetic field oriented perpendicular to the d.c. field. The weak field is produced from a radio-frequency (rf) generator and a coil placed around the sample, as schematically shown in Fig. 2.40. The radio-frequency of the rf generator could be changed until the

frequency of the oscillating field matches the "flipping" frequency between the two spin states. When the nuclei in the sample meet the resonance condition, the flip of spins absorb energy from the rf field of the coil, and this changes the Q of the circuit in which the coil is included. Maximum energy loss corresponds to the resonance condition.

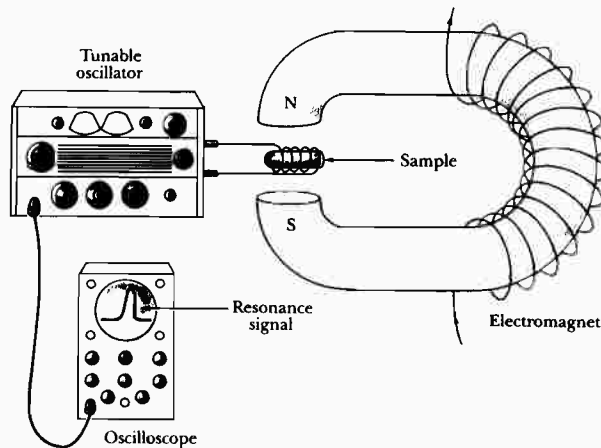


Fig. 2.40. An experimental arrangement for nuclear magnetic resonance. The rf magnetic field of the coil is provided by the variable frequency generator.

The magnetic resonance absorption function changes from place to place inside the body during the computerized tomography CT. This function is directly proportional to the integrated magnetic resonance absorption. Feeding this information to a computer used to drive an imaging algorithm, will establish an image for the test body. This method is called **MRI**, magnetic resonance imaging. It is better than using X-rays or gamma rays in CT tomography, because the photons associated with rf signals in MRI have energies of about 10^{-7} eV. In the human body, molecular bond strengths are much larger, and of the order of 1 eV, thus the rf photons cause little cellular damage. In comparison X-rays and gamma rays have energies ranging from 10^4 to 10^6 eV and have capability of causing

considerable cellular damage. This is the main advantage of MRI over other imaging techniques in medical diagnosis. Fig. 2.41 demonstrates a MRI taken of the human head.

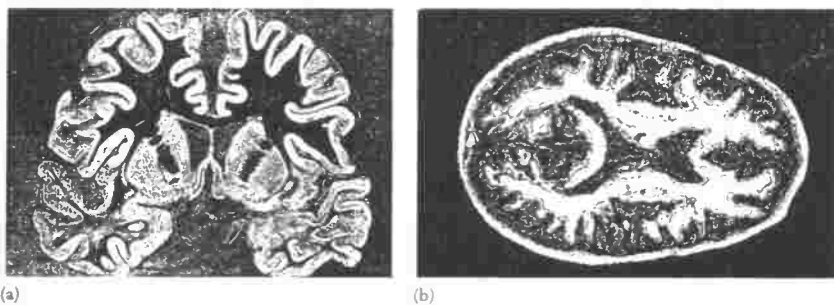


Fig. 2.41. Magnetic resonance images MRI, taken for the human head.

a) Sagittal view (horizontal section).

b) Coronal view (vertical section).

The slice images are of 1 cm thickness.