

Eghbaltalab

Unit 1

Section One: Reading Comprehension

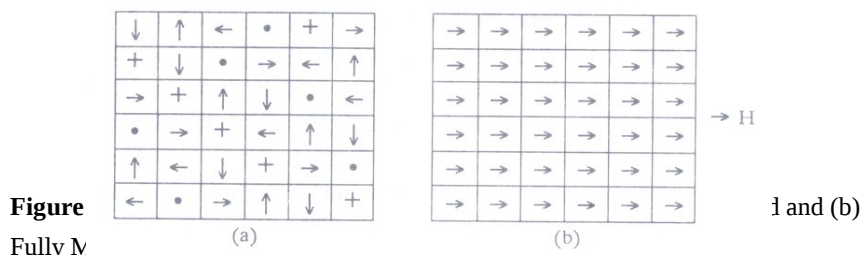
Theory of Magnetism

To understand the magnetic behavior of materials, it is necessary to take a microscopic view of matter. A suitable starting point is the composition of the atom, which Bohr described as consisting of a heavy nucleus and a number of electrons moving around the nucleus in specific orbits. Closer investigation reveals that the atom of any substance experiences a torque when placed in a magnetic field; this is called a *magnetic moment*. The resultant magnetic moment of an atom depends upon three factors—the positive charge of the nucleus spinning on its axis, the negative charge of the electron spinning on its axis, and the effect of the electrons moving in their orbits. The magnetic moment of the spin and orbital motions of the electron far exceeds that of the spinning proton. However, this magnetic moment can be affected by the presence of an adjacent atom. Accordingly, if two hydrogen atoms are combined to form a hydrogen *molecule*, it is found that the electron spins, the proton spins, and the orbital motions of the electrons of each atom oppose each other so that a resultant magnetic moment of zero should be expected. Although this is almost the case, experiment reveals that the relative permeability of hydrogen is not equal to 1 but rather is very slightly less than unity. In other words, the molecular reaction is such that when hydrogen is the medium there is a slight decrease in the magnetic field compared with free space. This behavior occurs because there is a precessional motion of all rotating charges about the field direction, and the effect of this precession is to set up a field opposed to the applied field regardless of the direction of spin or orbital motion. Materials in which this behavior manifests itself are called *diamagnetic* for obvious reasons. Besides hydrogen, other materials possessing this characteristic are silver and copper.

Continuing further with the hydrogen molecule, let us assume next that it is made to lose an electron, thus yielding the hydrogen ion. Clearly, complete neutralization of the spin and orbital electron motions no longer takes place. In fact, when a magnetic field is applied, the ion is so oriented that its net magnetic moment aligns itself with the field, thereby causing a slight increase in flux density. This behavior is described as *paramagnetism* and

is characteristic of such materials as aluminum and platinum. Paramagnetic materials have a relative permeability slightly in excess of unity.

So far we have considered those elements whose magnetic properties differ only very slightly from those of free space. As a matter of fact the vast majority of materials fall within this category. However, there is one class of materials-principally iron and its alloys with nickel, cobalt, and aluminum-for which the relative permeability is very many times greater than that of free space. These materials are called *ferromagnetic* and are of great importance in electrical engineering. We may ask at this point why iron (and its alloys) is so very much more magnetic than other elements. Essentially, the answer is provided by the *domain* theory of magnetism. Like all metals, iron is crystalline in structure with the atoms arranged in a space lattice. However, domains are subcrystalline particles of varying sizes and shapes containing about 10^8 atoms in a volume of approximately cubic centimeters. *The distinguishing feature of the domain is that the magnetic moments of its constituent atoms are all aligned in the same direction* Thus in a ferromagnetic material, not only must there exist a magnetic moment due to a nonneutralized spin of an electron in an inner orbit, but also the resultant spin of all neighboring atoms in the domain must be parallel.



It would seem by the explanation so far that, if iron is composed of completely magnetized domains, then the iron should be in a state of complete magnetization throughout the body of material even without the application of a magnetizing force. Actually, this is not the case, because the domains act independently of each other, and for a specimen of unmagnetized iron these domains are aligned haphazardly in all directions so that the net magnetic moment is zero over the specimen. Figure 1-1 illustrates the situation diagrammatically in a simplified fashion. Because of the crystal

lattice structure of iron the 'easy' direction of domain alignment can take place in any

one of six directions-left, right, up, down, out, or in-depending upon the direction of the applied magnetizing force. Figure 1-1(a) shows the unmagnetized configuration. Figure 1-1(b) depicts the result of applying a force from left to right of such magnitude as to effect alignment of all the domains. When this state is reached the iron is said to be *saturated*-there is no further increase in flux density over that of free space for further increases in magnetizing force.

Large increases in the temperature of a magnetized piece of iron bring about a decrease in its magnetizing capability. The temperature increase enforces the agitation existing between atoms until at a temperature of 750°C the agitation is so severe that it destroys the parallelism existing between the magnetic moments of the neighboring atoms of the domain and thereby causes it to lose its magnetic property. The temperature at which this occurs is called the *curie point*.

Part I. Comprehension Exercises

A. Put “T” for true and “F” for false statements. Justify your answers.

-1. With his atomic theory, Bohr contributed to the understanding of the magnetic behavior of materials.
-2. The atoms of a substance, if placed in a magnetic field, are subject to a torque.
-3. Platinum is a diamagnetic material.
-4. In ferromagnetic materials, the magnetic moments of large groups
-5. In an unmagnetized ferromagnetic material, the domains are aligned in different direction.
-6. The magnetic properties of iron increase with an increase in temperature.

B. Choose a, b, c, or d which best completes each item.

- 1. Permeability of silver is less than unity
 - a. because of its atoms setting up a field against the applied field
 - b. because of its molecules rotating about the applied field
 - c. due to the precessional spin of its positive charges
 - d. due to the orbital motions of its negative charges

2. It is true that
 - a. paramagnetic materials provide a small penetration of the magnetic field
 - b. paramagnetic materials provide a great penetration of the magnetic field
 - c. the resultant magnetic moment of an atom depends on its spinning axis
 - d. the resultant magnetic moment of an atom depends on the nucleus spinning on its axis
3. According to the text,
 - a. two atoms of hydrogen, if combined, pronounce a permeability greater than 1
 - b. two atoms of hydrogen, if combined, give rise to a high magnetic moment
 - c. diamagnetic materials have magnetic properties more than those of free space
 - d. diamagnetic materials have magnetic properties less than those of free space
4. Paramagnetism is based on the fact that the magnetic moment of a paramagnetic material, when placed in a magnetic field,
 - a. results in a decrease in flux density
 - b. lines up with the field
 - c. is equal to 1
 - d. is low compared with free space
5. The magnetic properties of diamagnetic and paramagnetic materials those of free space.

a. are greater than	b. are smaller than
c. differ slightly from	d. differ greatly from
6. The abnormal magnetic properties of iron may be caused by
 - a. the magnetic moment resulting from an inner orbital spin of a nonneutralized electron
 - b. the parallelism of the resultant spin of all neighboring atoms in the domain
 - c. the domains oriented at random with their axes pointing in various directions
 - d. both a and b

C. Answer the following questions orally.

1. What is called a magnetic moment?
2. What does the resultant magnetic moment of an atom depend on?
3. How do adjacent atoms affect the magnetic moment of each other?
4. How does the magnetic behavior of materials differ?
5. Why does platinum have the characteristic of paramagnetism?
6. What forms the domains in a ferromagnetic material?
7. What causes the alignment of the magnetic domains in iron?
8. What is called the curie point?

Part II Language Practice

A. Choose a, b, c, or d which best completes each item.

1. Copper is material, therefore, it exhibits a relative permeability slightly less than unity.
a. a paramagnetic b. a diamagnetic
c. a permeable d. a neutral
2. Iron provides a great penetration of the magnetic field, that is, its is many times greater than that of free space.
a. magnetic flux b. atomic composition
c. relative permeability d. magnetic moment
3. Elements and metals which have slight magnetic properties are called materials.
a. magnetic b. metallic
c. diamagnetic d. paramagnetic
4. Iron and some of its alloys have an appreciable magnetic permeability. These materials are called
a. ferromagnetic b. diamagnetic
c. paramagnetic d. magnetic
5. The state of is reached when all the magnetic domains are aligned in one direction.
a. magnetization b. saturation
c. flux density d. neutralization

B. Fill in the blanks with the appropriate form of the words given.

1. Magnet

- a. Maxwell showed that some of the properties of may be compared to a flow.

- b. Lines of flux are conventionally said to leave a material at the north pole and re-enter at the south pole,
- c. If the field is produced by a solenoid, we will have the same representation of lines of flux, but with the solenoid taking the place of a

2. Permeate

- a. Relativeis a pure number that is the same in all unit systems; the value and dimension of absolutedepend upon the system of units employed.
- b. A is an apparatus used for determining corresponding values of magnetizing force and flux density in a test specimen.

3. Move

- a. When a conductor is through a magnetic field in such a way as to cut the magnetic lines, an emf is generated in the conductor.
- b. A moving - conductor microphone is a microphone the electric output of which results from the of a conductor in a magnetic field.
- c. In a moving - conductor loudspeaker, the conductor is in the form of a coil connected to the source of electric energy.

4. Rotate

- a. The most important parts of a dc motor are the, the stator, and the brushgear .
- b. Aconverter combines both motor and generator action in one armature winding connected to both a commutator and slip rings, and is excited by one magnetic field.
- c. A rotary generator is an alternating-current generator adapted to be by a motor or prime mover.

5. Saturate

- a. A magnetic-core reactor operating in the region of saturation without independent control means is known as reactor.
- b. A sleeve is a flexible tubular product made from cotton and coated with an electrical insulating material.
- c. Saturation induction is sometimes referred to as flux density.

C. Fill in the blanks with the following words.

inductance	element	circuit	way
changing	treated	flux	it
discovered	current	from	

Inductance is a characteristic of magnetic fields, and it was first.....by

Faraday in his renowned experiments of 1831. In a general inductance can be characterized as that property of a circuit by which energy is capable of being stored in a magnetic field. A significant and distinguishing feature of inductance, however, is that makes itself felt in a circuit only when there is a/ancurrent or flux. Thus, although a circuit element may haveby virtue of its geometrical and magnetic properties, its presence in the is not exhibited unless there is a time rate of change of This aspect of inductance is particularly stressed when we consider it the circuit viewpoint. However, for the sake of completeness, inductance is also..... from an energy and a physical view.

D. Put the following sentences in the right order to form a paragraph. Write the corresponding letters in the boxes provided.

- a. Transformers are to be found in such varied applications as radio and television receivers and electrical power distribution circuits.
- b. An understanding of electromagnetism is essential to the study of electrical engineering because it is the key to the operation of a great part of the electrical apparatus found in industry as well as the home.
- c. Similarly, static transformers provide the means for converting energy from one electrical system to another through the medium of a magnetic field.
- d. Other important devices—for example, circuit breakers, automatic switches, relays, and magnetic amplifiers—require the presence of a confined magnetic field for their proper operation.
- e. All electric motors and generators, ranging in size from the fractional horsepower units found in home appliances to the 25,000-hp giants used in some industries, depend upon the electromagnetic field as the coupling device permitting interchange of energy between an electrical system and a mechanical system and vice versa.

1	2	3	4	5
				

Section Two: Further Reading

Magnetism

The main experimental facts underlying magnetism are the following:

The ancient Greeks knew that the mineral *loadstone* or *magnetite* (Fe_3O_4) attracts pieces of iron at some zones on its surface. Magnetite in fact is a natural magnet.

If a piece of magnetite is brought near a bar of hard iron, this too acquires the property of attracting iron filings; it has become an artificial magnet. This process is known as magnetic induction.

When iron bars become magnetized, the quality of attracting pieces of iron is found at two regions at the ends of the bars. These regions are called the poles of the magnet.

If we bring two magnets together, with one magnet fixed and the other free to turn, we see that the first magnet exerts some forces on the second. A magnet produces a magnetic field in the space around it. In a similar way, we have seen that electric charges produce an electric field.

The fact that a magnetic bar or a compass needle comes to rest in a roughly north-south direction when freely suspended near the surface of the earth is evidence that the earth itself acts as a magnet. By convention we give the name north pole to that pole of the magnetic bar or needle which seeks the geographic north, the other pole being known as the south pole.

If we take a given pole of a magnet and place it first at one and then at the other pole of a second magnet, in one case the two poles will attract each other and in the other case they will repel each other. It is found that unlike poles attract whereas like poles repel each other.

In a magnet it is not possible to separate the north pole from the south pole. In fact, if we break a magnet in two we find a south pole at the broken end of the part that had the original north pole, and a north pole at the broken end of the part that had the original south pole.

Of all the metals or elements, only iron, cobalt and nickel and some of their alloys have pronounced magnetic properties. These materials are known as *ferromagnetic* materials. Other elements and metals have slight magnetic properties, and they are called *paramagnetic* materials. There is a third series of materials that have magnetic properties less than those of a vacuum, and these are called *diamagnetic* materials.

Ousted in 1820 showed that a current flowing in an electric circuit exerts forces on

a nearby magnet and so demonstrated that a magnetic field is generated around an electric current. Consequently, if we place an electric circuit in a magnetic field, the circuit is subject to forces.

The fact that a magnetic field can be produced either by a magnet or by an electric current may seem strange. But we must remember that in matter we have microscopic circuits due to the movement of electrons, and these circuits are responsible for the magnetic effects of ferromagnetic materials. However, the causes which underlie the magnetic forces produced by electric circuits are not fully understood (just as there are still problems in our understanding of the forces between electric charges and the nature of the force of gravity), although we know the laws that govern their actions and can therefore use them. We know that atoms consist of a heavy central positive nucleus and a number of electrons, in either circular or elliptical orbits, around the nucleus. Recently there has been added the concept that each electron itself is spinning about an axis through its centre, this motion being known as *electron spin*. Here, it is impossible to offer a complete explanation of this and we must limit ourselves to saying that the fundamental magnetic particles in ferromagnetic materials are the spinning electrons. These electrons occupy definite shells in the atom, and some spin in one direction and some in the other. Their magnetic effects tend to neutralize each other partially but not wholly. The excess of those spinning in one direction over those spinning in the other causes each atom as a whole to act as a small permanent magnet. Moreover, in ferromagnetic materials there is the existence of some kinds of interatomic forces that cause the alignment of all magnetic effects of large groups of atoms to give highly magnetic domains. In an unmagnetized ferromagnetic substance these domains are oriented at random with their magnetic axes pointing in various directions, so that the resultant magnetic effect is zero. The application of an external field lines up the domain axes, thereby giving rise to the magnetic effect of a ferromagnetic material.

In hard iron the domains do not easily return to their previous positions when the external field is removed, while in soft iron this occurs fairly readily. Paramagnetic and diamagnetic materials, on the other hand, are substances in which the arrangement of the spinning electrons does not give appreciable magnetic properties. When the temperature of a ferromagnetic material is raised beyond a certain value (known as the *Curie point*), thermal agitation

destroys the alignment within the domains and the materials lose their ferromagnetic properties. These properties return when the materials are cooled. The Curie point for iron is of the order of 700°C . As in the case of an electric field, a magnetic field at each point may be defined by its field strength. This is represented by the vector H . The direction is that in which a north pole subjected to this field tends to move. Because the magnetic field may be produced by a current, the strength can be defined in terms of current. In order to do this we consider a solenoid, i.e., a coil of wire wound uniformly on a cylindrical former. If the solenoid is long compared with its radius, we can consider that a uniform magnetic field is produced inside the coil, parallel to its axis. If N is the number of turns, l the length of the solenoid and I the current that flows in the coil, we have $H = NI/l$. The magnitude of H is measured in amperes per metre, and the quantity NI is expressed in amperes.

Comprehension Exercises

A. Choose a, b, c, or d which best completes each item.

1. If we break a magnetic bar into two pieces, the two poles at the point of breakage will
 - a. be two north poles
 - b. be a north pole and a south pole
 - c. pronounce greater attraction
 - d. pronounce smaller attraction

2. It is true that
 - a. the circuits in matter produce magnetic forces
 - b. the circuits in an electric current produce magnetic forces
 - c. paramagnetic materials have smaller magnetic properties than diamagnetic materials
 - d. diamagnetic materials have greater magnetic properties than ferromagnetic materials

3. The factors bringing about the magnetic properties of materials are the spinning

a. nuclei	b. atoms
c. neutrons	d. electrons

4. Paragraph ten mainly discusses

a. the magnetic field	b. the electric field
c. the theory of magnetism	d. the theory of gravity

5. When the temperature of cobalt is below the Curie point
 - a. all magnetism disappears
 - b. some magnetism disappears

- c. the metal has appreciable magnetic properties
 - d. the alignment of the magnetic domains is destroyed
6. The vector H representing the field strength of a magnetic field may be expressed as the product of
- a. the number of turns in a coil and the current in amperes which flows through it
 - b. the number of turns in a coil and the current in amperes which flows through it per unit length
 - c. the current flowing through a coil and the length of the coil
 - d. the current flowing through a coil per unit length
7. In a bar magnet, the magnetic domains
- a. neutralize each other
 - b. repel each other
 - c. are at random
 - d. are aligned
8. A magnet and an electric current in a circuit produce a magnetic field by virtue of
- a. the position of the magnetic domains
 - b. the orientation of the atomic nuclei
 - c. the movement of the electrons
 - d. the alignment of the interatomic forces

B. Write the answers to the following questions.

1. How have the poles of a magnetic bar been initially named?
2. How do you describe the process of magnetic induction?
3. How is an electric field compared with a magnetic field?
4. How may a magnetic field be demonstrated?
5. How are paramagnetic materials different from diamagnetic materials?
6. What did Oersted prove in 1820?
7. What is the difference between the soft iron and the hard iron?
8. How do you describe the vector H ?



Section Three: Translation Activities

A. Translate the following passage into Persian.

Magnetostatics

Amongst the oldest and easiest to observe scientific phenomena are those of

magnetism. The subject of bar magnets and magnetic poles is given the name of magnetostatics by analogy with electrostatics. Magnetic phenomena, such as poles and the fields they produce, can be explained in terms of fields due to electric currents. Thinking on the atomic scale, the electrons which circulate round the heavy central positive nucleus constitute a current. So each orbiting electron produces a magnetic field. In general, the orbits of the electrons are disposed in random planes in space and so the net magnetic field is zero. Should a suitable stimulus be applied, the orbits can be aligned so that their magnetic fields are in the same direction. In some materials the orbits, once aligned, stay that way and these are the materials which produce permanent magnets. In other materials the orbits return to their random dispositions once the stimulus is removed—these are the materials used as electromagnets. For a given stimulus, they produce a greater field than the materials used for permanent magnets. It should be mentioned at this stage that ability to produce a high field is not the only factor to be considered when deciding upon a material to be used for an electromagnet; there are problems of energy loss to be considered if the state of magnetization is to be changed frequently.

B. Find the Persian equivalents of the following terms and expressions and write them in the spaces provided.

- | | |
|---------------------|-------|
| 1. agitate | |
| 2. alloy | |
| 3. apparatus | |
| 4. attract | |
| 5. brush gear | |
| 6. circuit | |
| 7. commutator | |
| 8. compass needle | |
| 9. compose | |
| 10. coupling device | |
| 11. crystalline | |
| 12. depict | |
| 13. diamagnetic | |
| 14. dispose | |
| 15. domain | |
| 16. electrostatic | |
| 17. exert on | |
| 18. ferromagnetic | |

19. gravity	
20. lattice	
21. magnetic amplifier	
22. magnetic induction	
23. magnetic moment	
24. manifest	
25. mineral	
26. neutralize	
27. nucleus	
28. orbit	
29. permeability	
30. precession	
31.random	
32. repel	
33. saturate	
34. sleeve	
35. specimen	
36.spin	
37. stator	
38. stimulus	
39.suspend	
40.torque	
41. transformer	
42.vacuum	

Unit 2

Section One: Reading Comprehension

Power Stations

There are five sources of energy which together account for nearly all the world's electricity. They are coal, oil, natural gas, hydroelectric power and nuclear energy. Coal, oil and nuclear plants use the steam cycle to turn heat into electrical energy, in the following way. The steam power station uses very pure water in a closed cycle. First it is heated in the boilers to produce steam at high pressure and high temperature, typically 150 atmospheres and 550°C in a modern station. This high-pressure steam drives the turbines which in turn drive the electric generators, to which they are directly coupled. The maximum amount of energy will be transferred from the steam to the turbines only if the latter are allowed to exhaust at a very low pressure, ideally a vacuum. This can be achieved by condensing the outlet steam into water. The water is then pumped back into the boilers and the cycle begins again. At the condensing stage a large quantity of heat has to be extracted from the system. This heat is removed in the condenser which is a form of heat exchanger. A much larger quantity of cold impure water enters one side of the condenser and leaves as warm water, having extracted enough heat from the exhaust steam to condense it back into water. At no point must the two water systems mix. At a coastal site the warmed impure water is simply returned to the sea at a point a short distance away. A 2 GW station needs about 60 tons of sea water each second. This is no problem on the coast, but inland very few sites could supply so much water all the year round. The alternative is to recirculate the impure water. Cooling towers are used to cool the impure water so that it can be returned to the condensers, the same water being cycled continuously. A cooling tower is the familiar concrete structure like a very broad chimney and acts in a similar way, in that it induces a natural draught. A large volume of air is drawn in round the base and leaves through the open top. The warm, impure water is sprayed into the interior of the tower from a large number of fine jets, and as it falls it is cooled by the rising air, finally being collected in a pond under the tower. The cooling tower is really a second heat exchanger where the heat in the impure water is passed to the atmospheric air; but unlike the first heat exchanger, the two fluids are

allowed to come into contact and as a consequence some of the water is lost by evaporation.

The cooling towers are never able to reduce the impure water temperature right down to the ambient air temperature, so that the efficiency of the condenser and hence the efficiency of the whole station is reduced slightly compared with a coastal site. The construction of the cooling towers also increases the capital cost of building the power station. The need for cooling water is an important factor in the choice of sites for coal, oil and nuclear plants. A site which is suitable for a power station using one type of fuel is not necessarily suitable for a station using another fuel.

Coal-Fired Power Stations

Early coal-burning stations were built near the load they supplied. A station of 2 GW output, consumes about 5 million tons of coal in a year. In Britain where most power station coal is carried by rail, this represents an average of about 13 trains a day each carrying 1000 tons. This means that large coal-fired stations need a rail link unless the station is built right at the pit head.

Oil-Fired Power Stations

Power station oil can be divided into crude oil which is oil as it comes from the well, and residual oil which remains when the more valuable fractions have been extracted in the oil refinery. The cost of moving oil by pipeline is less than that of moving coal by rail, but even so stations burning crude oil are often sited near deep-water berths suitable for unloading medium-sized tankers. Stations burning residual oil need to be sited near to the refinery which supplies them. This is because residual oil is very viscous and can only be moved through pipelines economically if it is kept warm.

Nuclear Power Stations

In contrast to coal and oil the cost of transporting nuclear fuel is negligible because of the very small amount used. A 1 GW station needs about $4\frac{1}{2}$ tons of uranium each week. This compares very favourably with the 50,000 tons of fuel which would be burnt each week in a comparable coal-fired power station. Present nuclear stations use rather more cooling water than comparable coal-fired or oil-fired plants due to their lower efficiency. All nuclear stations in Britain, with one exception, are situated on the coast and use sea water for cooling.

Hydroelectric Power Stations

Hydroelectric power stations must be sited where the head of water is available, and as this is often in mountainous areas, they may need long transmission lines to carry the power to the nearest load center or link up with the grid. All hydroelectric schemes depend on two fundamental factors: a flow of water and a difference in level or head. The necessary head may be obtained between a lake and a nearby valley, or by building a small dam in a river which diverts the flow through the power station, or by building a high dam across a valley to create an artificial lake.

Part I. Comprehension Exercises

A. Put "T" for true and "F" for false statements. Justify your answers.

- 1. Gas and nuclear plants use the steam cycle to turn heat into electricity.
- 2. Condensers remove the heat from the outlet steam.
- 3. The steam power station uses pure water in an open cycle.
- 4. Steam pressure affects the generators directly.
- 5. Having cooled off the exhaust steam, the warmed impure water may be recirculated.
- 6. Natural air is forced through the cooling tower.
- 7. Large coal-fired stations situated far from the pit head need a rail link.
- 8. Oil-fired power stations consume certain constituents of crude oil.
- 9. Nuclear power stations use less cooling water than comparable coal-fired or oil-fired plants.
- 10. Hydroelectric power stations have to be built where there is enough water pressure.

B. Choose a, b, c, or d which best completes each item.

1. In steam power stations, the turbine efficiency will increase if
 - a. the steam pressure is kept constant
 - b. the outlet steam is condensed into water
 - c. the steam temperature is not varied
 - d. the outlet water is pumped back into the boilers
2. The steam power station uses pure water
 - a. to produce the steam required to drive the turbines
 - b. to produce the steam required to activate the generators

- c. to create the vacuum space necessary for the system
- d. to create the pressure and temperature needed
- 3. The heat of the steam is removed by the condenser.
 - a. the recirculation of cold pure water in
 - b. the flow of natural air in one side of
 - c. the recirculation of the steam in
 - d. the flow of cold water through one side of
- 4. Prior to recirculation, impure water must be cooled
 - a. in broad concrete structures
 - b. in broad metal chimneys
 - c. at the bottom of the tower
 - d. at the top of the tower
- 5. The cooling factor in a cooling tower is the tower.
 - a. the pond under
 - b. the interior of
 - c. the water inside
 - d. the air passing through
- 6. Systems recirculating impure water, compared with those on the coast,
 - a. decrease the efficiency of the station
 - b. increase the capital cost of building the station
 - c. reduce the impure water temperature to the required level
 - d. both a and b
- 7. The first paragraph mainly discusses
 - a. the structure of a condenser compared with that of a cooling tower
 - b. the mechanism of the steam power station
 - c. the main sources of energy which account for electricity
 - d. the cooling water as a deciding factor in the choice of sites for coal, oil, and nuclear plants

C. Answer the following questions orally.

1. What are the five sources of energy used for the generation of electrical energy?
2. What are the two water systems used in the condenser?
3. What is the water resulted from steam condensation used for?
4. How much sea water does a 2 GW station need each second?
5. How is the mechanism of a cooling tower similar to that of a chimney?
6. How do you describe the mechanism of a cooling tower?
7. What are the two heat exchangers used in the system?
8. How much coal does a 2 GW station consume every year?
9. Why should stations burning residual oil be sited near to the refinery which supplies them?

10. Why is the cost of transporting nuclear fuel negligible compared with coal and oil?

Part II. Language Practice

A. Choose a, b, c, or d which best completes each item.

1. The energy of water may be converted to work by hydraulic
 - a. turbines
 - b. generator
 - c. boilers
 - d. towers
2. Gas oil must be and then used.
 - a. isolated
 - b. heated
 - c. refined
 - d. vapourized
3. In the condenser, the outlet steam is and recirculated.
 - a. exchanged
 - b. condensed
 - c. depressurized
 - d. purified
4. Cooling towers cause water to be
 - a. condensed
 - b. exhausted
 - c. evaporated
 - d. recycled
5. Air pump suction must be applied to the lowest pressure point or points within a condenser which are normally at the inlet tube plate where rate and hence steam side pressure drop are greatest.
 - a. the condensation
 - b. the temperature
 - c. the cooling
 - d. the evaporation

B. Fill in the blanks with the appropriate form of the words

given.

1. Exchange

- The atomic movements of materials are said to be held in parallel or antiparallel by exchange forces, thought to be due to the sharing or of electrons between neighbouring atoms in the crystal structure of the material.
- Coupling forces, similar to the forces of the atom, exist between the molecules of a compound.
- Cooling towers and condensers are two kinds of heat

2. Circulate

- a. A register retains data by inserting it into a delaying means and regenerating and reinserting the data into the register.

- b. A constant flow of electrolyte through a cell to facilitate the maintenance of uniform conditions of electrolysis is known as of electrolyte.
- c. A.....magnetic wave is a traverse magnetic wave for which the lines of magnetic force form concentric circles.

3. Couple

- a. Water heated in the boilers of the steam power station produces steam at high pressure which drives the turbines to generators.
- b. Typical oscillators are in practice amplifiers in which power is fed into the grid circuit from the plate circuit by means of either electrostatic or electromagnetic between these circuits.

4. Condense

- a. Condensed-mercury temperature is the temperature measured on the outside of the tube envelope in the region where the mercury is in a glass tube or at a designated point on a metal tube.
- b. Steam can be into water.
- c. A is a form of heat exchanger.

5. Drive

- a. A is an electronic circuit that supplies input to another electronic circuit.
- b. Grid power is the average of the product of the instantaneous values of the alternating components of the grid current and the grid voltage over a complete cycle .
- b. A system consisting of one or several electric motors and of the entire electric control equipment designed to govern the performance of these motors is called the electric

C. Fill in the blanks with the following words.

generally	quality	same	oil
produce	natural	feed	used
situated	heat		

Any steam power station burning coal or could fairly easily be converted to burn gas. Such stations must, of course, be near a large gas main. However, it is felt that natural gas is too high a/an fuel and too valuable as an industrialstock and home heating fuel, to be..... in power stations. The point is that gas burnt to produce electricity

which might then be used for home heating, would produce at about 33 percent efficiency, whereas the gas burnt in a domestic boiler would heat at up to 80 percent efficiency.

**D. Put the following sentences in the right order to form a paragraph.
Write the corresponding letters in the boxes provided.**

- a. The flame temperature is clearly much higher than the steam temperature, but the thermodynamic efficiency of a conventional station depends on the steam temperature not the flame temperature.
- b. Firstly, work associated with existing coal- and oil-burning power stations where efforts are being made to utilize the inherent thermodynamic efficiency of the very high flame temperatures of burning oil or pulverised coal.
- c. Generators have been constructed to convert some of the energy in the flame, which is a moving ionised gas, directly into electricity.
- d. Experiments and design studies are being carried out to develop new ways of generating electricity.
- e. Secondly, work is being done to try to convert solar energy into electricity.
- f. These fall broadly into three groups.
- g. Thirdly, we have what is sometimes called the nuclear alternative,
- h. These are known as magnetohydrodynamic generators.

1	2	3	4	5	6	7	8

Section Two: Further Reading

Electric Power Stations

Characteristics Influencing Generation and Transmission

There are four main characteristics of electricity supply which, however obvious, have a profound effect on the manner in which it is engineered. They are as follows:

(a) Electricity, unlike gas and water, cannot be stored and the supplier has little control over the load at any time. The control engineers endeavor to keep the output from the generators equal to the connected load at the specified voltage and frequency.

(b) There is a continuous increase in the demand for power. Although in industrialized countries the rate of increase has declined in recent years, even the modest rate entails massive additions to the existing systems. A large and continuous process of adding to the system thus exists. Networks are evolved over the years rather than planned in a clear-cut manner and then left untouched.

(c) The distribution and nature of the *fuel* available. This aspect is of great interest as coal is mined in areas not necessarily the main load centres; hydroelectric power is usually remote from the large load centres. The problem of station siting and transmission distances is an involved exercise in economics. The greater use of nuclear energy will tend to modify the existing pattern of supply.

(d) In recent years environmental considerations have assumed major importance and influence the siting, construction cost, and operation of generating plants. Planning is also affected because of delays in making a start to projects because of legal proceedings, etc. Of particular importance at the present time is the question of the environmental impact of nuclear plants, especially the proposed fast breeder reactor.

Energy Conversion Employing Steam

The combustion of coal or oil in boilers produces steam at high temperatures and pressures which is passed to steam turbines. Oil has economic advantages when it can be pumped from the refinery through pipelines direct to the boilers of the generating station. The use of energy resulting from nuclear fission is being progressively extended in electricity generation; here also the basic energy is used to produce steam for turbines. The axial-flow type of turbine is in common use with several cylinders on the same shaft.

The steam power station operates on the Rankine cycle, modified to include muted superheating, feed-waterheating, and steam reheating. Increased thermal efficiency results from the use of steam at the highest possible pressure and temperature. Also, for turbines to be economically constructed 500 MW and over are now being used. With steam turbines of 100 MW capacity and over the efficiency is increased by reheating the steam after it has been partially expanded, by an external heater. The reheated steam is then

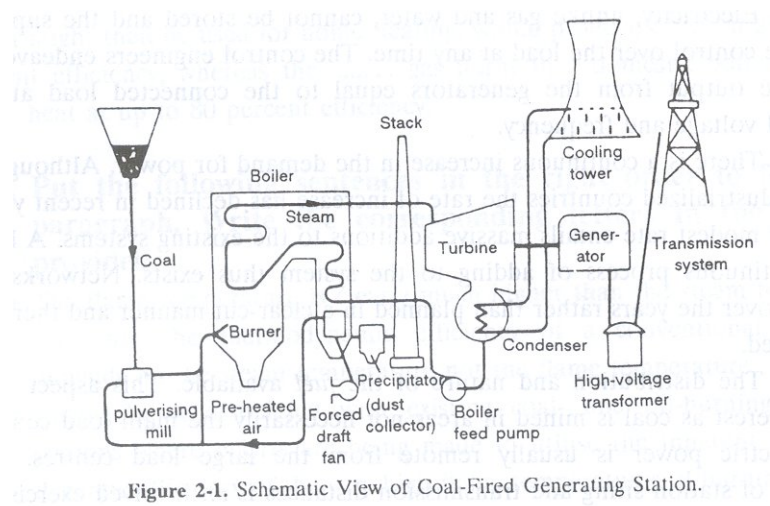
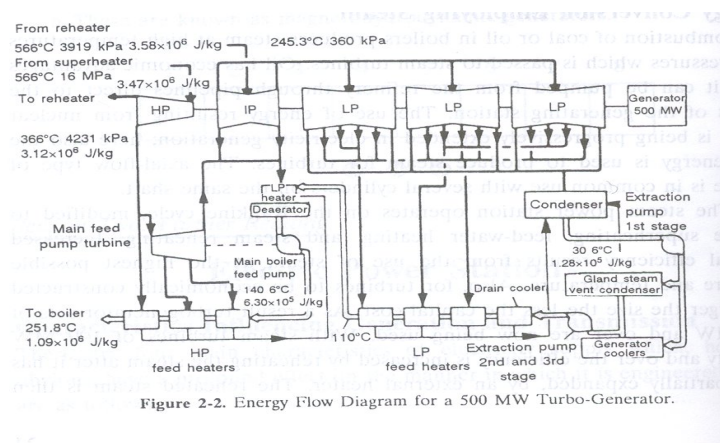


Figure returned to the turbine where it is expanded through the final stages of blading. A schematic diagram of a coal-fired station is shown in Figure 2-1. In Figure 2-2, the flow of energy in a modern steam station is shown. Despite continual advances in the design of boilers and in the development of improved materials, the nature of the steam cycle is such that efficiencies are comparatively low and vast quantities of heat are lost in the condensate. However, the great advances in design and materials in the last few years have increased the thermal efficiencies of coal stations to about 40 percent.



In coal-fired stations, coal is conveyed to a mill and crushed into fine powder i.e., pulverised. The pulverized fuel is blown into the boiler where it mixes with a supply of air for combustion. The exhaust from the LP turbine is cooled to form condensate by the passage through the condenser of large quantities of sea- or river-water. Where this is not possible cooling towers are used.

Fluidized-Bed Boilers. For typical coals, combustion gases contain 0.2-0.3 percent sulphur dioxide by volume. If the gas flow-rate through the granular bed of a great-type boiler is increased the gravity pull is balanced by the upward gas force and the fuel-bed takes on the character of a fluid. In a travelling grate this increases the heat output and temperature. The ash formed conglomerates and sinks into the grate and is carried to the ash pit. The bed is limited to the ash-sintering temperature of 1050-1200°C. Secondary combustion occurs above the bed where CO burns to CO_2 and H_2S to SO_2 . This type of boiler is undergoing extensive development and is attractive because of the lower pollutant level and better efficiency.

Energy Conversion Using Water

Perhaps the oldest form of energy conversion is by the use of water power. In the hydroelectric station the energy is obtained free of cost. This attractive feature has always been somewhat offset by the very high capital cost of construction, especially of the civil engineering works. Today, however, the capital cost per kilowatt of hydroelectric stations is becoming comparable with that of steam stations. Unfortunately, the geographical conditions necessary for hydro-generation are not commonly found. In most highly developed countries hydroelectric resources are used to the utmost.

An alternative to the conventional use of water energy, pumped storage, enables water to be used in situations which would not be amenable to conventional schemes. The utilization of the energy in tidal flows in channels has long been the subject of speculation. The technical and economic difficulties are very great and few locations exist where such a scheme would be feasible. An installation using tidal flow has been constructed on the La Rance estuary in northern France where the tidal height range is 9.2 m (30 ft) and the tidal flow is estimated at 18,000 m³/s.

Before discussing the types of turbine used, a brief comment on the general modes of operation of hydroelectric stations will be given. The vertical difference between the upper reservoir and the level of the turbines is known head. The water falling through this head gains kinetic energy which it

then imparts to the turbine blades. There are three main types of installation as follows:

- (a) *High. Head or Stored*-the storage area or reservoir normally fills in over 400 h;
- (b) *Medium Head or Pondage*-storage fills in 200-400 h;
- (c) *Run of River*-storage fills in less than 2 h and has 3-15 m head.

A schematic diagram for type (c) is shown in Figure 2-3.

Associated with these various heights or heads of water level above the turbines are particular types of turbine. These are:

- (a) *Pelton*. This is used for heads of 184-1840 m (600-6000 ft) and consists of a bucket wheel rotor with adjustable How nozzles.
- (b) *Francis*. Used for heads of 37-490 m (120-1600 ft) and is of the mixed flow type.
- (c) *Kaplan*. Used for run of river and pondage stations with heads of up to 61 m (200 ft). This type has an axial-flow rotor with variable-pitch blades.

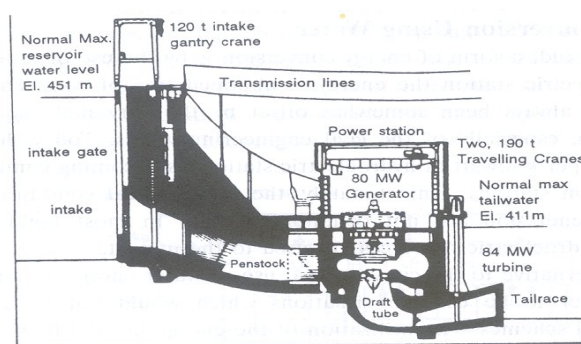


Figure 2-3. Hydroelectric Scheme —Kainji, Nigeria. Section through the intake dam and power house. The scheme comprises an initial four 80 MW Kaplan turbine sets with the later installation of eight more sets. Running speed 115.4 rev/min. This is a large-flow scheme with penstocks 9 m in diameter. (*Permission of Engineering.*)

Typical efficiency curves for each type of turbine are shown in Figure 2-4. As the efficiency depends upon the head of water which is continually

fluctuating, often water consumption in cubic meters per kilowatt-hour is used and is related to the head of water. Hydroelectric plant has the ability to start UP

quickly and the advantage that no losses are incurred when at a standstill. It has great advantages, therefore, for generation to meet peak loads at minimum cost, working in conjunction with thermal station. By using remote control of the hydro sets, the time from the instruction to start up to the actual connexion to the power network can be as short as 2 min

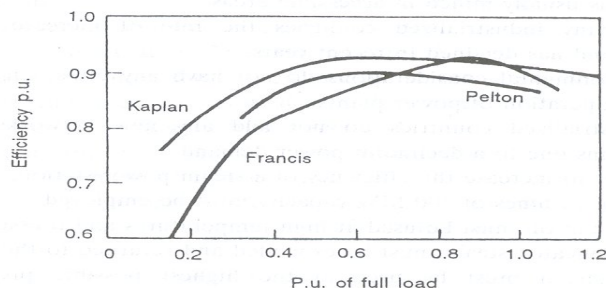


Figure 2-4. Typical Efficiency Curves of Hydraulic Turbines.

Gas Turbines

The use of the gas turbine as a prime mover has certain advantages over steam plant, although with normal running it is less economical to operate. The main advantage lies in the ability to start and take up load quickly. Hence the gas turbine is coming into use as a method for dealing with the peaks of the system load. A further use for this type of machine is as a synchronous compensator to assist with maintaining voltage levels. Even on economic grounds it is probably advantageous to meet peak loads by starting up gas turbines from cold in the order of 2 min rather than running spare steam plant continuously.

Comprehension Exercises

Choose a, b, c, or d which best completes each item.

1. We may deduce from the text that
 - a. since electricity cannot be stored, enough electricity must be generated at all times to meet the variations in demand

- b. since the supplier does not have control over the load, variations in demand have to be limited to certain degrees
 - c. gas, water, and electricity can be stored to satisfy the unexpected increases in demand
 - d. gas, water, and power systems are not amenable to ordinary energy requirements
2. It is true that
- a. coal is usually mined in accessible areas
 - b. in many industrialized countries the rate of increase in power demand has declined in recent years
 - c. environmental considerations do not have any effect on the siting and operation of power plants
 - d. industrialized countries do not add any new networks to their systems due to a decline in power demand
3. In order to increase the efficiency of a steam power station,
- a. steam turbines of 100 MW capacity must be employed
 - b. coal and oil must be used at high temperatures and pressures
 - c. the reheated steam must be expanded and returned to the turbine
 - d. the steam must be used at the highest possible pressure and temperature
4. It is true that
- a. the advances in the design of boilers have not affected the efficiency of coal stations
 - b. coal stations have low efficiencies because of the heat lost in the steam cycle
 - c. systems operating on the steam cycle have high efficiencies
 - d. the MW capacity of all the steam turbines used today is over 500
5. In fluidized-bed boilers
- a. the upward gas force causes the fuel-bed to take the character of a fluid
 - b. the fluid characteristic of the fuel-bed increases the heat output
 - c. CO burns to CO₂ and H₂ S to SO₂.
 - d. all of the above
6. The use of energy in tidal flows
- a. has greatly replaced the conventional use of water energy
 - b. has enabled man to make use of water energy wherever he likes
 - c. may be an alternative to the conventional use of water energy
 - d. is a common way of using water energy without any difficulties

7. The last paragraph mainly discusses
- how the gas turbine deals with the peaks of the system load
 - how the gas turbine is used as a synchronous machine
 - the advantages of the gas turbine
 - the mechanism of the gas turbine

B Write the answers to the following questions.

- 1 What are the four main characteristics of electricity supply?
2. What are coal and nuclear energy used for?
3. What **is** the function of an external heater?
4. What process does the coal go through in coal-fired stations?
5. When are cooling towers used in coal-fired systems?
6. Why are fluidized-bed boilers called so?
7. What are the advantages of fluidized-bed boilers?
8. What is the most prominent feature of hydroelectric power stations?
9. What are the initial requirements for hydro-generation?
10. How does water obtain the energy required to impart to the turbine blades?
11. What are the three types of installation?
12. What is the head of water?
13. How is Kaplan turbine different from Francis and Pelton turbines?



Section Three: Translation Activities

A. Translate the following passage into Persian.

Magnetohydrodynamic (MHD) Generation

In conventional power generation, fuel such as oil or coal is burned. The burning fuel heats boilers to produce steam. The steam is used to drive turbo-alternators. The MHD process generates electricity without requiring a boiler or a turbine.

MHD generation works on the principle that when a conductor cuts a Magnetic field, a current flows through the conductor. In MHD generation the conductor is an ionized gas. Small amounts of metal are added to the gas to improve its conductivity. This is called seeding the gas. The seeded gas is then

pumped at a high temperature and pressure through a strong magnetic field. The electrons in the gas are collected at an electrode. This movement of electrons constitutes a current flow.

Two methods of MHD generation can be used: the open-cycle and the closed-cycle. In the open-cycle method the hot gas is discharged. In the closed-cycle method it is recirculated.

The open-cycle method uses gas from burning coal or oil. The gas is seeded and then passed through a magnetic field to generate current. The seeding elements are recovered and the gas can then be used to drive a turbine before being allowed to escape.

The closed-cycle method uses an inert gas, such as helium, which is heated indirectly. The gas is circulated continually through the MHD generator.

MHD generation is still in its early stages but already an efficiency rate of 60% has been reached. This compares with a maximum of 40% from conventional power stations.

B. Find the Persian equivalents of the following terms and Expressions and write them in the spaces provided.

- | | |
|--------------------|-------|
| 1. ambient | |
| 2. breeder reactor | |
| 3. chimney | |
| 4. concrete | |
| 5. condense | |
| 6. conglomerate | |
| 7. convey | |
| 8. crude oil | |
| 9. decline | |
| 10. diven | |
| 11. efficiency | |
| 12. email | |
| 13. exchange | |
| 14. exhaust | |
| 15. extract | |
| 16. fluctuate | |
| 17. grate type | |

18. incur	
19. inland	
20. magnetohydrodynamic generator	
21. outlet	
22. pit head	
23. prominent	
24. pulverize	
25. seeded gas	
26. speculate	

Unit 3

Section One: Reading Comprehension

Electrical Insulation

Insulation is required to keep electrical conductors separated from each other and from other nearby objects. Ideally, insulation should be totally nonconducting, for then currents are totally restricted to the intended conductors. However, insulation does conduct some current and so must be regarded as a material of very high resistivity. In many applications, the current flow due to conduction through the insulation is so small that it may be entirely neglected. In some instances the conduction currents, measured by very sensitive instruments, serve as a test to determine the suitability of the insulation for use in service.

Although insulating materials are very stable under ordinary circumstances, they may change radically in characteristics under extreme conditions of voltage stress or temperature or under the action of certain chemicals. Such changes may, in local regions, result in the insulating material becoming highly conductive. Unwanted current flow brings about intense heating and the rapid destruction of the insulating material. These insulation failures account for a high percentage of the equipment troubles on electric-power systems. The selection of proper materials, the choice of proper shapes and dimensions, and the control of destructive agencies are some of the problems of the insulation-system designer.

Many different materials are used as insulation on electric-power systems. The choice of material is dictated by the requirements of the particular application and by cost. In residences, the conductors used in branch circuits and in the cords to appliances may be insulated with rubber or plastics of several different kinds. Such materials can withstand necessary bending, are relatively stable in characteristics, and are inexpensive. They are subjected to relatively low electrical stress.

High-voltage cables are subjected to extreme voltage stress; in some cases several hundred kilovolts are impressed across a few centimeters of insulation. They must be manufactured in long sections, and must be sufficiently flexible as to permit pulling into ducts of small cross section. The insulation may be oil-impregnated paper, varnished cambric, or synthetic materials such as polyethylene.

The coils of generators and motors may be insulated with tapes of various kinds. Some of these are made of thin sheets of mica held together by a binder, and other are of fiber glass impregnated with insulating varnish. This insulation must be capable of withstanding quite high operating temperature-mechanical

forces, and vibration.

The insulation on power-transformer windings is commonly paper tape and pressboard operated under oil. The oil saturates the paper, greatly increasing its insulation strength, and, by circulating through ducts, serves as an agent for carrying away the heat generated due to I^2R losses and core losses in the transformer. The transformer insulation is subjected to high electric stress and to large mechanical forces. The shape and arrangement of conducting metal parts is of particular concern in transformer design.

Overhead lines are supported on porcelain insulators. Between the supports air serves as insulation. Porcelain is chosen because of its resistance to deterioration when exposed to the weather, its high dielectric strength, and its ability to wash clean in rain.

Part I Comprehension Exercises

A. Put "T" for true and "F" for false statements. Justify your answers.

- 1. The higher the insulation the less the loss of power.
- 2. In order to avoid insulation failures, very expensive materials are used in power systems.
- 3. Insulation failures do not affect the electric equipment.
- 4. Voltage and temperature variations may bring about insulation failures.
- 5. Rubber and plastic insulating materials are preferred to other kinds because of their cost.
- 6. Polyethylene and mica have different applications in electrical-power systems.

B. Choose a, b, c, or d which best completes each item.

- 1. The first paragraph mainly discusses
 - a. electrical conductors b. nonconducting materials
 - c. the purpose of insulation d. the application of insulators
- 2. As we understand from the text,
 - a. perfect insulation is not possible
 - b. stable insulators are not available

- c. chemicals do not affect good insulators
- d. insulators may never change to temporary conductors
- 3. The second paragraph mainly discusses
 - a. the problems caused by the insulation-system designer
 - b. the factors resulting in insulation failures
 - c. the characteristics of insulating materials
 - d. the rapid destruction of insulating materials
- 4. Tapes of insulating fiber glass are commonly used to insulate
 - a. ordinary conductors
 - b. the windings of power transformer
 - c. high-voltage cables
 - d. the coils of generators and motors
- 5. Insulating tapes
 - a. cannot withstand high electrical stress
 - b. can withstand high temperatures
 - c. are used to insulate ducts of small cross section
 - d. are used to stop deterioration caused by the weather

C. Answer the following questions orally.

1. What is the purpose of insulation?
2. What is one way of deciding on the suitability of insulators used for different purposes?
3. What does the choice of an insulating material depend on?
4. What is the application of porcelain insulators?
5. What are insulating tapes used for?

Part II. Language Practice

A. Choose a, b, c, or d which best completes each item.

1. In order to keep electrical conductors separated from each other, materials must be used.
 - a. capacitive
 - b. resistive
 - c. insulating
 - d. conducting
2. The of metals increases with increase of temperature.
 - a. conductivity
 - b. resistivity
 - c. solubility
 - d. durability
3. Voltage stress may affect of insulating materials.
 - a. the sensitivity
 - b. the suitability
 - c. the stability
 - d. the conductivity

4. Certain insulating materials are impregnated with oil; that is, they are oil.

- | | |
|-------------------|-----------------|
| a. saturated with | b. covered with |
| c. deprived of | d. made of |

5. Porcelain has high resistance to deterioration; in other words, it does not..... quickly.

- | | |
|-------------|---------------|
| a. deflect | b. degenerate |
| c. decrease | d. decline |

B. Fill in the blanks with the appropriate form of the words given.

1. Arrange

- Modern types of air blast or oil circuit breakers have frequently been fitted with trip-free mechanisms in which the tripping instructions are to override the closing instructions.
- Different insulation have different characteristics.
- On important lines of high lightning incidence, it is accepted practice to seek to prevent direct strokes to phase conductors by one or more shielding wires above the phase conductors to intercept lightning strokes and conduct them to the ground.

2. Operate

- The design of a power system should be such that, when breakdowns are inevitable, they are confined to locations where they cause minimum damage and the least disturbance to
- The whole of the electrical and mechanical quantities that characterize the work of a machine at a given time is known as conditions.
- A switch can be by a lever or other operating means.

3. Subject

- If the insulation were to the normal operating voltage which varies within quite narrow limits, there would be no problem.
- It has been known for many years that an insulator surface to high voltage (dc or pulsed) in vacuum can acquire a large positive Charge.
- The parameters which determine the lightning performance of a transmission line are to large variations according to frequency distribution laws which can only be determined by field observations.

4. Measure

- a. The gas pressure can be by means of a standard pressure gauge.
- b. Most branches of science and technology rely on electrical for the control of processes and machines as well as for information.
- c. The role played by electricalinstruments is vital to all modern laboratories and factories.
- d. The current is the value read on the microammeter during a direct high-voltage test of insulation.

5. Insulate

- a. For the economic transmission of power over considerable distances the voltage must be high, although with higher voltages the cost rises.
- b. When any object is said to be , it is understood to be in suitable manner for the conditions to which it is subjected.
- c. An insulated joint is used to adjacent pieces of conduits, pipes, rods, or bars.
- d. The solid generally are of the form of annular discs and truncated cones.

C. Fill in the blanks with the following words.

manufactured	moistened	causes	due to
shattered	cracked	normal	perfect
traceable	current		

Modern porcelain insulator are designed and in such a fashion that in themselves they are almost in operation. Flashover of line insulators is almost always to the breakdown of the air around them overvoltage from lightning or other Insulators whose surfaces are contaminated and then by light rain or fog may flash over even under-operating-voltage conditions. If an insulator is or porous and permits lighting or power- frequency, to pass through the body of the insulator, it may be with the resultant dropping of the line.

D. Put the following sentences in the right order to form a paragraph. Write the corresponding letters in the boxes provided.

- a. Under voltag stress, an increase in temperature causes an increase of the conductance of the insulating material.

- b. The increased temperature causes an increase in the conductance of that local area, and more current flows, thereby increasing the temperature even more.
- c This is sometimes called *thermal breakdown*.
- d. Of course, the increase of temperature will cause increased conduction of heat away from the region and a stable condition may result.
- e. Conduction current flow implies a release of energy in the insulation.
- f. However, if the voltage is further increased, the temperature may continue to rise and the current may continue to increase until chemical changes destroy the insulation and puncture results.
- g. Under electrical stress near the puncture value, a local region may increase in temperature because heat is released faster than it is carried away.



Section Two: Further Reading

Insulation Behavior

When insulation is placed between two metallic conductors A and B connected to a voltage source (Figure 3-1), several phenomena associated with the insulation may be identified. The insulation is a dielectric, influences the capacitance between the plates, a current of low magnitude flows through the body of the insulation, a leakage current flows over the surface of the insulation, and if the voltage is great enough sudden changes in the body of the insulation may make it highly conductive.

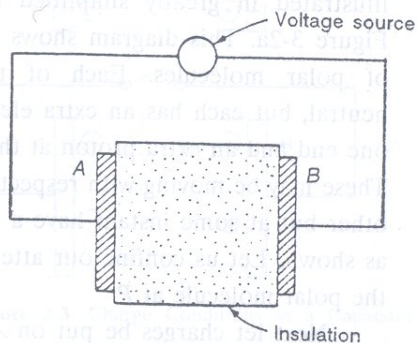


Figure 3-1
Stress Between Two Metal Electrodes

Capacitance and Dielectric Hysteresis

As is well known, the presence of a dielectric between two conducting plates increases the measured capacitance between these plates. The behavior of the electrons and protons comprising the dielectric accounts for this capacitance increase. This phenomenon is worth investigating, for it explains some other characteristics of insulation of interest.

All matter is made up of protons, electrons, and neutrons. In the normal state, these particles are grouped as atoms or molecules in which the number of electrons and the number of protons are equal, therefore, each group is electrically neutral. However, each of these particles experiences a force due to its interaction with any charges placed on nearby plates. We say that the charged particles respond to the electric field set up between the plates.

In a perfect insulator, the electrons protons are held together in the atoms molecules and are not free to drift from one to another. However, in the presence of an electric field, they may move very slight distances, electrons toward the anode, the protons toward the cathode. This situation is illustrated in gross simplified form in Figure 3-2a. This diagram shows a group of polar molecules. Each of them is neutral, but each has an extra electron at one end and an extra proton at the other. These molecules are moving with respect to each other but at any instant have a position as shown. Let us concentrate our attention to the polar molecule at *P*.

Next let charges be put on *A* and *B*. In connection to a voltage source as shown in Figure 3-2b. The molecule *P* rotates, taking up a new position with

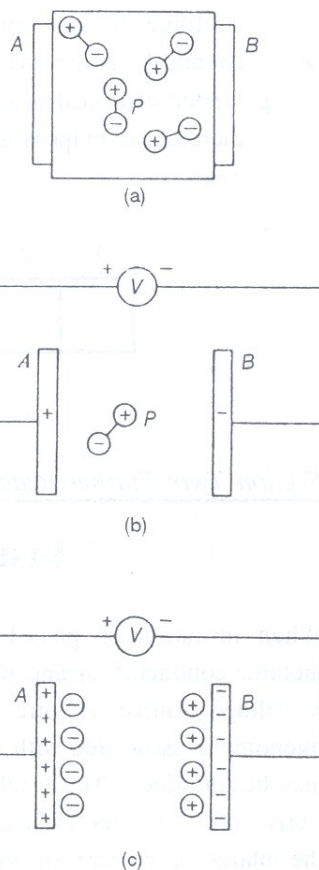


Figure 3-2. The Effect of Electrical Stress of Insulation. (a) Some polar molecules under normal conditions. (b) The polar molecule *P* changes position under stress. (c) The charge layers adjacent to the electrodes.

the electron displaced toward *A* and the proton toward *B*, under influence of the electric field.

The effect of the many, many electron-proton pairs in the dielectric volume,

moving as shown by the example *P*, is to produce an effect shown in Figure 3-2c. Adjacent to the anode *A* there are an excess number of electrons in the dielectric; near the cathode *B* are an excess number of protons in the dielectric. These charges partially neutralize the effect of the charge originally placed on the plate and additional charges move from the voltage source to *A* and *B* as the charges in the dielectric take on their new positions. Hence, for the same voltage between the plates, the charges that have moved from the source to the plates have been increased as a result of the presence of the dielectric. The capacitance between the plates is greater, therefore, than it would be without the dielectric.

Whenever a system of particles is moved from one position to another (such as system *P* from the position shown in Figure 3-2a to that of Figure 3-2b), there are forces which restrict the motion, and time is required to make the change. Such is the case in dielectrics. In some materials the change is made in a fraction of a microsecond; in others it may take several hours. During the period of change, the capacitance appears to increase and current flows in the external circuit. It is sometimes stated that charge is 'soaking' into the dielectric. The phenomenon is known as dielectric absorption.

When the voltage source is disconnected from plates *A* and *B* and the voltage between them is made zero by a short-circuiting connection (Figure 3-3a), the displaced particles in the dielectric tend to go back to their normal state. However, if it took a long time to get them oriented, it will take a long time to get them back into the normal state. Hence the condition shown in Figure 3-3a may persist for some time, a charge remaining on each plate equal in effect to the charge remaining in the dielectric adjacent to the plates

Next, suppose that the short

Figure 3-3. Charge Conditions in a Capacitor That has Been Energized, (a) Immediately after being short-circuited, (b) After the short circuit is removed.

circuit is removed. Forces continue to restore the dielectric to its neutral state. With the circuit open, the charges held on the plates cannot be removed. As the dielectric returns to its normal condition, the trapped charges on the plates produce a voltage between A and B. This voltage may be serious hazard to a workman who expected the capacitance between the plates to be discharged by the short-time application of a short circuit. This hazard is particularly serious on equipment of high capacitance such as high-voltage cables, static capacitors, and generator windings. For this reason, it is always desirable to keep such equipment continuously short-circuited when workmen are to be in physical contact with the presumably deenergized equipment.

Referring again to Figure 3-2a and b, the movement of particles, such as the polar molecule *P*, may result in the movement of other nonpolar molecules. If the molecular motion is increased, the temperature of the material is increased. If the power supply is an ac source, each reversal of voltage will tend to cause a reversal of the position of the polar molecules and electrical energy from the source will be converted to heat in the insulation. This loss is known as *dielectric hysteresis*. It increases with frequency and with applied voltage. It must be considered in high-voltage cable design.

Conduction Currents

When voltage is applied between two plates separated by a dielectric (Figure 3-1), those few free electrons that are present in the insulation drift from cathode to anode. This is termed a conduction current (from anode to cathode) and represents power loss into the insulation. In insulation, the number of free electrons is low, and as a result the resistivity of the material is high. The number of free electrons may be increased by an increase in temperature.

Surface Leakage Currents

Leakage currents flow along paths between electrodes over the surface of the insulating material. The magnitude of these currents is in no way related to the resistivity of the material itself. The value of the leakage current depends on the applied voltage, the insulation material, the surface contamination, and the moisture content of the air. On seriously contaminated high-voltage line insulator surfaces leakage currents may be as much as 100 milliamperes.

Insulation Breakdown

Insulation may undergo a very sudden change in characteristics in a process known as

breakdown. Consider the arrangement shown in Figure 3-4. Two parallel-plane electrodes *A* and *B* are separated by a sheet of dielectric of thickness *t*. A variable voltage source *V* provides a difference of potential between *A* and *B*. Suppose the voltage is slowly raised. At first the conduction current is very low, perhaps measurable in microamperes. With increased applied voltage, the current suddenly increases, and the insulation takes on the character of a metallic conductor. This is termed insulation breakdown. On examination, a small damaged place may be found extending through the insulation sheet. Perhaps there will be some charring and perhaps there will be a hole.

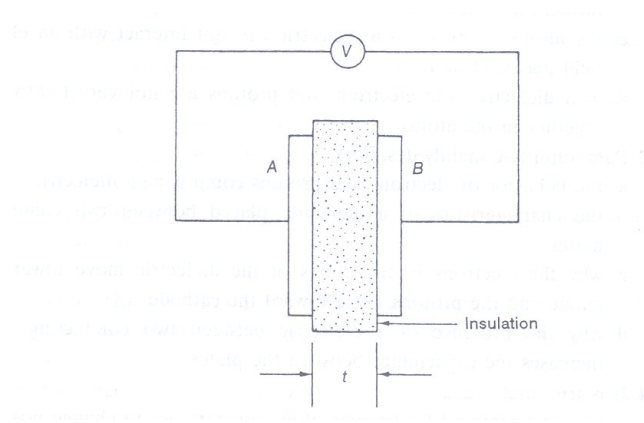


Figure 3-4. Insulation of Thickness *t* Being Stressed by Applied Voltage *V*.

The voltage at which such breakdown occurs is called the *breakdown*, or *puncture voltage*, V_{Ps} and the electric field intensity ϵ_p at that point is known as the *breakdown gradient* or *puncture strength* of the insulation, where *t* is the insulation thickness. The puncture strength of a particular sample is not a constant but varies with the thickness of the insulation, the shape and geometry of the electrodes, and the *rate* of application of voltage.

Comprehension Exercises

A. Choose a, b, c, or d which best completes each item.

1. The first paragraph mainly describes
 - a. the phenomenon of capacitance affected by a dielectric
 - b. the phenomenon of leakage current produced by a dielectric
 - c. the dielectric behavior when placed between two metallic conductors connected to a voltage source
 - d. the dielectric connecting a voltage source to two metallic conductors
2. It is true that
 - a. the presence of an electric field causes the electrons and protons of a dielectric to move
 - b. in a perfect insulator, the electrons and protons will never be influenced by any external factor
 - c. the atomic particles of a dielectric will not interact with an electric field placed close to it
 - d. in a dielectric, the electrons and protons are not very tightly held together in the atoms
3. Paragraph five mainly describes
 - a. the behavior of electrons and protons comprising a dielectric
 - b. the characteristics of a dielectric placed between two conducting plates
 - c. why the electrons in the atoms of the dielectric move toward the anode and the protons move toward the cathode
 - d. why the presence of a dielectric between two conducting plates increases the capacitance between the plates
4. It is true that
 - a. the time required for systems of atomic particles to change positions varies from one dielectric to another
 - b. some internal forces usually cause the systems of atomic particles to move and change position
 - c. the capacitance between the plates does not change as the atomic particles change positions
 - d. some insulating materials when disconnected from the voltage source do not return to their normal state
5. Having removed the short circuit, equipments of high capacitance will
 - a. discharge the trapped charges on the plates very quickly and cause the dielectric to return to its normal state

- b. be seriously dangerous to anybody in physical contact with them
 - c. displace the particles in the dielectric and cause a serious hazard to the workmen
 - d. be unable to remove the charges on the plates and cause the dielectric to go back to its normal state
6. As we understand from the text,
- a. the number of free electrons in an insulating material decreases with an increase in temperature
 - b. the resistivity of an insulating material decreases with an increase in temperature
 - c. extreme temperature will permanently lower insulation resistance
 - d. dielectric hysteresis will cause loss of heat in the dielectric
7. It is true that
- a. the lower the resistivity of a dielectric the higher the magnitude of leakage currents
 - b. the higher the resistivity of a dielectric the lower the magnitude of leakage currents
 - c. the value of the leakage current depends on factors such as surface contamination and air moisture
 - d. the value of the leakage depends on factors such as temperature and pressure

B. Write the answers to the following questions.

1. How does a dielectric affect the capacitance between two conducting plates?
2. What may cause the electrons and protons of an insulator to move?
3. What is the phenomenon of dielectric absorption?
4. What happens if the voltage source is disconnected from the plates and a short-circuiting connection is made?
5. What will happen if the short circuit is removed?
6. Why does the temperature of a dielectric between two conducting plates increase?
7. What is dielectric hysteresis?
8. How does power loss into the dielectric occur?
9. How does temperature affect an insulator as compared with a metal?
10. What is insulation breakdown?



Section Three: Translation Activities

A. Translate the following passage into Persian.

Dielectric Heating

Dielectric heating is a method of heating a nonconducting material, a dielectric, by high-frequency voltages. The material is placed between metal plates across which a high-frequency supply is connected as shown in Figure 3-5. The dielectric and the plates then form a capacitor and an electrostatic field is set up in the dielectric. As very high frequencies are used, up to 200 MHz, the movement of electrons in the dielectric becomes rapid. This causes considerable heat in the substance.

Dielectric heating has two great advantages over other forms of heating: it provides rapid heat, and the heat is produced uniformly throughout the material. In other words, the inside of the material

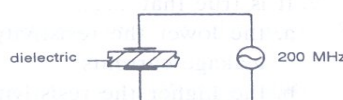


Figure 3-5.

gets hot at the same time as the surface. In addition, dielectric heating can be easily controlled and it is predictable. Accurate heating times can be calculated knowing the dielectric properties of the materials to be heated.

Dielectric heating has many different uses, from the manufacture of plastic raincoats to baking biscuits. It is especially used in plastics, wood-working, and food industries.

A typical use is the manufacture of plywood. In the past, the layers of wood and glue were steam-heated under pressure until the glue melted and the wood was firmly bonded. The heat took a long time to penetrate the wood, the glue did not melt uniformly and it dried unevenly. With dielectric heating, because of the difference in dielectric properties, the glue melts before the wood heats. It heats uniformly and it dries evenly. Using the dielectric process, a single press can prepare 100 3-ply, 1 cm thick sheets of plywood in about 30 minutes.

B. Find the Persian equivalents of the following terms and expressions and write them in the spaces provided.

1. absorption
2. annular disc
3. breakdown gradient

4. circumstance

5. contaminate	
6. deprive	
7. destruct g. deteriorate	
9. dielectric	
10. disturbance	
11. duct	
12. hysteresis	
13. insulate	
14. leakage	
15. microammeter	
16. oil-impregnated paper 17. overvoltage	
18. puncture	
19. shatter	
20. shielding wire	
21. soak	
22. solubility	
23. stroke	
24. trace	
25. truncated cone	
26. varnished cambric	

Unit 4

Section One: Reading Comprehens

The Distribution System

Although there is no 'typical' electric power system, a diagram including the several components that are usually to be found in the makeup of such a system is shown in Figure 4-1; particular attention should be paid to those

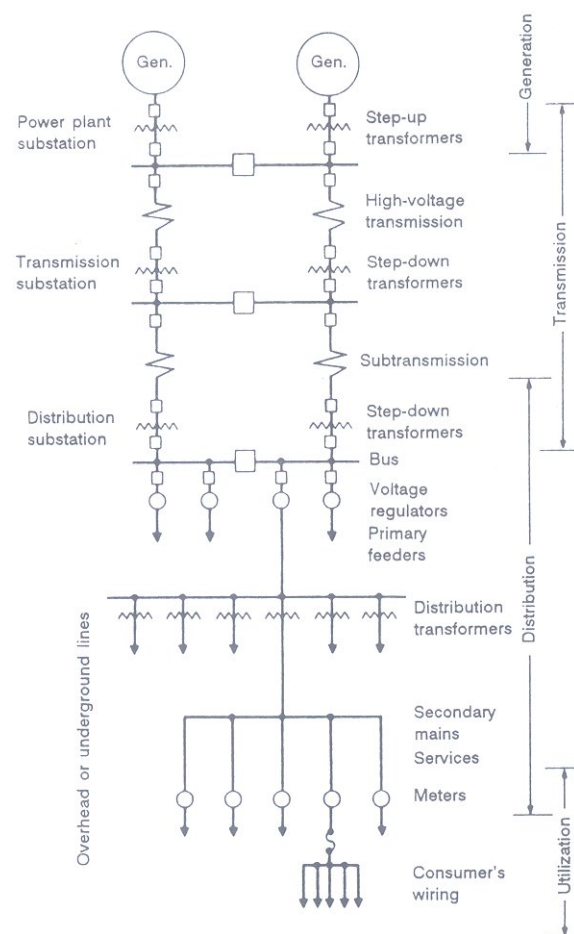


Figure 4-1. Typical Electric System Showing Operational Divisions. Note overlap of divisions.

elements which will make up the component under discussion, the distribution system.

While the energy flow is obviously from the power generating plant to the consumer, it may be more informative for our purposes to reverse the direction of observation and consider events from the consumer back to the generating source.

Energy is consumed by users at a nominal utilization voltage that may range generally from 110 to 125 V, and from 220 to 250 V, the nominal figures are 277 and 480 V. It flows through a metering device that determines the billing for the consumer, but which may also serve to obtain data useful later for planning, design, and operating purposes. The metering equipment usually includes a means of disconnecting the consumer from the incoming supply should this become necessary for any reason.

The energy flows through conductors to the meter from the secondary mains (if any); these conductors are referred to as the consumer's *service*, or sometimes also as the *service drop*.

Several services are connected to the secondary mains; the secondary mains now serve as a path to the several services from the distribution transformers which supply them.

At the transformer, the voltage of the energy being delivered is reduced to the utilization voltage values from higher primary line voltages that may range from 2200 V to as high as 46,000 V.

The transformer is protected from overloads and faults by fuses or so-called weak links on the high-voltage side; the latter also usually include circuit-breaking devices on the low-voltage side. These operate to disconnect the transformer in the event of overloads or faults. The circuit breakers (where they exist) on the secondary, or low-voltage, side operate only if the condition is caused by faults or overloads in the secondary mains, services, or consumers' premises; the primary fuse or weak link, in addition, operates in the event of a failure within the transformer itself.

If the transformer is situated on an overhead system, it is also protected from lightning or line voltage surges by a surge arrester, which drains the voltage surge to ground before it can do damage to the transformer.

The transformer is connected to the primary circuit, which may be a lateral or spur consisting of one phase of the usual three-phase primary main. This is done usually through a line or sectionalizing fuse, whose function is to disconnect the lateral from the main in the event of fault or overload in the lateral. The lateral conductors carry the sum of the energy components

flowing through each of the transformers, which represent not only the energy used by the consumers connected thereto, but also the energy lost in the lines and transformers to that point.

The three-phase main may consist of several three-phase branches connected together, sometimes through other line or sectionalizing fuses, but sometimes also through switches. Each of the branches may have several single-phase laterals connected to it through line or sectionalizing fuses.

Where single-phase or three-phase overhead lines run for any considerable distance without distribution transformer installations connected to them, surge arresters may be installed on the lines for protection.

Some three-phase laterals may sometimes also be connected to the three-phase main through *circuit reclosers*. The recloser acts to disconnect the lateral from the main should a fault occur on the lateral, much as a line or sectionalizing fuse. However, it acts to reconnect the lateral to the main, reenergizing it one or more times after a time delay in a predetermined sequence before remaining open permanently. This is done so that a fault which may be only of a temporary nature, such as a tree limb falling on the line, will not cause a prolonged interruption of service to the consumers connected to the lateral.

The three-phase mains emanate from a *distribution substation*, supplied from a *bus* in that station. The three-phase mains, usually referred to as a *circuit* or *feeder*, are connected to the bus through a protective circuit breaker and sometimes a voltage regulator. The voltage regulator is usually a modified form of a transformer and serves to maintain outgoing voltage within a predetermined band or range on the circuit or feeder as its load varies. It is sometimes placed electrically in the substation circuit so that it regulates the voltage of the entire bus rather than a single outgoing circuit or feeder, and sometimes along the route of a feeder for partial feeder regulation. The circuit breaker in the feeder acts to disconnect that feeder from the bus in the event of overload or fault on the outgoing or distribution feeder.

The substation bus usually supplies several distribution feeders and carries the sum of the energy supplied to each of the distribution feeders connected to it. In turn, the bus is supplied through one or more transformers and associated circuit breaker protection. These substation transformers step down the voltage of their supply circuit, usually called the *subtransmission* system, which operates at voltages usually from 23,000 to 138,000 V.

The subtransmission systems may supply several distribution substations and may act as *tie feeders* between two or more substations that are either of

the *bulk power* or *transmission* type or of the distribution type. They may also be tapped to supply some distribution load, usually through a circuit breaker,

for a single consumer, generally an industrial plant or a commercial consumer having a substantially large load.

The transmission or bulk power substation serves much the same purposes as a distribution substation, except that, as the name implies, it handles much greater amounts of energy: the sum of the energy individually supplied to the subtransmission lines and associated distribution substations and losses. Voltages at the transmission substations are reduced to outgoing subtransmission line voltages from transmission voltages that may range from 69,000 to upwards of 750,000 V.

The transmission lines usually emanate from another substation associated with a power generating plant. This last substation operates in much the same manner as other substations, but serves to step up to transmission line voltage values the voltages produced by the generators. Because of material and insulation limitations, generator voltages may range from a few thousand volts for older and smaller units to some 20,000 volts for more recent, larger ones. Both buses and transformers in these substations are protected by circuit breakers, surge arresters, and other protective devices.

In all the systems described, conductors should be large enough that the energy loss in them will not be excessive, nor the loss in voltage so great that normal nominal voltage ranges at the consumers' services cannot be maintained.

In some instances, voltage regulators and capacitors are installed at strategic points on overhead primary circuits as a means of compensating for voltage drops or losses, and incidentally help in holding down energy losses in the conductors.

In many of the distribution system arrangements, some of the several elements between the generating plant and the consumer may not be necessary. In a relatively small area, such as a small town, that is served by a power plant situated in or very near the service area, the distribution feeder

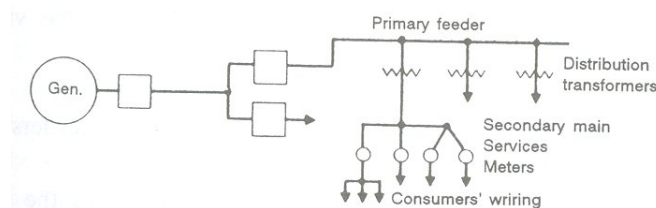


Figure 4-2. 'Abbreviated' Electric System.

may emanate directly from the power plant *bus*, and all other elements may be eliminated, as indicated in Figure 4-2. This is perhaps one extreme; in many other instances only some of the other elements may not be necessary; e.g., a similar small area somewhat distant from the generating plant may find it necessary to install a distribution substation supplied by a transmission line of appropriate voltage only.

In the case of areas of high load density and rather severe service reliability requirements, the distribution system becomes more complex and more expensive. The several secondary mains to which the consumers' services are connected may all be connected into a mesh or network. The transformers supplying these secondary mains or network are supplied from several different primary feeders, so that if one or more of these feeders is out of service for any reason, the secondary network is supplied from the remaining ones and service to the consumers is not interrupted. To prevent a feeding-back from the energized secondary network through the transformers connected to feeders out of service (thereby energizing the primary and creating unsafe conditions), automatically operated circuit breakers, called *network protectors*, are connected between the secondary network and the secondary of the transformers; these open when the direction of energy flow is reversed.

Part I. Comprehension Exercises

A. Put "T" for true and "F" for false statements. Justify your answer

- 1. The text describes the distribution elements used between the power generating plant and the consumer.
- 2. The metering device is mainly used to offer data useful for design and operating purposes.
- 3. The three-phase main may consist of several three-phase branches which in turn may consist of several single-phase laterals.
- 4. Any power system must have secondary mains in order to supply the consumer with energy through the services.
- 5. Primary feeders are connected to the substation bus via voltage regulators or protective circuit breakers.
- 6. The voltage regulator is the same as the transformer.
- 7. Subtransmission systems may be used as tie feeders between different types of substations.
- 8. The closer the power plant to the area it serves the fewer the elements between the generating plant and the consumer.

9. Distribution transformers may directly supply the consumer with

energy.

10. The elemental arrangement in a complex distribution system is so that service interruption is improbable.

B. Choose a, b, c, or d which best completes each item.

1. It is true that
 - a. circuit breakers disconnect the high-voltage side of the transformer in the event of overloads
 - b. fuses and circuit breakers are identical devices
 - c. fuses are weak links always installed on the secondary side of the transformer
 - d. circuit breakers protect the high-voltage side of the transformer in the event of overloads
2. The distribution transformer
 - a. is connected to the primary circuit through a line or sectionalizing fuse
 - b. is connected to the primary main through a line or sectionalizing fuse
 - c. helps to disconnect the lateral from the main in the event of fault or overload
 - d. helps to disconnect the consumers from the services in the event of fault or overload
3. A circuit recloser is used to
 - a. connect a three-phase lateral to the three-phase main
 - b. disconnect the lateral from the main if a fault occurs on the lateral
 - c. reconnect the lateral to the main after a predetermined time delay
 - d. all of the above
4. It is true that
 - a. the substation bus supplies substation transformers
 - b. substation transformers supply the substation bus
 - c. the substation bus is supplied by the distribution feeders
 - d. substation transformers are the same as subtransmission systems
5. The bulk power substation handles the energy supplied to
 - a. the subtransmission lines and associated substations
 - b. the substation bus and the primary feeders
 - c. the distribution transformers and the meters
 - d. the secondary mains and associated services

6. According to the text,
- a. material and insulation limitations do not allow the generators to work at their full capacities
 - b. material and insulation limitations have resulted in the use of various protective devices to protect the generators
 - c. the voltages produced by the generators are stepped up in the substation associated with the power plant
 - d. the voltages produced by the generators are reduced to usable voltages in the substation associated with the power plant
7. The last paragraph mainly describes
- a. consumers' service interruption b. consumers' service reliability
 - c. a complex distribution network d. a complex distribution system

C. Answer the following questions orally.

1. What is called the service?
2. What is the function of a surge arrester?
3. What does a lateral refer to?
4. What is the function of a voltage regulator?
5. What part does the substation bus play in the distribution system?
6. Where do the transmission lines originate from?
7. What is the function of a capacitor installed on an overhead primary circuit?
8. When does the distribution system become more complex?
9. What is a network?
10. How do network protectors help a distribution system?

Part II Language Practice

A. Choose a, b, c, or d which best completes each item.

1. The deliver electric energy from the secondary distribution or street main, or other distribution feeder, or from the transformer, to the wiring system of the premises served.
 - a. meters b. buses
 - c. services d. feeders
2. The function of is to interrupt circuit faults.
 - a. a line b. a service
 - c. a main d. a transformer
3. A serves as a common connection for two or more circuits.
 - a. fuse b. switch
 - c. lateral d. bus

4. Two or more generating systems, substations, or feeding points may be connected together by
 - a. a tie wire
 - b. a tie trunk
 - c. a tie feeder
 - d. a tie line
5. To automatically disconnect a transformer from a secondary network in response to predetermined electric conditions on the primary feeder or transformer, are employed.
 - a. network relays
 - b. network protectors
 - c. circuit reclosers
 - d. circuit analyzers

B. Fill in the blanks with the appropriate form of the words given.

1. Connect

- a. A connection diagram shows the of an installation or its component devices, controllers, and equipment.
- b. A network is if there exists at least one path composed of branches of the network, between every pair of nodes of the network.
- c. A low voltage or secondary network is a continuous secondary main or grid fed by a number of transformers to the same primary feeder.

2. Protect

- a. To ensure maximum , the system must possess a high degree of electricity.
- b. equipment should be used against vibrations of voltage.
- c. A differential relay responds to the difference between incoming and outgoing electrical quantities associated with the apparatus.

3. Limit

- a. The function of a relay is to prevent or damage during faults.
- b. The inrush current of the rectifier transformer is generally the factor.
- c. Hard limiting is a limiting action with negligible variation in output in the range where the output is
- d. A bridge is a bridge circuit used as a limiter circuit.

4. Regulate

- a. The substation may or may not require voltage equipment.
- b. The circuit on the output side of the is known as the voltage
- c. The voltage may be held constant at any selected point on the circuit.

- d. A voltage relay is used on an automatically operated voltage regulator to control the voltage of the regulated circuit.

5. Distribute

- a. Electric power is received from substations and is to the consumers at voltage levels and degrees of continuity that are acceptable to various types of consumers.
- b. For a transverse electromagnetic wave on a two-conductor transmission line, the constants are series resistance, series inductance shunt conductance, and shunt capacitance per unit length of line.
- c. A distribution switchboard is used for the distribution of electric energy at voltages common for such in a building.
- d. A duct installed for occupancy of distribution mains is known as a duct

C. Fill in the blanks with the following words.

short-circuit	mechanical	failure	apart
interruption	electrical	simply	pull
conductors	ordinary		

Since a failure of a conductor results in a complete to a circuit, it is imperative that the causes of such be minimized. The failure may occur from causes where the stresses and strains imposed are too great and the conductors literally tear More often, however, the cause may initially be a/an failure which then affects the conductor mechanically. Overloads or currents, for example, may cause heating of the to the point where they begin to liquefy and mechanical stresses can no longer be sustained and the conductors apart, perhaps vaporizing in the process.

D. Put the following sentences in the right order to form a paragraph. Write the corresponding letters in the boxes provided.

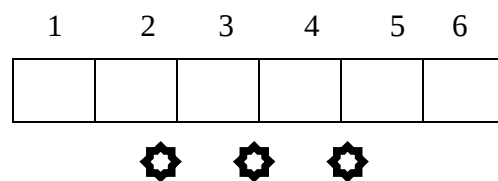
- a. This accessibility and inaccessibility should prevail even under advert or contingency conditions,
- b. Moreover, the conductors and equipment on the poles should be so situated that they can be handled safely by the people working on them
- c. For example, a distribution pole line should be so located that free and easy access to the facilities is available at all times, yet it should not

interfere with pedestrian and vehicular traffic, nor intrude into areas (such as playgrounds) where its presence may constitute a particular hazard.

d. Safe methods include the use of protective equipment, the use of live-line tools and equipment, and the deenergization and grounding of the facilities on which work is to be performed.

e. As much as practical, the utility's facilities have to be both accessible (to the workers) and inaccessible (to the public).

f. This not only implies providing sufficient working space, but includes considerations of how the work may be performed safely



Section Two: Further Reading

Types of Delivery Systems

The delivery of electric energy from the generating plant to the consumer may consist of several more or less distinct parts that are nevertheless somewhat interrelated. The part considered 'distribution', i.e., from the bulk supply substation to the meter at the consumer's premises, can be conveniently divided into two subdivisions:

1. Primary distribution, which carries the load at higher than utilization voltages from the substation (or other source) to the point where the voltage is stepped down to the value at which the energy is utilized by the consumer.
2. Secondary distribution, which includes that part of the system operating at utilization voltages, up to the meter at the consumer's premises.

Primary Distribution

Primary distribution systems include three basic types:

Radial Systems. The radial-type system is the simplest and the one most commonly used. It comprises separate feeders or circuits 'radiating' out of the station or source, each feeder usually serving a given area. The feeder may

be considered as consisting of a main or trunk portion from which there radiate spurs or laterals to which distribution transformers are connected, as illustrated in Figure 4-3.

The spurs or laterals are usually connected to the primary main through fuses, so that a fault on the lateral will not cause an interruption to the entire feeder. Should the fuse fail to clear the line, or should a fault develop on the feeder main, the circuit breaker back at the substation or source will open and the entire feeder will be deenergized.

To hold down the extent and duration of interruptions, provisions are made to sectionalize the feeder so that unfaulted portions may be reenergized as quickly as practical. To maximize such reenergization, emergency ties to adjacent feeders are incorporated in the design and construction; thus each part of a feeder not in trouble can be tied to an adjacent feeder. Often spare capacity is provided for in the feeders to prevent overload when parts of an adjacent feeder in trouble are connected to them. In many cases, there may be

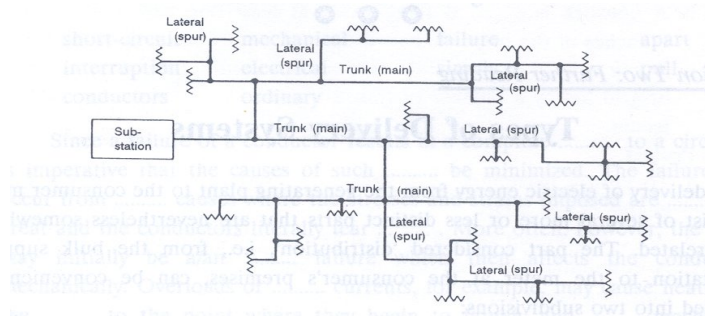


Figure 4-4. Primary Feeder Schematic Diagram Showing Trunk or Main Feeds and Laterals or Spurs .

enough diversity between loads on adjacent feeders to require no extra capacity to be installed for these emergencies.

Loop Systems. Another means of restricting the duration of interruption employs feeders designed as loops, which essentially provide a two-way primary feed for critical consumers. Here, should the supply from one direction fail, the entire load of the feeder may be carried from the other end, but sufficient spare capacity must be provided in the feeder. This type of system may be operated with the loop normally open or with the loop normally closed.

Primary Network Systems. Although economic studies indicated that under some conditions the primary network may be less expensive and more reliable than some variations of the radial system, relatively few primary network systems have been put into actual operation and only a few still remain in service.

This system is formed by lying together primary mains ordinarily found in radial systems to form a mesh or grid. The grid is supplied by a number of power transformers supplied in turn from subtransmission and transmission lines at higher voltages. A circuit breaker between the transformer and grid, controlled by reverse-current and automatic reclosing relays, protects the primary network from feeding fault current through the transformer when faults occur on the supply subtransmission lines. Faults on sections of the primaries constituting the grid are isolated by circuit breakers and fuses. See Figure 4-4.

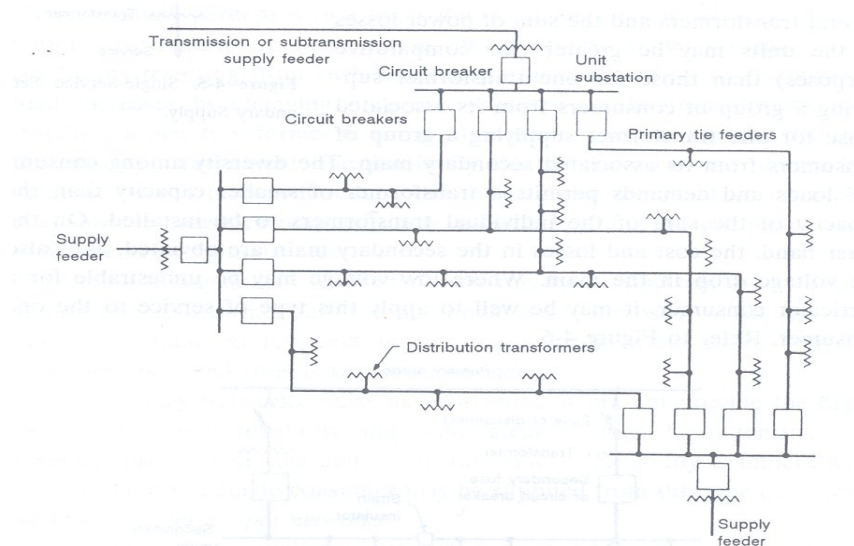


Figure 4-4. Primary Network. Sectionalizing devices on feeders not shown.

This type of system eliminates the conventional substation and long primary trunk feeders, replacing them with a greater number of 'unit' substations strategically placed throughout the network. The additional sites

necessary are often difficult to obtain. Moreover, difficulty is experienced in maintaining proper operation of the voltage regulators (where they exist) on the primary feeders when interconnected.

Secondary Distribution

Secondary distribution systems operate at relatively low utilization voltages and, like primary systems, involve considerations of service reliability and voltage regulation. The secondary system may be of four general types:

Individual Transformers-Single Service. Individual-transformer service is applicable to certain loads that are more or less isolated, such as in rural areas where consumers are far apart and long secondary mains are impractical, or where a particular consumer has an extraordinary large or unusual load even though situated among a number of ordinary consumers.

In this type of system, the cost of the several transformers and the sum of power losses in the units may be greater (for comparative purposes) than those for one transformer supplying a group of consumers from its associated secondary main. The diversity among consumers' loads and demands permits a transformer of smaller capacity than the capacity of the sum of the individual transformers to be installed. On the other hand, the cost and losses in the secondary main are obviated, as is also the voltage drop in the main. Where low voltage may be undesirable for a particular consumer, it may be well to apply this type of service to the one consumer. Refer to Figure 4-6.

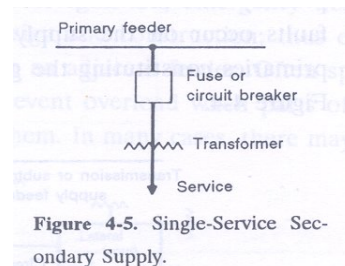


Figure 4-5. Single-Service Secondary Supply.

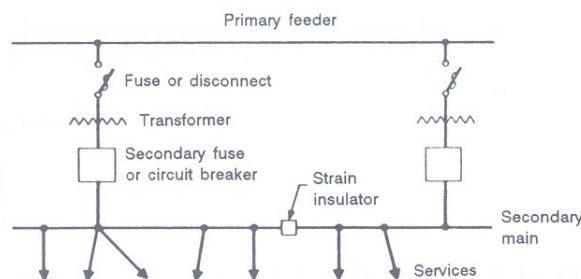


Figure 4-6. Common-Secondary-Main Supply.

Figure 4-6. Common-Secondary-Main Supply.

Common Secondary Main. Perhaps the most common type of secondary system in use employs a common secondary main. It takes advantage of

diversity between consumers' loads and demands, as indicated above. Moreover, the

larger transformer can accommodate starting currents of motors with less resulting voltage dip than would be the case with small individual transformers. See Figure 4-6.

Banked Secondaries. The secondary system employing banked secondaries is not very commonly used, although such installations exist and are usually limited to overhead systems.

This type of system may be viewed as a single-feeder low-voltage network, and the secondary may be a long section or grid to which the transformers are connected. Fuses or automatic circuit breakers located between the transformer and secondary main serve to clear the transformer from the bank in case of failure of the transformer. Fuses may also be placed in the secondary main between transformer banks. See Figure 4-7.

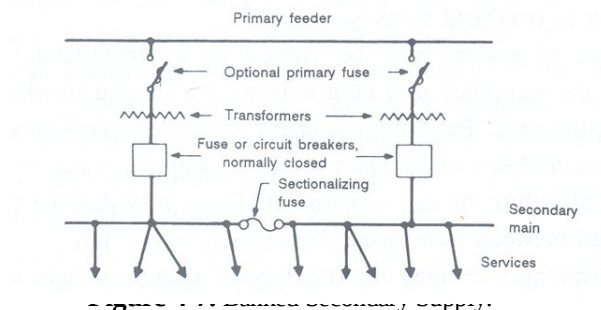
Some advantages claimed for this type of system include uninterrupted service, though perhaps with a reduction in voltage, should a transformer fail; better distribution of load among transformers; better normal voltage conditions resulting from such load distribution; an ability to accommodate load increases by changing only one or some of the transformers, or by installing a new transformer at some intermediate location without disturbing the existing arrangement; the possibility that diversity between demands on adjacent transformers will reduce the total transformer load; more capacity available for inrush currents that may cause flicker; and more capacity as well to burn secondary faults clear.

Some disadvantages associated with this type of system are as follows: should one transformer fail, the additional loads imposed on adjacent units may cause them to fail, and in turn their loads would cause still other transformers to fail (this is known as *cascading*).

Secondary Networks. Secondary networks at present provide the highest degree of service reliability and serve areas of high load density, where revenues justify their cost and where this kind of reliability is imperative. In some instances, a single consumer may be supplied from this type of system by what are known as *spot networks*.

In general, the secondary network is created by connecting together the secondary mains fed from transformers supplied by two or more primary feeders. Automatically operated circuit breakers in the secondary connection between the transformer and the secondary mains, known as *network protectors*, serve to disconnect the transformer from the network when its primary feeder is deenergized; this prevents a back feed from the secondary

into the primary feeder. This is especially important for safety when the primary feeder is deenergized from fault or other cause. The circuit breaker or protector is backed up by a fuse so that, should the protector fail to operate, the fuse will blow and disconnect the transformer from the secondary mains. See Figure 4-7.



The number of primary feeders supplying a network is very important. With only two feeders, only one feeder may be out of service at a time, and there must be sufficient spare transformer capacity available so as not to overload the units remaining in service; therefore this type of network is sometimes referred to as a *single-contingency* network.

Most networks are supplied from three or more primary feeders, where the network can operate with the loss of two feeders and the spare transformer capacity can be proportionately less. These are referred to as *second-contingency* networks.

Comprehension Exercises

A. Choose a, b, c, or d which best completes each item.

1. A radial distribution system
 - a. consists of a source from which mains and laterals radiate
 - b. consists of the trunks from which laterals equipped with transformers radiate.
 - c. includes that part of the system that employs feeders designed as open loops.
 - d. includes that part of the system that employs feeders designed as closed loops.
2. The circuit breaker at the source deenergizes the entire feeder
 - a. if a fault develops on the feeder main

- b. as soon as the fuses installed at the lateral and main intersections begin to work

- c. if any of the fuses installed at each lateral and main intersection fails to operate
 - d. both a and c
3. In order to minimize the duration of interruptions,
- a. feeders are sectionalized
 - b. adjacent feeders are provided with emergency ties
 - c. spare capacity is provided in the feeder
 - d. all of the above
4. It is true that
- a. primary network systems are more popular than radial systems
 - b. voltage regulators installed on the primary network feeders do not always operate properly
 - c. primary network systems are usually less reliable than radial systems
 - d. circuit breakers between the grid and the transformers do not function very well
5. Individual-transformer service is applied
- a. where an individual consumer has an extraordinary large load
 - b. where an individual consumer is not distant from the source
 - c. to increase power reliability in thinly-populated areas
 - d. to reduce the sum of power losses in the units
6. One crucial disadvantage of a low-voltage network is that
- a. distribution of load among transformers is not uniform
 - b. additional transformers are difficult to be installed
 - c. if one transformer fails, other transformers may also fail
 - d. if flickers appear, they cannot be prevented
7. A network protector in a secondary network is employed
- a. to serve as a path between the transformer and the secondary mains
 - b. to deenergize the primary feeder and disconnect the transformer from the network
 - c. to blow the fuse and disconnect the transformer from the secondary mains
 - c. to prevent a back feed from the secondary mains into the primary feeder

B. Write the answers to the following questions.

1. What are the two subdivisions of the distribution system?
2. What does the primary system include?

3. What is the function of the fuses that connect the laterals to the mains?
4. How does the loop system restrict the duration of interruption?
5. What part does the spare capacity provided in the loop system play?
6. What is the function of a circuit breaker installed between the transformer and the grid in a primary network system?
7. How are faults on the sections of the primaries isolated?
8. What replaces conventional substations in the primary network system?
9. What type of secondary distribution system is applied where long secondary mains are impractical?
10. What is the advantage of a large transformer over a small one?
11. What is cascading?
12. What characteristics should the transformers employed in a low-voltage network have?
13. How does a low-voltage network handle the diversity of demands?
14. What are the single-contingency and second-contingency networks?



Section Three: Translation Activities

A. Translate the following passage into Persian.

Maintainability

Each item selected, and its place in the distribution system, must be viewed in light of its possible failure or malfunction for whatever reason. The design of the distribution system, therefore, should take into consideration the method of maintaining each of the several elements making up the distribution system.

The Distribution System. In general, the safest means is to be able to deenergize and ground the particular item requiring maintenance, preferably without affecting the remainder of the circuit. Circuits, both primary and secondary, are arranged so that small sections may be deenergized by interconnecting the remaining portions to other sources, by means of some sort of switches. Smaller pieces may be deenergized by means of hot-line or live-line clamps.

Where deenergization is not practical, work may be carried out by insulating the worker. This is accomplished by protective gear, such as rubber gloves, sleeves, blankets, line hose, insulator hoods, and other similar devices

Another method insulates workers from ground by having them work from insulated

platforms or insulated buckets mounted on line trucks. Still another means involves the handling of the energized facilities by tools having sufficient insulation properties which, properly handled, enable the worker to accomplish required maintenance by essentially a remote operation of the tools; this is referred to as live-line, or hot-line, maintenance. To facilitate such operations, appropriate details and modifications are included in the design of distribution systems, e.g., wider spacing of conductors; hot-line ties that hold conductors to the insulators; 'unnecessary' extensions of primary and secondary mains so that mains butt each other, permitting their temporary connection during contingencies by means of jumpers and the arrangement of terminals at substations to accommodate portable substations.

B. Find the Persian equivalents of the following terms and expressions and write them in the spaces provided.

- | | |
|--------------------------------|-------|
| 1. circuit fault | |
| 2. compensate | |
| 3. emanate | |
| 4. facility | |
| 5. flicker | |
| 6. inrush current | |
| 7. lateral | |
| 8. limb | |
| 9. malfunction | |
| 10. obviate | |
| 11. overload | |
| 12. premise | |
| 13. radial system | |
| 14. regulator | |
| 15. sectionalizing fuse | |
| 16. service drop | |
| 17. short circuit | |
| 18. single-contingency network | |
| 19. spare capacity | |
| 20. spot network | |
| 21. spur | |
| 22. substation bus | |
| 23. sustain | |
| 24. trunk | |

Unit 5

Section One: Reading Comprehension

Protective Devices

For the distribution system to function satisfactorily, faults on any part of it must be isolated or disconnected from the rest of the system as quickly as possible; indeed, if possible, they should be prevented from happening. The principal devices to accomplish this include fuses, automatic sectionalizers, reclosers, circuit breakers, and lightning or surge arresters. Success, however, depends on their coordination so that their operations do not conflict with each other. Figure 5-1 indicates where these devices are connected on the system.

Fuses

Time-Current Characteristic. A fuse consists basically of a metallic element that melts when ‘excessive’ current flows through it. The magnitude of the excessive current will vary inversely with its duration. This time-current characteristic is determined not only by the type of metal used and its dimensions (including its configuration), but also on the type of its enclosure and holder. The latter not only affect the melting time, but in addition, affect the arc clearing time. The *clearing time* of the fuse, then, is the sum of the melting time and the arc clearing time. Refer to Figures 5-1 and 5-2. Note that for curve *b* in Figure 5-3, the clearing time for a certain value of current is less than for curve *a*; the fuse with the characteristic *b* is therefore referred to as a ‘fast’ fuse, compared with the fuse of curve *a*.

Fuses are rated in terms of voltage, normal current-carrying ability, and interruption characteristics usually shown by time-current curves. Each curve actually represents a band between a minimum and a maximum clearing time for a particular fuse.

Fuse Coordination. The number, rating, and type of the interrupting devices shown in Figure 5-1 depend on the system voltage, normal current, maximum fault current, the sections and equipment connected to them, and other local conditions. The devices are usually located at branch intersections and at other key points. When two or more such devices are employed in a circuit, they will be coordinated so that only the faulted portion will be deenergized.

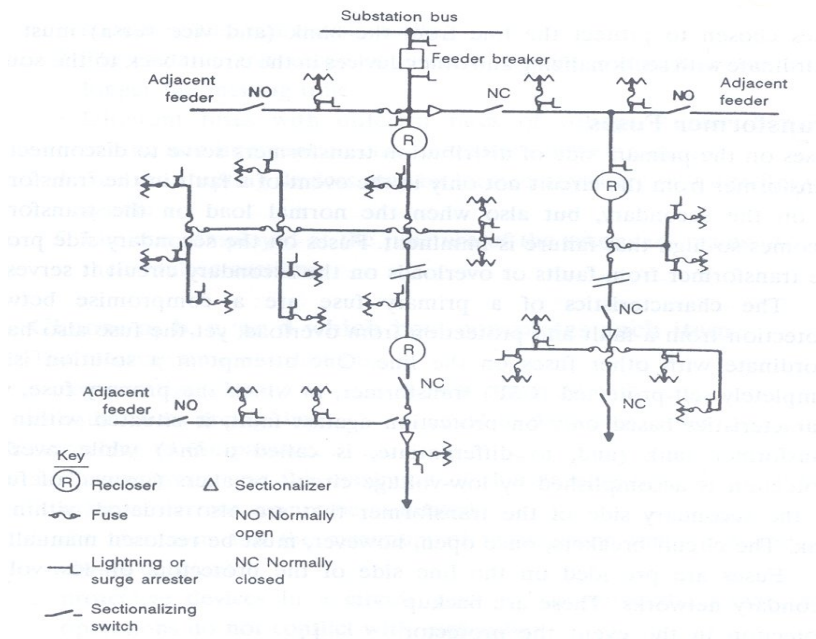


Figure 5-1. Radial Primary Feeder Showing Location of Protective Devices.

Repeater Fuses

Line fuses are sometimes installed in groups of two or three (per phase), known as repeater fuses, having a time delay between each two fuse units. When a fault occurs, the first fuse will blow and the second fuse will be mechanically placed in the circuit by the opening of the first; if the fault Persists, the second fuse will blow; if There a third fuse, the process is repeated. If the fault is permanent, all of the fuses will blow and the faulted part of the circuit will be deenergized. new fuses must be installed to restore the line to normal.

Where capacitors are applied to feeders for power factor correction,

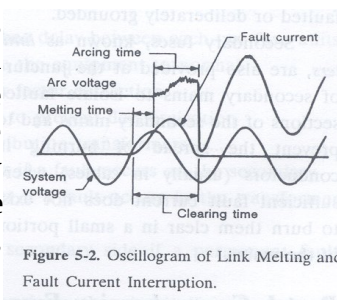


Figure 5-2. Oscillogram of Link Melting and Fault Current Interruption.

fuses chosen to protect the line from the bank (and vice versa) must also coordinate with sectionalizing and other devices in the circuit back to the source .

Transformer Fuses

Fuses on the primary side of distribution transformers serve to disconnect the transformer from the circuit not only in the event of a fault in the transformer or on the secondary, but also when the normal load on the transformer becomes so high that failure is imminent. Fuses on the secondary side protect the transformer from faults or overloads on the secondary circuit it serves.

The characteristics of a primary fuse are a compromise between protection from a fault and protection from overload, yet the fuse also has to coordinate with other fuses on the line. One attempt at a solution is the completely self-protected (CSP) transformer, in which the primary fuse, with characteristics based only on protection against fault, is situated within the transformer tank (and, to differentiate, is called a *link*) while overload protection is accomplished by low-voltage circuit breakers (instead of fuses) on the secondary side of the transformer that are also situated within the tank. The circuit breakers, once open, however, must be reclosed manually.

Fuses are provided on the line side of the protectors on low-voltage secondary networks. These are backup protection in the event the protector fails to open during back feed from the network into the primary when it is faulted or deliberately grounded.

Secondary fuses, known as *limi-ters*, are also provided at the juncture of secondary mains to isolate faulted sections of the secondary mains and to prevent the spread of burning in conductors (usually in cables) where sufficient fault current does not exist to burn them clear in a small portion of the mains.

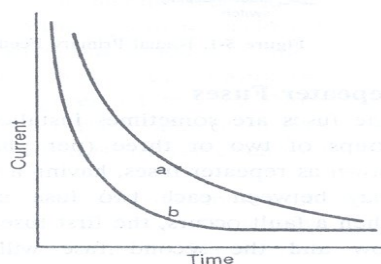


Figure 5-3. Typical Time-Current Characteristic for Fuses.

Part I. Comprehension Exercises

A. Put T for true and "F" for false statements. Justify your answers.

..... 1. Protective devices connected on a distribution system cause the

system to function satisfactorily.

- 2. The greater the excessive current flowing through a fuse, the longer the melting time.
- 3. Different fuses with different rates of voltage, current-carrying ability, and interruption characteristics can be produced.
- 4. The interrupting devices may be located anywhere in a distribution system.
- 5. Limners, employed at the juncture of the secondary mains, isolate their faulted sections.

B. Choose a, b, c, or d which best completes each item.

1. The clearing time of a fast fuse is
 - a. comparatively low b. comparatively high
 - c. equal to its arcing time d. equal to its melting time
2. According to the passage,
 - a. the minimum and the maximum of the clearing time for a certain value of current cannot be evaluated
 - b. the minimum and the maximum of the clearing time of a fuse are always constant
 - c. protective devices in a circuit must be coordinated so that their operations do not conflict with each other
 - d. protective devices in a circuit must be adjusted to deficiencies resulting from manufacturing problems
3. Repeater fuses
 - a. are installed in groups with a time delay between each two fuse units
 - b. are installed in series to restore the equipment to normal
 - c. may act as capacitors for power factor correction
 - d. may all operate simultaneously to prevent deenergization
4. Fuses on the primary side of distribution transformers
 - a. will not protect the transformer if a fault occurs on the secondary
 - b. will not protect the transformer if a fault occurs in the transformer itself
 - c. will protect the fuses on the secondary side if a permanent fault occurs
 - d. will disconnect the transformer from the circuit if it is seriously overloaded
5. It is true that
 - a. a transformer fuse is basically designed for fault protection

- b. a transformer circuit breaker is manually opened and closed
- c. a self-protected transformer is equipped with links and circuit breakers
- d. a self-protected transformer is fully automatic

C. Answer the following questions orally.

1. What kinds of protective devices are used in a distribution system?
2. What does the time-current characteristic depend on?
3. How important are the enclosure and the holder of the fuse?
4. What do the number, rating, and type of the interrupting devices applied in a circuit depend on?
5. What can a primary fuse do?
6. How are fuses used as backup protection on low-voltage secondary network?

Part II Language Practice

A. Choose a, b, c, or d which best completes each item.

1. Devices called are designed to open when a fault occurs on that part of the main in which they are connected.
 - a. regulators
 - b. fuses
 - c. reclosers
 - d. relays
2. Some are designed to open in air, with special provisions for handling the arc that follows when the contacts are opened.
 - a. fuses
 - b. arresters
 - c. line sectionalizers
 - d. circuit breakers
3. In liquid-filled construction, link is enclosed in a tube that is filled with a fire-extinguishing fluid such as carbon tetrachloride.
 - a. the fuse
 - b. the switch
 - c. the fault-counting relay
 - d. the circuit breaker
4. Switches are installed in the main of the feeder, enabling the main to be sectionalized, isolating the fault between two switches or other devices.
 - a. connecting
 - b. sectionalizing
 - c. energizing
 - d. feeding
5. Like the secondary circuit, the design of the primary is based on the maximum voltage variation permissible at the farthest consumer.
 - a. recloser
 - b. limiter
 - c. conductor
 - d. feeder

B. Fill in the blanks with the appropriate form of the words given.

1. Energize

- a. In a dependent time delay relay, the time delay varies with the value of the quantity.
- b. A relay is said to 'pick up' when it changes from the unenergized position to the position.
- c. Relay functioning time is the time between and operation or between deenergization and release.
- d. The elapsed time after the coil has been to the time required to seat the armature of the relay is the relay seating time.

2. Blow

- a. In a transformer circuit, the fuse is chosen so as to carry the inrush transient without
- b. If a fault persists in a circuit, the fuses may and the faulted part of the circuit will be deenergized.
- c. A blade is an active element of a fan.
- d. In liquid-filled fuses, a spring is used to hold the fuse under tension so that, when it the resultant arc is quickly lengthened and quenched in the fluid; the gas formed is inert and helps in out the arc.

3. Melt

- a. Melting-speed ratio refers to the ratio of the current magnitudes required to the current-responsive element at two specified melting times.
- b. The time required for overcurrent to sever the current-responsive element is known as the time.
- c. In a sand-filled fuse, the heat and gases generated when the fuse are absorbed by the sand, which tends to squelch the arc.

4. Interrupt

- a. An interrupted continuous wave is a wave that is at a constant audio-frequency rate.
- b. An is designed to interrupt specified currents under specified conditions.
- c. The loads that can be in the event of a capacity deficiency on the supplying system are loads.
- d. There will be the loss of service to one or more consumers or other facilities if occurs.

5. Inverse

- a. The dependent time delay relay, known as inverse time delay relay, has an operating time which is an function of the electrical characteristic quantity.
- b. The inverse time delay relay with definite minimum (I.D.M.T) is a relay in which the time delay varies with the characteristic quantity up to a certain value, after which the time delay becomes substantially independent.

C. Fill in the blanks with the following words.

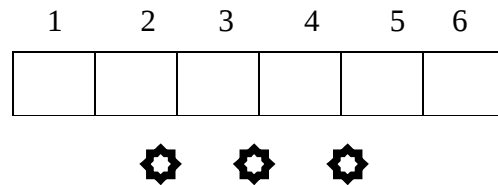
supplied	several	feeders	which
substations	auxiliary	load	adds
minimum	service		

The number and sources of supply subtransmission to the distribution substation will depend not only on the to be served, but also on the degree of reliability sought. Some rural substations may be from only one subtransmission feeder, serving urban and suburban areas have a/an of two supply feeders and may have more. Each additional incoming feeder, however, to the bus and switching requirements, including devices for their protection, all of add to costs.

D. Put the following sentences in the right order to form a paragraph. Write the corresponding letters in the boxes provided.

- a. The higher the distribution voltage, the farther apart substations may be located, but they also become larger in capacity and in the number of customers served.
- b. In general, it should be situated as close to the load center to be served as practical.
- c. The difficulty in obtaining substation sites is an important factor in selecting the distribution voltage, both in original designs and in later conversions.
- d. Thus, the problem of the number and location of distribution substations involves not only the study of transmission and subtransmission designs, but more emphasis on service reliability and consideration of additional costs that may be justified.

- e. This implies that all loads can be served without undue voltage regulation, including future loads that can be expected in a reasonable of time.
- f. Perhaps the first consideration regarding a distribution substation is its location.



Section Two: Further Reading

Automatic Line Sectionalizers

Automatic line sectionalizers are connected on the distribution feeder in series with line and sectionalizing fuses; they are also in series with and electrically farther from the source than reclosers or circuit breakers with reclosing cycles. These devices are decreasing in usage, but many exist on distribution systems.

When a fault occurs on the circuit beyond the sectionalizer, the fault current initiates a *fault-counting relay* that is coordinated with the characteristics of the fuses and other devices. Each time the circuit is deenergized (from reclosers or circuit breakers), the relay moves toward the trip position; just before the final operation that will lock out the recloser or circuit breaker if the fault persists, the sectionalizer will trip (while no fault current is flowing) and open the circuit at that point, removing the fault and permitting the circuit breaker or recloser to close and reset into its normal position; service is thus restored to the rest of the circuit up to the location of the sectionalizer. If the fault is of a temporary nature and is cleared before the reclosing devices complete their operations, the sectionalizer will reset to its position after the circuit is reenergized.

Sectionalizers are rated on continuous current-carrying capacity, min-tripping and counting current, and maximum momentary fault current, as well as for maximum system voltage, load-break current, and impulse voltage or basic insulation level (BIL).

More than one sectionalizer can be connected in series with a reclosing

device. The sectionalizer nearest the reclosing device can be set to operate after (say) three operations while the more remote one is set for (say) two such operations.

Sectionalizers are relatively low-cost devices; they are not required to interrupt fault current although fault current flows through them. They may be operated manually and are considered the same as load-break switches.

Reclosers

Reclosers are essentially circuit breakers of lower capacity, both as to normal current and interrupting duty. They are usually installed on major branches of distribution feeders in series with other sectionalizing devices; they perform the same function as repeater fuses connected in the circuit or circuit breakers at the substation.

Reclosers are designed to remain open, or 'locked out', after a selected sequence of tripping operations. A fault will trip the recloser; if the fault is temporary in nature and no longer exists, the next tripping operation does not take place and the recloser returns to its normally closed position, ready for another incident. If the fault persists, the recloser will close and the operation will be repeated until the recloser locks out. The reclosers are usually set for three automatic reclosing operations before locking out; the first operation is usually 'instantaneous', i.e., occurring as quickly as the breaker contacts can open with no time delay; the second and third operations have time delays inserted, that for the second tripping smaller than that for the third; a fourth tripping will result in the recloser's remaining open until it is automatically or manually restored to normal, ready for the next incident.

Reclosers can operate on one or more time-current characteristic curves. The reclosing characteristics of the recloser for each operation are coordinated with those of the fuses at the coordinating points in the circuit and with those of the relays controlling the circuit breaker at the substation.

Circuit Breakers-Relays

Where the fault current is beyond the ability of a fuse or recloser to interrupt it safely, or where repeated operation within a short period of time makes it more economical, a circuit breaker is used. The circuit breakers must not only interrupt the normal load current, but must be mechanically able to withstand the forces resulting from the large magnetic fields created by the fault current flowing through them. Since the field will depend on the magnitude of the fault current, which in turn also depends on the voltage of the circuit, the

stresses that must be accommodated depend on both of these values. Their time-current

characteristics, however, are dependent on the protective relays associated with them and must be coordinated with those of down-line reclosers, fuses, and other protective devices.

Overcurrent Relays

Overcurrent relays close their contacts to actuate the circuit that causes the circuit breaker to open or close when the current flowing in them reaches a predetermined value.

Instantaneous. Without time delay deliberately added, the relay will close its contacts ‘instantaneously’, i.e., in a relatively short time, in the nature of 0.5 to perhaps 20 cycles. To prevent frequent operation of the breaker from transient, nonpersistent conditions, undesirably high settings may be applied to the relay.

Inverse Time. The operation of the relay may be made to vary approximately inversely with the magnitude of the current. The current setting may be varied and time delay introduced by varying the restraint on the movable element of the relay. Greater selectivity between relays and fuses in the circuit may thus be obtained.

Definite Time. A definite time delay can be introduced before the relay begins to operate, allowing greater selectivity to be achieved. This feature is often added to the inverse-time characteristic beyond a certain value of current after which the relay operation is completed after the fixed time delay. This inverse definite minimum time feature is employed in most Overcurrent relay applications.

The distribution circuit may be sectionalized with reclosers, automatic sectionalizers, and fuses, at which points faults may be isolated without affecting the entire circuit; fuses are also provided on the primary side of distribution transformers. The definite time characteristic of the relay associated with the circuit breakers at the substation is coordinated with the characteristics of reclosers and fuses on the distribution circuit.

Directional Relays

Directional relays are essentially Overcurrent relays to which an element similar to a wattmeter is added, both sets of contacts being in series. The Overcurrent element will operate to close its contacts regardless of the direction of flow of power in the line; the wattmeter element will tend to turn in one direction under normal flow of power and in the reverse direction

when power flows in the opposite direction. Hence, both sets of contacts must be closed and power flowing in a given direction before the relay will operate. Both elements may be combined into one so that only a single set of contacts is required.

This type of relay is used in primary or secondary network operations to open the protectors to prevent current from the network from energizing the high side of the transformers and their supply feeder during contingencies.

Differential Relays

Differential relays operate on the difference between the current entering the line or equipment being protected and the current leaving it. As long as the incoming current and the outgoing current are essentially equal, the relay will not operate. A fault within the line or equipment, however, will disturb this equilibrium, and the relay will operate to trip the supply circuit breaker or breakers on both sides of the line or equipment being protected. This type of relay is used to protect buses, transformers, and regulators at the substation. Since the voltages at which these operate may be high, current transformers installed on both sides of the equipment, with proper ratios in the case of transformers, supply the currents to the relay. Refer to Figure 5-4.

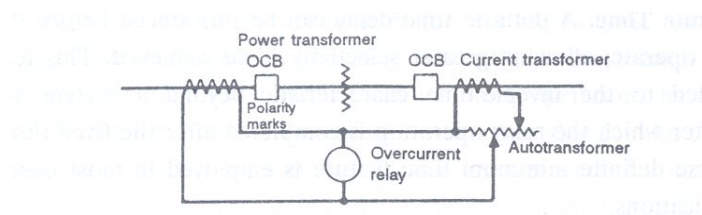


Figure 5-4. One-Line Diagram of Current-Differential Protection.

Comprehension Exercises

A. Choose a, b, c, or d which best completes each item.

1. An automatic line sectionalizer is
 - a. a self-controlled circuit opening device that locks out other protective devices and allows the circuit to be reenergized
 - b. a self-controlled circuit opening device that removes the fault and allows the circuit to be reenergized
 - c. employed in parallel with other sectionalizers to prevent deenergization of the circuit
 - d. employed in a circuit to open the reclosers and prevent deenergization of the circuit

2. Reclosers are designed so that they will after a selected sequence of tripping operations.

- a. not react to the fault
 - b. not act as fuses
 - c. remain open or locked out
 - d. allow the fault to pass through them
3. It is true that reclosers
- a. are set for three simultaneous operations
 - b. are set for indefinite automatic operations
 - c. can operate on only one time-current characteristic
 - d. can operate on one or more time-current characteristic curves
4. As we understand from the text,
- a. reclosers are not designed to withstand the forces resulting from severe faults
 - b. reclosers are designed to do the same function as circuit breakers in the circuit
 - c. circuit breakers are not able to withstand the stresses resulting from large magnetic fields
 - d. circuit breakers are independent from relays associated with them and other protective devices in the circuit
5. A definite-time relay indicates that
- a. a delay is purposely introduced in action regardless of the magnitude of the quantity that causes the action
 - b. a delay introduced in action causes the relay to operate inversely with the magnitude of the current
 - c. the relay may close its contacts instantaneously regardless of the addition of a time delay
 - d. the relay may be based on a time delay introduced by varying the restraint on its movable element
6. A differential relay is applied to the circuit
- a. to respond to the difference between the current entering the protected equipment and the current leaving it
 - b. to respond to the difference between the high voltage currents flowing through the protected apparatus
 - c. to stabilize the state of equilibrium between the current entering the protected line and the current leaving it
 - d. to stabilize the circuit breakers on both sides of the protected equipment with enough energy

B. Write the answers to the following questions.

1. How are automatic line sectionalizers employed in a circuit?
2. How are sectionalizers rated?
3. How does a sectionalizer close to a reclosing device function compared with one which is far from the reclosing device?
4. Where are reclosers used as repeater fuses?
5. How are reclosers compared with circuit breakers?
6. What is the function of an overcurrent relay?
7. What are the different types of overcurrent relays?
8. What part does a directional relay play in a distribution system?



Section Three: Translation Activities

A. Translate the following passage into Persian.

Surge or Lightning Arresters

The function of a surge or lightning arrester is to limit the voltage stresses on the insulation of the equipment being protected by permitting surges in voltage to drain to ground before damage occurs. The surges in voltage generally are caused by lightning (either by direct stroke or by induction from a nearby stroke) or by switching.

Arresters consist of two basic components: a spark gap and a nonlinear resistance element (for a valve type) or an expulsion chamber (for an expulsion type). When a surge occurs, the spark gap breaks down or sparks over, and permits current to flow through the resistance (or chamber) element to ground. Since the arrester at this point presents a low-impedance path, a large current, referred to as *60-cycle follow current*, flows through the arrester. The nonlinear resistance, at the higher voltages, will tend to restrict this current and eventually cause it to cease to flow; here, the magnitude of the follow current is independent of the system capacity. The expulsion chamber will confine the arc, build up pressures that eventually blow out the arc, and cause the follow current to cease to flow; here, the follow current is a function of the system capacity and the expulsion chamber must be suitably designed. After each such operation, the arrester must be capable of repeating this operating cycle.

74

B. Find the Persian equivalents of the following terms and expressions and write them in the spaces provided.

1.backup	
2. contingency	
3. deficiency	
4. elapse	
5. expulsion chamber	
6. fault-counting relay	
7. imminent	
8. instantaneous	
9. lighting arrester	
10. link	
11. permanent	
12.persist	
13. quench	
14.recloser	
15. repeater fuse	
16. restraint	
17. squelch	
18. tension	

Unit 6

Section One: Reading Comprehension

Relationship Between Voltage and Current

This is a long and interesting story. It is the heart of electronics. Crudely speaking, the name of the game is to make and use gadgets that have interesting and useful I versus V characteristics. Resistors (I simply proportional to V), capacitors (I proportional to rate of change of V), diodes (I only flows in one direction), thermistors (temperature-dependent resistor), photoresistors (light-dependent resistor), strain gauges (strain-dependent resistor), etc., are examples. We will gradually get into some of these exotic

devices; for now, we will start with the most mundane and most widely used circuit element, the resistor (Figure 6-1).



Figure 6-1.

It is an interesting fact that the current through a metallic conductor (or other partially conducting material) is proportional to the voltage across it. (In the case of wire conductors used in circuits, we usually choose a thick enough gauge of wire so that these 'voltage drops' will be negligible.) This is by no means a universal law for all objects. For instance, the current through a neon bulb is a highly nonlinear function of the applied voltage (it is zero up to a critical voltage, at which point it rises dramatically). The same goes for a variety of interesting special devices—diodes, transistors, light bulbs, etc.

A resistor is made out of some conducting stuff (carbon, or a thin metal or carbon film, or wire of poor conductivity), with a wire coming out each end. It is characterized by its resistance:

$$R = \frac{V}{I}$$

R is in ohms for V in volts and I in amps. This is known as Ohm's law. Typical resistors of the most frequently used type (carbon composition) come in values from 1 ohm (1Ω) to about 22 megohms ($22\text{ M}\Omega$). Resistors are also characterized by how much power they can safely dissipate (the most

commonly used ones are rated at $\frac{1}{4}$ or $\frac{1}{2}$ watt and by other parameters such as tolerance (accuracy), temperature coefficient, noise, voltage coefficient (the extent to which R depends on applied V), stability with time, inductance, etc. Capacitors (Figure 6-2) are devices that might be considered simply frequency-dependent resistors. They allow you to make frequency-dependent voltage dividers, for instance. For some applications (bypass, coupling) this is almost all you need to know, but for other applications (filtering, energy storage, resonant circuits) a deeper understanding is needed. For example, capacitors cannot dissipate power, even though current can flow through them, because the voltage and current are 90° out of phase.

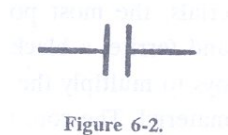


Figure 6-2.

A capacitor (the old-fashioned name was *condenser*) is a device that has two wires sticking out of it and has the property

$$Q = CV$$

A capacitor of C farads with V volts across its terminals will contain Q coulombs of stored charge.

Taking the derivative, you get

$$I = C \frac{dV}{dt}$$

Capacitors come in an amazing variety of shapes and sizes. The basic construction is simply two conductors near each other; in fact, the simplest capacitors are just that. For greater capacitance, you need more area and closer spacing; the usual approach is to plate some conductor onto a thin insulating material (called a dielectric), for instance, aluminized Mylar film rolled up into a small cylindrical configuration. Other popular types are thin ceramic wafers, metal foils with oxide insulators (electrolytics), and metallized mica. Each of these types has unique properties. In general, ceramic and Mylar types are used for most noncritical circuit applications; tantalum capacitors are used where greater capacitance is needed, and electrolytics are used for power supply filtering.

Inductors (Figure 6-3) are closely related to capacitors; the rate of current change in an inductor depends on the voltage applied across it, whereas the rate of voltage change in a capacitor depends on the current through it. The

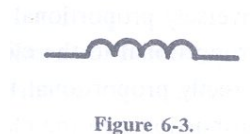


Figure 6-3.

defining equation for an inductor is

$$V = L \frac{dI}{dt}$$

where L is called the *inductance* and is measured in henrys (or mH, μ H, etc.). Putting a voltage across an inductor causes the current to rise as a ramp (for a capacitor, supplying a constant current causes the voltage to rise as a ramp).

The symbol for an inductor looks like a coil of wire; that is because, in its simplest form, that is all it is. Variations include coils wound on various core materials, the most popular being iron (or iron alloys, laminations, or powder) and ferrite, a black, nonconductive, brittle magnetic material. These are all ploys to multiply the inductance of a given coil by the ‘permeability’ of the core material. The core may be in the shape of a rod, a toroid (doughnut), or even more bizarre shapes, such as a ‘pot core’ (which has to be seen to be understood; the best description we can think of is a doughnut mold split in half, if doughnuts were made in molds).

Inductors find heavy use in radio frequency (RF) circuits, serving as RF ‘chokes’ and as parts of tuned circuits. A pair of closely coupled inductors form the interesting object known as a transformer.

Part I. Comprehension Exercises

A. Put T for true and "F" for false statements. Justify your answers.

- 1. The relationship between voltage and current is a crucial point in electronics.
- 2. Resistors and capacitors provide for useful current versus voltage.
- 3. Thermistors are light sensitive devices.
- 4. Capacitors are of different types.
- 5. Each capacitor has various applications.
- 6. In order to produce coils with different rates of inductance, various core materials are employed.

B. Choose a, b, c, or d which best completes each item.

- 1. Ohm’s law states that the current in a circuit is
 - a. inversely proportional to the resistance of the circuit and is directly proportional to the electromotive force in the circuit
 - b. directly proportional to the resistance of the circuit and is inversely proportional to the electromotive force in the circuit
 - c. directly proportional to the resistance and the electromotive force in the circuit

- d. inversely proportional to the resistance and the electromotive force in the circuit
2. It is true that resistors
- a. separate signals
 - b. generate waves
 - c. dissipate power
 - d. store energy
3. The current through a capacitor is
- a. independent of frequencies
 - b. proportional to the voltage
 - c. independent of the variations of voltage
 - d. proportional to the rate of change of voltage
4. It is true that
- a. resistors may be used in bypass applications
 - b. capacitors may be used for filtering
 - c. resistors are identical with condensers
 - d. capacitors are identical with thermistors
5. If you change the voltage across a farad by one volt per second, you are
- a. supplying an ampere
 - b. supplying a farad
 - c. increasing the voltage
 - d. increasing the current
6. We may deduce from the text that 1 volt across 1 henry produces a current that
- a. decreases at 1 amp per second
 - b. increases at 1 amp per second
 - c. is constant up to a critical point
 - d. is zero up to a critical point

C. Answer the following questions orally.

1. What is the function of a diode?
2. What is a resistor made up of?
3. How are capacitors basically made?
4. Which part of a capacitor is called a dielectric?
5. Why does aluminized Mylar film act as a capacitor?
6. How do you describe a pot core used in an inductor?
7. What are some applications of inductors?

Part II. Language Practic

A. Choose a, b, c, or d which best completes each item.

- 1-A voltage or current that varies at a constant rate is referred to as
- a. ramp
 - b. a rise
 - c. a drop
 - d. a tap

2. A is an electron device that makes use of the change of resistivity of a semiconductor with change in temperature.
 - a. thermocouple
 - b. thermoelement
 - c. thermostat
 - d. thermistor
3. A consists of two electrodes separated by a dielectric for introducing capacitance into an electric circuit.
 - a. capacitor
 - b. resistor
 - c. diode
 - d. strain gauge
4. A introduces relatively small insertion loss to waves in one or more frequency bands and relatively large insertion loss to waves of other frequencies.
 - a. bypass
 - b. filter
 - c. coupling
 - d. condenser
5. The property of an electric circuit by virtue of which a varying current induces an electromotive force in that circuit or in a neighboring circuit is called
 - a. conductance
 - b. capacitance
 - c. inductance
 - d. resistance

B. Fill in the blanks with the appropriate form of the words given.

1. Recognize

- a. Electrical signal such as voltages exist throughout a digital system in either one of two values and represent a binary variable equal to 1 or 0.
- b. Each digital logic family is by its basic NOR or NAND.

2. Rate

- a. Rated accuracy is the limit that errors will not exceed when an instrument is used under any combination of operating conditions.
- b. Rate-of-rise suppressors are devices used to control the of rise of current and/or voltage to the semiconductor devices in a semiconductor power converter.
- c. The of electric apparatus in general is expressed in volt- amperes, horsepower, kilowatts, or other appropriate units.

3. Resist

- a. A resistor introduces into an electric circuit.

- b. A as used in electric circuits for purposes of operation,

- protection, or control, commonly consists of an aggregation of units.
- c. When the resistivity of substance is known, the of any body composed of that substance can be calculated.
 - d. The resistance of a wire is directly proportional to the of the substance forming the wire.

4. Capacitor

- a. A parallel circuit consisting of an inductor in parallel with a is termed a parallel resonant or parallel tuned circuit when the resultant current taken from the supply is at its minimum value.
- b. A filter consists of an arrangement of resistors and inductive and elements.
- c. Capacitance current or component is a reversible component of the measured current on charge or discharge of the winding and is due to the geometrical , that is, the capacitance as measured with altering current of power or higher frequencies.

5. Depend

- a. The distinction between the two aspects of reliability, and security, is usually made when a communication channel is involved in the relay system and noise or extraneous signals are a potential hazard to the correct performance of the system.
- b. A dependent contact is a contacting member designed to complete any one of two or three circuits, on whether a two- or three-position device is considered.
- c. An operation solely by means of directly applied manual energy is referred to as manual operation.

C. Fill in the blanks with the following words.

inductor	resonant	tuned
parallel	capacitive	divided
composed	connected	circuit

A series tuned circuit is of a capacitor and an inductor in series. The frequency at which the and inductive reactances are equal is the frequency. The reactance value of the or capacitor at this frequency by the series resistance in the is the Q (quality factor). A Parallel circuit consists of an inductor in with a capacitor.

D. Put the following sentences in the right order to form a paragraph. Write the corresponding letters in the boxes provided

- a. Inductors usually consist of many adjacent turns of insulated wire, wound on a single support of laminated iron for low frequency inductors and ferrite for high frequencies.
- b. Resistors may be wire-wound to dissipate considerable heat, or they may be a thin film or a composition.
- c. At very high frequencies, air-core inductors are generally used.
- d. Resistors, inductors and capacitors are passive circuit components.
- e. Capacitors may be fixed, with solid dielectrics, or variable with usually an air dielectric.
- f. Rheostats and potentiometers are variable resistors with two and three terminals respectively.

1	2	3	4	5	6
					

Section Two: Further Reading

Passive Circuit Components-Resonant Circuits-Filters

A familiar electrical component in electronic circuits is the *resistor*. Resistors are circuit elements having specified values of resistance. Some resistors are made from a long, very fine wire wound on an insulating support; and these *wire-wound resistors* are generally used when it may be necessary to dissipate considerable heat.

Another common type of resistor is the *thin-film resistor*. This is made by depositing a thin film of metal on a cylindrical insulating support. High resistance values are a consequence of the thinness of the film.

A third type of resistor that has had very wide application is the *composition resistor*. In this type, the resistive element is a combination of finely divided carbon or graphite and a non-conducting inert material or filler such as talc, with synthetic resin used as a binder. These substances are

employed in such proportions as to give the correct resistance value in the finished

product.

The need often arises to vary the resistance of a resistor while it is permanently connected in a circuit. Components which enable this to be done are known as *variable resistors*, and their main element is a mechanical slider or arm which slides over the resistance element included in the circuit. A variable resistor with only two terminals is known as a *rheostat* while one with three is known as a *potentiometer*.

Electronic components with appreciable inductance are called *inductors*. They consist of many adjacent turns of wire wound on the same support. Large inductances for use at low frequencies are obtained by winding many hundreds of turns of wire on a core of a ferromagnetic material such as iron. When iron cores are used they are laminated in order to reduce eddy currents. Ferrite cores, made of high resistivity ferromagnetic materials, are used at high frequencies because their high resistance makes eddy currents negligible. They are not used at low frequencies, however, because their magnetic properties are not so favourable as those of iron. At very high frequencies, air-core inductors are employed. Despite the fact that variable inductances can be obtained by moving one portion of the winding with respect to the other, such components are not widely used.

The third major component we may consider is the *capacitor*. Capacitors may be variable or fixed. *Variable capacitors* most often use an air dielectric but sometimes the dielectric can be compressed gas or a liquid. Capacitance adjustment in variable capacitors is usually obtained by varying the effective plate area. Variable capacitors are also termed tuning capacitors and are generally used for varying the resonance frequency of a tuned circuit.

Fixed capacitors employ a wide range of dielectric materials and new types are continually finding valuable applications. Among the important types are those with solid dielectrics such as mica, plastic films, certain ceramics, paper and electrolytic films.

In a circuit composed of a capacitor and an inductor connected in series with a source of alternating current in which the frequency can be varied over range, at low frequencies the capacitive reactance of the circuit is large and the inductive reactance is small, while at high frequencies the inductive reactance is large and the capacitive reactance is small. Between these two extremes there is a frequency called the *resonant frequency* at which the capacitive and inductive reactances are exactly equal. As a result they completely cancel each other out and the current flow is determined wholly by the resistance of the circuit. Under these conditions the circuit is termed a

Series resonant circuit. At the resonant frequency the current has its largest Value, assuming the source voltage to be constant regardless of frequency.

The principle of resonance finds its most extensive application in radio frequency circuits, The value of the reactance of either the inductor or the capacitor at the resonant frequency of a series resonant circuit divided by the series resistance in the circuit (which is always present) is termed the *quality factor*, or more usually the *Q* factor or merely the *Q* of the circuit. In particular, for a given value of inductance, a circuit having a higher *Q* will have a smaller resistance and consequently will have a higher resonant current in comparison with the off-resonance current. Consequently (the curve of current versus frequency will be more sharply peaked and the circuit is said to be more sharply tuned.

A parallel circuit consisting of an inductor in parallel with a capacitor is termed a *parallel resonant or parallel tuned* circuit when the resultant current taken from the supply is at its minimum value. When a variable frequency source of constant voltage is applied to this parallel circuit there is a resonant effect similar to that in a series circuit, although in this case the source current is smallest at the frequency for which the inductive and capacitive reactances are equal.

A *filter* is a network that will permit the passage of electrical signals of a particular frequency or band of frequencies, while offering a much greater impedance to signals of higher or lower frequencies. Such a network may therefore be employed either to accept or reject signals of given frequencies. It consists of an arrangement of resistors and inductive and capacitive elements. There are three main types of filters; low-pass, high-pass, and band-pass. A low-pass filter is one that will permit all frequencies below a specified one-the cut-off frequency-to be transmitted with little or no loss, it will, however, attenuate all frequencies above the cut-off frequency. A high-pass filter is one with a cut-off frequency above which there is little or no loss in transmission, whereas below which there is considerable attenuation. A band-pass filter, on the other hand, will transmit a selected band of frequencies, attenuating all those either higher or lower than the desired band.

Comprehension Exercises

A. Choose s, b, c, or d which best completes each item.

1. The first three paragraphs mainly discuss wire-wound, thin-film, and composition resistors.
 - a. a kind of material identical in

- b. the rate of resistance offered by
 - c. a variety of materials used in
 - d. the structure and the application of

2. The resistance values of thin-film resistors vary inversely with of the film.
 - a. the thickness
 - b. the thinness
 - c. the length
 - d. the width
3. The fifth paragraph mainly discusses various
 - a. inductors and their deficiencies
 - b. inductors and their applications
 - c. core materials used in inductors
 - d. core materials used at low frequencies
4. To vary the resonant frequency of a tuned circuit are employed.
 - a. variable resistors
 - b. variable capacitors
 - c. fixed capacitors
 - d. composition resistors
5. The eighth paragraph mainly discusses
 - a. the capacitive reactance at low frequencies
 - b. the inductive reactance at high frequencies
 - c. the composition of a series resonant circuit
 - d. the application of a series resonant circuit
6. When a variable frequency source of constant voltage is applied to a parallel resonant circuit there is a resonant effect and
 - a. the source current is smallest at the frequency for which the inductive and capacitive reactances are equal
 - b. the source current is largest at the frequency for which the inductive and capacitive reactions are equal
 - c. the inductive and capacitive reactances cancel each other out and produce an impedance
 - d. the inductive and capacitive reactances cancel each other out and result in high resistance
7. The last paragraph mainly discusses
 - a. the cut-off frequency
 - b. the band-pass filter
 - c. the structure of a filter
 - d. the function of a filter

B. Write the answers to the following questions.

1. What is the application of wire-wound resistors?
2. What is a composition resistor composed of?
3. What is a variable resistor?
4. What is the difference between a rheostat and a potentiometer?

5. Why are ferrite cores for inductors not used at low frequencies?
6. Why are iron cores in inductors laminated?
7. What are some important types of dielectric materials that fixed capacitors use?
8. What is the advantage of the resonant frequency?
9. What is the Q of a circuit?
10. What are filter networks used for?



Section Three: Translation Activities

A. Translate the following passage into Persian.

Resistors

Resistors are truly ubiquitous. There are almost as many types as there are applications. Resistors are used in amplifiers as loads for active devices, in bias networks, and as feedback elements. In combination with capacitors they establish time constants and act as filters. They are used to set operating currents and signal levels. Resistors are used in power circuits to reduce volt-ages by dissipating power, to measure currents, and to discharge capacitors after power is removed. They are used in precision circuits to establish currents, to provide accurate voltage ratios, and to set precise gain values. In logic circuits they act as bus and line terminators and as 'pull-up' and 'pull-down' resistors. In high-voltage circuits they are used to measure voltages and to equalize leakage currents among diodes or capacitors connected in series. In radio frequency circuits they are even used as coil forms for inductors.

Resistors are available with resistances from 0.01 ohm through 10^{12} ohms standard power rating from $\frac{1}{8}$ watt through 250 watts, and accuracies from 0.005% through 20%. Resistors can be made from carbon-composition moldings, from metal films, from wire wound on a form, or from semiconductor elements similar to field-effect transistors (FETs). But, by far, the most familiar resistor is the $\frac{1}{4}$ or $\frac{1}{2}$ watt carbon-composition resistor. These are available in a standard set of values ranging from 1 ohm to 100 megohm, with twice as many values available for the 5% tolerance as for the 10% type. We prefer the Alien Bradley type AB ($\frac{1}{4}$ watt, 5%) resistor for general use because of its clear marking, secure lead seating, and stable properties.

B. Find the Persian equivalents of the following terms and

expression and write them in the space provided.

1. aggregate
2. attenuate
3. band-pass
4. binder
5. bypass
6. choke
7. coupling
8. cut-off frequency
9. dissipate
10. doughnut
11. electromotive force
12. field-effect transistor (FET)
13. inductance
14. laminate
15. parallel tune
16. passive
17. quality factor
18. rheostat
19. thermistor
20. tolerance
21. toroid
22. ubiquitous

Unit 7

Section One: Reading Comprehension

Important Features of a Computer

The development process of the computer over the last 150 years has resulted in all computers containing a number of fundamental features:

Stored Program Control. The computer program, which is a sequence of instructions executed one by one to perform the required data manipulation, must be stored within the computer. This has important advantages over external program storage.

Conditional Branching. One advantage of internal program storage is that the next instruction to be executed need not be the next in sequence since any instruction can be accessed as fast as any other (this is known as *random access*). The choice of which instruction to execute next can therefore be based upon the result of the previous operation or operations, giving the computer the ability to make decisions based upon the processing it performs.

Loops and Subroutines. The ability of a program to execute a particular set of instructions repetitively when required can produce enormous savings in the storage needed for the program. Conditional branches can be made to loop back and repeat a set of instructions a number of times, and commonly required subtasks within a program can be called up from any other part of the program as required, without needing to include the instructions of the subtask in the main program every time it is called.

Speed of Electronics. Even though the individual instructions available in a computer may be quite limited, because each instruction can be executed so fast, relatively powerful processing can be accomplished in what appears subjectively to be a very short time. (Compare this with the speed potential of Babbage's mechanical computer.)

Cost. The cost of computing power, and particularly the cost of computer memory is continually decreasing. It is now cheaper to store a computer instruction in an electronic memory than to store it on a card or a piece of paper tape.

Instructions Can Modify Themselves. Although this was one of von Neumann's original ideas embodied in the concept of stored program control, it has not been widely used since. One reason is that it is very difficult to keep

track of what the computer is doing once the computer program has been modified from that originally written by the programmer. In microprocessors, in particular, this concept is avoided, because the implementation of a microcontroller with a permanently fixed

program precludes any possibility of subsequent program changes.

A Simple Computer

Figure 7-1 shows the structure of a simple computer. The computer can be split into a number of separate components, though the components shown do not necessarily represent the physical division between components in a real computer. For example, the control unit and arithmetic and logic unit (ALU) are generally implemented as a single chip, the microprocessor, in microcomputers. Similarly, the input and output unit may be combined into a single chip in some microcomputers. Nevertheless, Figure 7-1 represents the conceptual structure of any computer from the smallest microcomputer to the largest mainframe computer.

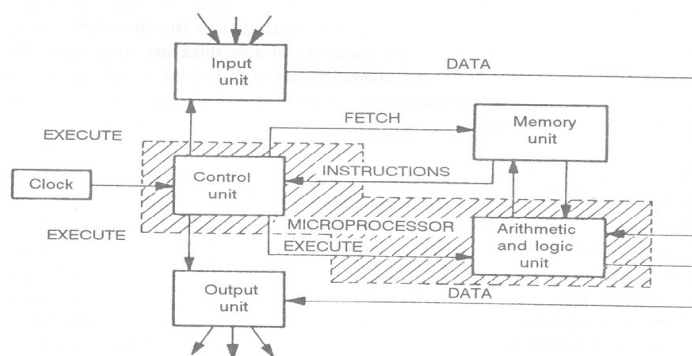


Figure 7-1. Conceptual Structure of a Computer.

The first requirement of any computer is a mechanism for manipulating data. This is provided by the ALU, which can perform such functions as adding or subtracting two numbers, performing logical operations, incrementing and decrementing numbers and left and right shift operations. From this very basic set of operations, more complex processing functions can be

generated by programming. Larger computers may provide additional more powerful instruction (for example, multiply and divide) within the computer instruction set.

Clearly, every computer must also include an input and an output unit. These provide the mechanism by which the computer communicates with the outside world. The outside world may consist of someone typing at a computer terminal and watching the response on a screen, or it may be some equipment, for example a washing machine, which is providing data inputs such as water temperature, water level and drum rotation speed, and is being controlled according to the program inside the computer, via computer outputs which switch on and off the water taps and heater, and alter the motor speed.

The computer must include an internal memory, which serves two functions. First, it provides storage for the computer program; second, it provides temporary storage for data which may be generated at some point during program execution by the ALU, but not be required until somewhat later. Such data *variables* must be able to be written into the memory unit by the computer, and subsequently read back when the data are required. The memory is organised as a one-dimensional array (or list) of *words*, and each instruction or data variable occupies one or more words in the memory. Each word is made up of a number of *bits* (binary digits) of storage in parallel.

The control unit of the computer controls the sequence of operations of all the components described above, according to the instructions in the computer program. Each instruction is fetched from the memory, and is then decoded by the control unit and converted into a set of lower-level control signals which cause the function specified by that instruction to be executed. When one instruction execution has been completed the next instruction is fetched and the process of decoding and executing the instruction is repeated. This process is repeated for every instruction in the program and only differs if a branch instruction is encountered. In this case, the next instruction to be fetched from the memory is taken from the part of the memory specified by the branch instruction, rather than being the next instruction in sequence.

The final component of the computer is a clock, or fixed-frequency oscillator, which synchronises the operation of all parts of the computer, and ensures that all operations occur in the correct sequence. The clock frequency defines the instruction execution speed of the computer and is constrained by the operating speed of the semiconductor circuits which make up the computer.

part I. Comprehension Exercises

A. Put "T" for true and "F" for false statements. Justify your answers.

- 1. At rundown access, the access time is independent of the physical location of the data.

- 2. Random access increases the speed of the computer.
- 3. At random access, the choice of which instruction to execute next depends on the result of the following operations.
- 4. Subsequent program changes is a new criterion.
- 5. Microprocessors are usually programmed so as to modify instructions.
- 6. The ALU is responsible for data manipulation.
- 7. The control unit selects, interprets, and sees to the execution of program instructions.

B. Choose a, b, c, or d which best completes each item.

1. The computer program needed to direct data manipulation
 - a. does not consist of sequentially arranged instructions
 - b. does not have the ability to execute instructions repeatedly
 - c. must be stored in the internal memory
 - d. must be stored in the external memory
2. The computer
 - a. has the ability to alter the sequence of program execution
 - b. lacks the ability of executing a sequence of instructions repeatedly
 - c. cannot call up subtasks from other parts of the program
 - d. cannot execute a subtask any time needed to do so
3. We understand from the text that the computer
 - a. can produce the instructions necessary to solve problems
 - b. can deviate from a top-down, structured design strategy
 - c. does not have direct access to the internal memory
 - d. does not make frequent uses of the internal memory
4. Modern computers are than their predecessors.
 - a. cheaper and less boring
 - b. cheaper and more powerful
 - c. faster but less effective
 - d. faster but more expensive
5. The clock is a regular time-keeping device used as a component of
 - a. the ALU
 - b. the memory
 - c. the arithmetic unit
 - d. the control unit

6. It is true that
 - a. the internal memory is directly under the control of the CPU
 - b. the internal memory is directly under the control of the input/output units
 - c. the ALU is responsible for choosing data from memory
 - d. the ALU is responsible for decoding the instructions fetched from memory

C. Answer the following questions orally.

1. What does a computer program consist of?
2. How do you describe random access?
3. Why is it cheaper to store instructions in an electronic memory than to store them on cards?
4. What are the basic components of a computer?
5. What is the function of the input/output units?
6. What comprises the outside world of the computer?

Part II. Language Practice

A. Choose a, b, c, or d which best completes each item.

1. By the computer is able to make logical decisions based upon the results of computation.
 - a. task specification
 - b. problem solving
 - c. program branching
 - d. program implementation
2. If knows a procedure that will solve the problem, the solution may then be coded in a selected language.
 - a. the programmer
 - b. the systems analyst
 - c. the operator
 - d. the systems designer
3. The main-control program specifies the order in which each in the program will be processed.
 - a. main task
 - b. major module
 - c. program design
 - d. subroutine module
4. The primary storage section can be designed to store a fixed number of characters or in each numbered address location.
 - a. words
 - b. records
 - c. files
 - d. bases
5. Primary memory is used transiently, which means that a program, or part of it, is kept in while the program is being executed.
 - a. the arithmetic-logic unit
 - b. the control unit
 - c. the internal storage
 - d. the external storage

B. Fill in the blanks with the appropriate form of the words

given.

1. Program

- a. The first step in implementation is to debug the program.
- b. A to be tested has generally demonstrated that it will run and produce results.
- c. Critical or lengthy operations that have been slowly carried out by software can be converted into microprograms and fused into read-only-memory chip.
- d. Applications programmers code modules that have been mapped out by the chief..... .

2. Communicate

- a. Transferring data from one location or operation to another, for use or for further processing, is data
- b. In a data system, workstations and other remote I/O devices are linked with one or more processors to capture input data and receive output information,
- c. A computer must be able to with the user.

3. Develop

- a. An effective approach in the programming analysis stage of program is to break down a large problem into a series of smaller tasks.
- b. The second generation of computers was in 1960.
- c. A microcomputer has into a most necessary information processing tool in the business today.

4. Process

- a. The heart of any computer system is the , which consists of primary storage, arithmetic-logic, and control elements.
- b. Data consists of three basic activities: capturing the input data, manipulating the data, and managing output results.
- c. Since processors of almost any size today can far more data in a second than a single set of I/O devices can supply or receive, it is common to overlap jobs.

5. Execute

- a. A load module which is the result of system routines linked with an object module is directly by the computer.
- b. The time necessary for a program is usually indicated on the computer print-out.

- c. Instructions and cycles are synchronized by a specific number of electric pulses produced by an electronic clock that is built in the processor.

C. Fill in the blanks with the following words.

appropriate operation move used
 necessary executed storage require

When a program instruction is to be, the control section first retrieves it from Next, the instruction is interpreted to determine the action. This could mean an arithmetic is needed-or a compare or a data or a branch. After this is determined, the part of the control section or ALU is to execute the instruction. Often this process will the use of registers.

**D. Put the following sentences in the right order to form a paragraph.
 Write the corresponding letters in the boxes provided.**

- a. Auxiliary storage devices generally provide *serial access* to data, in the same way as different pieces of music are stored serially on a cassette tape.
- b. A useful general-purpose computer normally has some type of *auxiliary* or back-up storage mechanism for long-term archiving of data or programs
- c. In order to access data from such devices, a large time penalty must be tolerated; as a result, such devices are not usually used when fast random access is required.
- d. The auxiliary storage can usually be completely detached from the computer, and is often some type of magnetic storage medium such as a tape or a floppy disk.

1	2	3	4
			

Primary Storage Components

Within any computer, both instructions and data are stored internally in memory as binary numbers. This is because the computer can only understand and execute instructions coded in the binary machine language code appropriate to that particular type of computer. It may appear to the user that the computer is executing a program written in some other programming language, such as BASIC or PASCAL. In fact, in order to execute a program written in any other language, the program must first be converted to the computers' own machine language. This can either be done once and for all using a *compiler*, or as the program is executed using an *interpreter*.

Each instruction in the machine code instruction set of the computer is normally represented by one or more *words* in the computer memory. Each word is represented physically by a number of binary digits (*bits*) in parallel. Data are also stored as words within the computer memory and the word-length of the computer therefore defines the number of binary digits which the computer can manipulate simultaneously, since arithmetic manipulations are generally performed upon one memory word at a time. The number of bits per word is chosen by the designer of the computer, and is one measure of the processing power of the computer. At the present time, microprocessors typically use 8, 16 or 32 bits per word, minicomputers 16-32 bits per word and mainframe computers 32-64 or more bits per word.

Computer memory is implemented in a number of different ways. Semiconductor memory is the dominant technology at present. **Semiconductor** storage elements are tiny integrated circuits. Both the storage cell circuits and the support circuitry needed for data writing and reading are packaged on chips of silicon. There are several semiconductor storage technologies currently in use. It is not necessary to consider the physics of these different approaches in any detail. It is enough just to mention that faster and more expensive *bipolar semiconductor* chips are often used in the arithmetic/logic and certain other sections of the processor while slower and less expensive chips that employ *metal-oxide semiconductor (MOS)* technology are usually used in the primary storage section. These primary storage components are often referred to as random-access memory (**RAM**) **chips** because any of the locations on a chip can be randomly selected and used to directly store and retrieve data and instructions.

RAM chips may be classified as dynamic and static. The storage cell circuits in **dynamic RAM chips** contain (1) a transistor that acts in much the same way as a mechanical on-off light switch and (2) a capacitor that is capable of storing an electric charge. Depending on the switching action of the transistor, the capacitor either contains no charge (0 bit) or does hold a charge (1 bit). Since the charge on the capacitor tends to 'leak off', provision is made to periodically 'regenerate' or refresh the storage charge. A dynamic RAM chip thus provides volatile storage; that is, the data stored are lost in the event of a power failure.

Static RAM chips are also volatile storage devices, but as long as they are supplied with power, they need no special regenerator circuits to retain the stored data. Since it takes more transistors and other devices to store a bit in a static RAM, these chips are more complicated and take up more space for a given storage capacity than do dynamic RAMs. Static RAMs are thus used in specialized applications, while dynamic RAMs are used in the primary storage sections of most computers. Because of the volatile nature of these storage elements, a backup **uninterruptible power system (UPS)** is often found in larger computer installations. Personal computer users can also invest a few hundred dollars and get a small battery-powered UPS. This device supplies current for a period long enough for users to save data on a disk and then shut down system in an orderly way.

Specialized Storage Elements in the Processor Unit

You know that every processor has a primary storage section that holds the active program(s) and data being processed. In addition to this *general-purpose* storage section, however, many processors also have built-in *specialized* storage elements that are used for specific processing and control purposes.

One element used during *processing* operations is a **high-speed buffer (or cache)** memory that is both faster and more expensive per character stored than primary storage. This high-speed circuitry is used as a 'scratch pad' to temporarily store data and instructions that are likely to be retrieved many times during processing. Processing speed can thus be improved. Data may be transferred automatically between the buffer and primary storage so that application programmers are unaware of its use. Once found only in larger systems, cache memory is now available in some of the tiny microprocessor chips used in personal computers.

Other specialized storage elements found in many processors are used for *control* purposes. The most basic computer functions are carried out by wired circuits. Additional circuits may then be used to combine these very

basic functions into somewhat higher-level operations (to subtract **values**, move data, etc.). But it is also possible to perform these same higher-level operations with a

series of special programs. These programs-called **microprograms** because they deal with low-level machine functions-are thus essentially substitutes for additional hardware.

Microprograms are typically held in the processor unit in special control storage elements called read-only memory (**ROM**) **chips**. Unlike RAM chips, which are volatile, ROM chips retain stored data when the power goes off. Microprogram control instructions that cause the machine to perform certain operations can be repeatedly read from a ROM chip as needed, but the chip will not accept any input data or instructions from computer users.

The most basic type of ROM chip is supplied by the computer manufacturer as part of the computer system, and it cannot be changed or altered by users. Such chips have found wide application as a program storage medium in video games and personal computers. Of course, it is possible for a user to 'customize' a system by choosing the machine functions that will be performed by microprograms and by then using a second type of ROM chip. For example, critical or lengthy operations that have been slowly carried out by software can be converted into microprograms and fused into a programmable read-only memory (**PROM**) **chip**. Once they are in a hardware form, these tasks can usually be executed in a fraction of the time previously required.

PROM chips are supplied by computer manufacturers and custom ROM vendors. Once operations have been written into a PROM chip, they are permanent and cannot be altered. There are other types of ROM control chips available, however, that *can* be erased and reprogramed. Since one type of erasable and programmable read-only memory (**EPROM**) **chip** needs to be removed from the processor and exposed for some time to ultraviolet light before it can accept new contents, it is hardly suitable for use by application programmers. Another type of electrically erasable programmable read-only memory (**EEPROM**) **chip** is also available that can be reprogramed with special electric pulses. Regardless of the type of ROM chip used, however, they all serve to increase processor efficiency by controlling the performance of a few specialized tasks.

Comprehension Exercises

A. Choose a, b, c, or d which best completes each item.

- 1-The first paragraph mainly discusses
- a. the language of the computer
 - b. the language of the user
 - c. the use of a compiler
 - d. the use of an interpreter

2. As we understand from the text,
 - a. the user does not have access to the arithmetic/logic memory
 - b. the user does not have access to the main memory
 - c. bipolar semiconductor chips cannot be used in the primary storage
 - d. metal-oxide semiconductor chips cannot be used in the ALU
3. The third paragraph mainly describes
 - a. random-access memory chips used to store and retrieve data and instructions when necessary
 - b. semiconductor storage elements including the types used in the primary storage and those used in other sections of the processor
 - c. modern computers with semiconductor storage elements used in their primary storage section
 - d. bipolar semiconductor chips used in the arithmetic/logic and other sections of the computer
4. The fourth paragraph mainly describes
 - a. what dynamic RAM chips contain
 - b. how dynamic RAM chips are refreshed
 - c. the volatility of dynamic RAM chips
 - d. the characteristics of dynamic RAM chips
5. An uninterruptible power system
 - a. acts as a transistor being either on or off
 - b. acts as a capacitor storing either 0s or 1s
 - c. is used as a backup to compensate for the volatility of RAM chips
 - d. is used as a backup to personal computers to raise their storage capacity
6. It is true that
 - a. the most developed processors have either specialized or general-purpose storage elements
 - b. the most basic computer functions are carried out by either high-speed circuitry or by special regenerator circuits
 - c. computer users have no control over specialized storage elements
 - d. computer users are conscious of the operations performed by cache memory
7. As we understand from the text,
 - a. a microprogram does not efficiently perform high-level operations
 - b. a microprogram is a sequence of elementary instructions residing in the processor
 - c. ROM chips are interchangeable with static RAMs
 - d. EPROM chips cannot be reprogrammed by application programmers

B. Write the answers to the following questions.

1. What is the function of a compiler?
2. What is a word ?
3. Who decides on the number of bits per word?
4. What does RAM stand for?
5. What is the major advantage of random-access memory chips?
6. What are the similarities and differences between the dynamic and static RAM chips?
7. What is the function of cache memory?
8. What are microprograms used for?
9. How do ROM chips differ from RAM chips?
10. How does the computer user customize a system?



Section Three: Translation Activities

A. Translate the following passage into Persian.

Organizing Data for Processing

Data is the raw material to be processed by a computer. Such material can be letters, numbers, or facts-such as grades in a class, baseball batting averages, or light and dark areas in a photograph. Processed data becomes **information**-data that is organized, meaningful, and useful.

To be processed by the computer, raw data must be organized into characters, fields, records, files, and data bases.

A character is a letter, number, or special character (such as \$, ?, or *). One or more characters comprise a field.

A field contains an item of data. For example, suppose a sports club was making address labels for direct mailing. For each person, it might have a date-joined field, a name field, a street address field, a city field, a state field, and a postal code field.

A record is a collection of related fields. Thus, on the sports club list, one person's date-joined, name, address, city, state, and postal code would comprise a record.

A file is a collection of related records. The entire list of address labels for the sports club would be a file.

A data base is a collection of interrelated data stored together with

Minimum redundancy. Specific data items can be retrieved for various applications. For instance, the sports club data could be obtained according to state or postal code or alphabetically by last name.

Whenever a change is to be made to stored data, a record is generated containing the new data. The record is called a **transaction**. Whenever files are changed to reflect new information, the process is called **updating**. Files of records are stored on some form of medium, usually magnetic disk or magnetic tape, so they can be read into main computer storage for processing.

A file can be a **transaction file**, one that contains modifications to existing records. For example, in our address label list, a transaction would be a change in a label (a new address), an added label (a new member), or a deleted label (a member resigns). Or a file can be a **master file**, which contains relatively permanent data-the master address label list, in this case-that is updated by a transaction file.

B. Find the Persian equivalents of the following terms and expressions and write them in the spaces provided.

- | | |
|-------------------------------|-------|
| 1. cache | |
| 2. compiler | |
| 3. customize | |
| 4. encounter | |
| 5. EPROM | |
| 6. execute | |
| 7. fetch | |
| 8. interpreter | |
| 9. leak | |
| 10. mainframe | |
| 11. manipulate | |
| 12. metal-oxide semiconductor | |
| 13. microprogram | |
| 14. MOS | |
| 15. PROM | |
| 16. random access | |
| 17. redundance | |
| 18. static RAM | |
| 19. transaction | |
| 20. ultraviolet | |
| 21. vendor | |
| 22. volatile storage | |

Unit8

Section One: Reading Comprehension

Data Processing

The arithmetic and logic unit (ALU) of a computer may be thought of as the heart of a computer since this is the component which performs the data manipulation operations which are essential in any data processing task. Similarly the control unit may be likened to the brain of the computer since its function is to control the program execution according to the sequence of instructions encoded as the computer program. The control unit does this by *fetching* the instructions one by one from the memory and then *obeying* the data manipulation which each instruction specifies. Thus each instruction execution may be seen to require a two-beat cycle, where the first beat always fetches the instruction and the second beat executes the function specified by the instruction. The fetch operation is invariant for all instructions, but the execute operation varies according to the instruction fetched and may, for example, cause the addition of the contents of two registers, the clearing of a memory location or the transfer of data between a register and a memory location. The computer central processing unit or microprocessor unit must therefore perform three tasks, each of which is examined in more detail below:

CPU Communication With Memory

Data are communicated between the CPU and the memory of a computer using *buses*; similarly, communication between the registers in a computer CPU occurs via buses internal to the CPU. Conceptually there is no difference between these two types of bus, but in microprocessors the *external buses* connecting to memory are brought out to the pins of the microprocessor while the *internal buses* exist only on the silicon chip, and are inaccessible externally.

In order for the CPU to be able to access instructions or data stored in the main memory, the CPU must first of all supply the address of the required memory word using an *address bus*. When this address has been specified, the required instruction or data can be read into the CPU using the *data bus*.

In the case of data it may also be necessary to *write* data into the

memory from time to time as well as to *read* data *from* the memory. One way to do this would be to have two separate data buses, one for reading data from the memory into the CPU, and one to write data back into memory. To minimize the number of pins required on a microprocessor CPU chip, however, it is common to multiplex the data read and data write buses into a single bus with a separate *read/write control line* which specifies the direction of data transfer.

A further requirement of the CPU is a register to keep track of which instruction is being executed, so that instructions can be executed in sequence. This register is called the *program counter* (also sometimes known as the instruction counter), and acts as a pointer to the next instruction to be executed. It is incremented immediately after each instruction is accessed from memory, and also after accessing every instruction operand or operand address if the operand is stored in a separate word of memory. Each time that it is required to read the next instruction word from memory, the output of the program counter is connected via the address bus to the memory, thus supplying the address of the next instruction. After allowing an appropriate time for propagation delays the instruction word can be read into the CPU from the data bus.

Another register called the accumulator is also provided for short-term data storage in the CPU.

Instruction Execution-Data Manipulation

If it were only possible to use a computer to transfer data backwards and forwards between registers and memory, the computer would be of very little practical use; its power lies in its ability to manipulate and process information so that new information can be generated from data which already exists. The arithmetic and logic unit (ALU) is the component of a computer system which performs the task of data manipulation. Generally it is adequate to consider the ALU as a black box with two one-word-wide data inputs, and one one-word-wide data output. In addition, the ALU contains a number of control inputs which specify the data manipulation function to be performed. The ALU is a combinational logic circuit, whose output is an instantaneous function of its data and control inputs; it has no storage capability. Thus the result of any ALU operation must be stored in an accumulator.

Instruction Interpretation and Control

A final aspect of the CPU to be considered in this unit is the mechanism for

converting the instruction opcodes into the low-level control signals which cause the operation specified by the opcode to be executed. This is the function of the *control unit* of the computer. Because the instruction sets of different types of computer differ, so, too, do the control units differ in the way that they decode the operation codes to elemental control signals for the rest of the computer. Nevertheless, some observations may be made about the general structure of computer control units.

When the instruction opcode is loaded into the CPU from the computer's memory, it is stored in the *instruction register*. This register provides temporary storage for the opcode while it is applied to the inputs of an *instruction decoder*, which converts the opcode to the required low level control signals. The instruction decoder is a combinational logic component, and may be implemented using random logic, a PLA, or even using a ROM, though this tends to be rather inefficient in this application.

The *combinational logic control signals* are connected to all the combinational components in the computer, such as the ALU function and mode select inputs, and multiplexer and demultiplexer control inputs. The timing of the signals applied to these units is unimportant, so long as sufficient time is allowed for data and control signals to propagate through all components.

The *register control* outputs of the instruction decoder are connected to the write (latch data) inputs of all the CPU registers, and to the memory read/write line of the control bus. The timing of these signals is important because it controls when data are written into the CPU registers or memory.

Part I. Comprehension Exercises

A. Put "T" for true and "F" for false statements. Justify your answers.

-1. The control unit directs and coordinates the computer system in executing stored program instructions.
-2. The control unit applies a separate fetch operation to each instruction.
-3. The execute operation performed by the control unit is constant for all instructions.
-4. The data read and the data write buses are multiplexed in order to reduce the number of pins required on a microprocessor.
-5. As we understand from the text, data and information are the same.
-6. The instruction decoder interprets the opcode received.

B. Choose a, b, c, or d which best completes each item.

1. The first paragraph mainly
 - a. introduces the control unit as the most active component of a computer
 - b. introduces the ALU and the control units as the most important features of a computer
 - c. describes the mechanism of the ALU
 - d. describes the fetch and the execute operations
2. In order for the CPU to read data from or to write data into memory
 - a. the registers in the CPU must supply the required data
 - b. the registers in the CPU must be cleared of data
 - c. data and address buses must be used respectively
 - d. address and data buses must be used respectively
3. The fifth paragraph mainly discusses
 - a. the mechanism of the program counter
 - b. the mechanism of the central processing unit
 - c. the order of instructions stored in memory
 - d. the order of instructions fetched from memory
4. Data manipulation is performed in
 - a. the control unit
 - b. the arithmetic-logic unit
 - c. the main memory
 - d. the accumulator
5. It is true that
 - a. the combinational logic control signals are connected to certain combinational components in the computer
 - b. the combinational logic control signals are time dependent because they control data written into the CPU registers
 - c. the instruction register provides temporary storage for the operation part of an instruction received from memory
 - d. the instruction decoder receives the operation part of the instruction interpreted in the instruction register

C. Answer the following questions orally.

1. What is the function of the ALU?
2. Why is a two-beat cycle needed for each instruction execution?
3. How does the CPU communicate with memory?
4. What are the data and the address buses?
5. What is the function of a read/write control line?
6. What is a program counter?

A. Choose a, b, c, or d which best completes each item.

- B. Fill in the blanks with the appropriate form of the words given.**

- a.** A is an electronic symbol-manipulating system designed and organized to automatically accept and store input data, process them, and produce output results under the direction of a detailed step by step stored program of instructions.
- b.** Addition, subtraction, multiplication, and division are that the computer can perform. ,
- c.** The computer component that performs the mathematical operations required for problem solving is called the element.

a. Symbolic addressing is the practice of expressing an, not in

terms. of its absolute numeric location, bur rather in terms of symbols convenient to the programmer.

b. The computer memory contains locations.

c. A 16-line bus is built into a microprocessor chip to determine the primary storage locations of the needed instructions and data.

d. To improve the data handling and capabilities of their products, microprocessor suppliers introduced improved chips in the early 1980s. !

3 .Differ

a. There is not a very big in flowcharting for a program to be written in Cobol or Fortran.

b. There are manycomputers manufactured today, and a buyer must be able to compare the advantages and disadvantages of each.

c. The opinions of programmers as to the best way of solving a problem often greatly.

4. Logic

a. A program must be..... organized if successful results are to be expected.

b. There are three basic kinds ofoperations: equal to, less than, and greater than.

5. Store

a. In a personal computer system, operating-system programs are commonly..... on a floppy disk.

b. Generally, the larger the system, the greater its ...capacity is.

c. Small magnetic tapes and floppy disks are typically used for off-line secondary

d. One of the basic types of processing a computer can do is the of data.

C. Fill in the blanks with the following words.

Contributed computers application important electronic
mainframe minicomputersmicrocomputers
Programmable inclusion rather low

Although microcomputers have the general characteristics of digital, a notable property of microcomputers is their relativelycost and small size. This has greatlyto their popularity and success. While largecomputers and minicomputers have more computational power than

..... , this power is not always needed in every Furthermore, the cost of large computers and..... has often precluded their adoption into certain..... systems which would have otherwise benefited from theirThe microprocessor has now made it possible to use a/andevice in a logic system where cost constraintsthan speed and computational power are

D. Put the following sentences in the right order to form a paragraph. Write the corresponding letters in the boxes provided.

- a. This was a major advance in the electronics industry, but it only paved the way for greater things: the evolution of the many analog circuits and the more complex digital circuits which contained several transistors and other components.
- b. The number of components which can be stacked on one chip determines the complexity of the circuit.
- c. Because these MSI. circuits are more reliable, cost less, and use .less power than the transistor equivalents, they are easier and cheaper to use.
- d. The microprocessor evolved from the transistor, which was the first major semiconductor device.
- e. Techniques were developed to put more and more complex circuits on the same chip, evolving into medium scale integration (MSI).
- f. Several transistors were then put on one semiconductor substrate, and the integrated circuit (IC) evolved.

1	2	3	4	5	6



Section Two: Further ReadinK

System Software

A vital part of any general-purpose computer is the system software, or software tools which are used in conjunction with the computer hardware.

Without the system software the computer is rather like a car without petrol; although the basic mechanics of the system exist, there is no way of actually using it. This unit concentrates on the software tools which are required to turn a general-purpose computer into a useful computer system for applications programming, for microprocessor applications development, or for use as a business or administrative system. Some insight into the types of programs required can be gained by considering the various tasks which might be undertaken using a general-purpose computer.

Loader

The first requirement when a computer is switched on is some kind of *loader* program which can be used to load any other program from the backing storage medium into memory prior to execution. In most modern computers the loader would be a rather large program stored in ROM and designed to read programs from a disk, as shown in Figure 8-1.

The concept of using one piece of software to enable another (generally more complex) piece of software to be run is called *bootstrapping*. Thus the loader program provides a method of bootstrapping other more complex programs. This approach is taken to minimize the use of the computer system memory; the loader program which resides permanently in ROM typically occupies only about 1 kbyte of memory, and therefore frees most of the

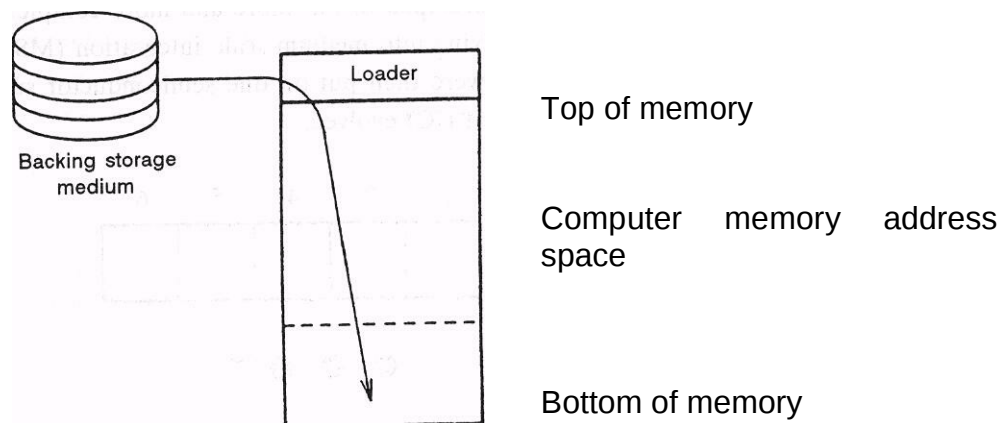


Figure 8-1. Operation of a Loader Program.

remaining memory for other programs. The loader program is also often combined with a simple *monitor* or *debugger* program which may be used to debug machine code programs, and also to verify operation of the computer

hardware without requiring access to a backing storage medium.

A different type of computer may also be used for bootstrapping, by making use of *cross-software*. For example, a program may be written for a microprocessor such as MC6809 on a minicomputer or mainframe computer, using an assembler written in a language which is available on the large computer, such as PASCAL, C, or FORTRAN. The assembler program which runs on the minor mainframe computer is known as a *cross-assembler* because it produces machine code for the microcomputer and not for the computer on which the assembler is run. The machine code program can then be transferred to the target computer either in PROMs or via any of the other storage media, or it can be downloaded directly via a serial or parallel communication link.

Disk Operating System

Software which enables programs to be loaded from (and stored on) backing storage media can be combined with facilities which handle the display terminal(s) and other computer peripherals such as printers, plotters, and so on, to provide a general-purpose control program. This *operating system* program generally makes use of floppy or hard disks as the main backing storage medium, because of the ease of randomly accessing different areas of disk containing different files. Hence the program is known as a *disk operating system* (DOS).

A disk operating system provides two fundamental facilities. First, it provides a mechanism for communication between the computer and the user by handling input and output to the user's console and executing the commands specified by the user. Second, it provides a mechanism for program storage and retrieval, though generally in a rather more flexible and sophisticated form than the simple loader program. Through the medium of the operating system users are able to access source files and call up text editors to make changes to them. They can then assemble or compile the source code (as appropriate) to produce machine code programs, and the machine code can be loaded into memory and executed, all under the control of the operating system.

Files, which may contain machine code, ASCII coded source text, data or any other information, are stored on the disk in a format defined by the operating system, and are accessed via *filenames*. The filenames, which are simply mnemonics chosen by the user to reflect the contents of a file, are stored in a directory which can be examined by the user. Manipulation of files

can then be performed simply by referring to the file using its filename. In a similar way, all the other system programs described in this unit are accessed through an operating system.

The operating system also provides a number of housekeeping utility programs which are necessary so that the user can maintain his files on the disk in an orderly fashion. Typically, the operating system includes the following:

- (a) *Directory listing* so that the user can determine what files are on the disk, their size, when they were created and other useful information.

- (b) *File erasing* so that unwanted files may be removed from the disk to free space for other purposes.

- (c) *File renaming* so that filenames may be changed if required.

- (d) *File transfer* so that files may be copied to another disk for backup purposes or for duplication.

- (e) *File listing* so that the contents of text files may be printed.

- (f) *File execution* so that machine code files may be loaded into memory and executed.

In addition, operating systems may include many other more advanced and specialized facilities which are beyond the scope of this discussion.

Finally, the operating system provides a straightforward mechanism to enable the user to access terminals, printers and storage devices connected to the computer from within his application programs. Access is achieved by executing subroutine calls to standard subroutines within the operating system which control these devices. This saves the user from needing a detailed knowledge of the interface characteristics of the peripheral devices connected to the computer if he simply wishes to run application programs which access the standard computer peripherals under the control of the operating system.

Text Editor

Given a loader for bootstrapping the operating system when the computer is switched on, and the ability to manipulate files stored on a disk storage system, the user's next requirement is generally a facility for developing application programs in a high-level language or assembly language. In either case, the application program is initially written as a source code text file using a *text editor*. The editor program must therefore provide a mechanism for inputting the source code program and saving it on disk, and subsequently facilities must also be available for making changes, corrections and additions to the source program.

Three types of activity are involved in using a text editor. First, the user must identify the text which is to be manipulated, or the place at which text is to be inserted in the file. This requires a *pointing* operation to specify where the change is to be made. Second, the user must specify the *operation* or *command* which is to be performed. Typical editing operations will include insertion of new text, deletion of text, and substitution of new text for old. In each case, the change may be made on a character, word, or line basis. More sophisticated editors also include commands to search a complete file for occurrences of a specified character string, and optionally replace this with a new string. They may also provide facilities for merging separate files or subdividing files. Finally, having specified the required editor command, the user must type in the replacement text or new text if the command requires it.

Computer text editors can be categorized into two broad types: *line-based* and *screen-based*. Screen-based editors are the simplest to use because the display terminal represents a *window* on the text file, and the pointing operation is achieved by using cursor keys on the terminal to position the cursor where a change is to be made. If the user wishes to move outside the range of the current display window, the screen is scrolled up or down to position the window at the required place in the file. Commands are generally displayed using a simple menu positioned on a fixed 'status' area of the screen, and selected by typing the first letter of the command or using special command function keys. Thus the text editor is simple and straightforward to use, and the effect of executing a command is immediately obvious on the screen.

If a screen-based editor cannot be used, the alternative is to use a line-based editor. All operations in this case are referred to a specified line in the source file. The line may be specified using a line number, or it may be identified using an *invisible cursor*. In this latter case, commands are executed to move the invisible cursor through the file, and the current cursor position within the file can be determined by causing the editor to list the current line on the terminal. Line-based editors generally require much greater memory on the part of the user and hence require significant training before competence is achieved.

Assemblers and Compilers

Once a source file has been prepared using a text editor, the next stage is to convert the program to machine code using an assembler (if the program is written in assembly language) or a compiler (if in a high-level language). For

short programs, the facilities introduced so far prove quite adequate; as larger and more complex programs are written, however, some disadvantages begin to become apparent.

As the size of the source program increases, the time required to assemble or compile the program also increases as does the editing time (since it takes longer to locate the required part of the program within the file). Thus the efficiency of the programmer begins to fall. One way to resolve this problem is to split the program into a number of separate source modules, each of which can be edited and assembled or compiled separately. A mechanism is then required to build the various program segments together into a final machine code program. Another related problem also may become apparent; this is a need to be able to segment and separate different parts of a program physically in memory. As an example, in developing a microprocessor application, it may be required to specify one memory area for program and data constants which will be stored in ROM, and a separate non-contiguous area for data variables stored in RAM. Furthermore, at the time the program is written, the final allocation of memory may not be known. Thus a simple mechanism for *relocating* the program to different memory addresses is needed.

Debuggers and Simulators

Once a program has been written and assembled or compiled, the next stage is to run the program and verify its operation. Even programs written by experienced and expert programmers very seldom work correctly to begin with. It is very difficult to foresee all the ways in which a program algorithm is executed and invariably algorithm faults are shown up when the program is first tested. To cater for this, debuggers and simulators are required to test a program.

If the program is written in the computer's native assembly language, then it is often possible to load the machine code program into the computer's memory and execute it under the control of a *monitor* program of the type already discussed. If, however, the program has been written using a cross-assembler, then it cannot be executed directly, but instead a *simulator* software package is often provided which interprets each target machine code instruction and modifies the memory locations in the host computer to simulate the action of the program on the target computer registers, memory and input/output circuits. If the configuration of the application computer is very different from that of development computer, it may not be possible fully

to debug the software within the development system. In this case, in-circuit emulation facilities are used to complete the debugging process.

Comprehension Exercises

A. Choose a, b, c, or d which best completes each item.

1. A loader program is designed
 - a. to read the programs executed in the main memory into the secondary storage
 - b. to read the program to be executed from the secondary storage into the main memory
 - c. to bootstrap pieces of simpler programs in order to back up the computer memory
 - d. to bootstrap more complex programs in order to verify their validity
2. A program written for a microprocessor may be assembled on a mini or mainframe and then transferred to the target microprocessor by
 - a. a loader
 - b. a monitor
 - c. a text editor
 - d. a cross-assembler
3. As we understand from the text,.....
 - a. operating system programs are commonly stored on disks
 - b. tapes are easier to use as the backing storage medium
 - c. a disk operating system is not of much help to the user
 - d. a disk operating system is not as efficient as a loader
4. An operating system enables the users
 - a. to communicate with the computer
 - b. to make changes to the source files
 - c. to have access to all the system programs
 - d. all of the above
5. It is true that
 - a. the user can refer to files via filenames
 - b. the user cannot manipulate the source files
 - c. the operating system makes the system more difficult to use
 - d. the operating system does not control the overall operations of the computer
6. The text editor
 - a. controls the execution of other programs
 - b. allows a user to enter and store programs
 - c. translates each source language statement into a sequence of machine instructions

d. inputs the source code programs, saves it on disk, and reviews and alters its materials

7. Using a text editor, the user must

- a. identify the text to be manipulated
- b. specify the operation to be performed
- c. type in the new text if required
- d. all of the above

8. According to the text,

- a. a window is a portion of the visual display screen used to show the current status of an application of interest
- b. a cursor is positioned where a change is to be made when a line- based editor is used
- c. line-based editors require smaller memory compared with screen- based editors
- d. screen-based editors are hard to use since the user cannot move outside the range of the current display window

9. It may be inferred from the text that

- a. screen-based editing is not necessarily carried out under the user's control
- b. screen-based editing requires a high communication rate between the console and the computer
- c. the concept of a cursor seems meaningless on a display terminal
- d. the concept of a cursor is only meaningful when a change is made on a character

10. An application program, once prepared, has to be

- a. converted into a high-level language
- b. segmented physically in memory
- c. split into a number of modules
- d. assembled or compiled .

B. Write the answers to the following questions.

- 1. What is bootstrapping?
- 2. Why does bootstrapping minimize the use of the computer system memory?
- 3. What is a directory?
- 4. What does an operating system include?
- 5. What is a filename?
- 6. Why does the user not need to know enough about the interface characteristics of the peripheral devices connected to the computer?

7. What is a pointing operation used for?
8. What do editing operations include?
9. What problems arise as the size of the source program increases?
10. What are debuggers and simulators used for?



Section Three: Translation Activities

A. Translate the following passage into Persian.

The Microprocessor

One of the results of the advancement in solid-state technology is the capability of fabricating very large numbers of transistors (say, 1000 and over) within a single silicon chip. This is known as *large-scale integration*. A direct consequence of large-scale integration is the microprocessor. In general, a *microprocessor* is a programmable logic device fabricated according to the concept of large-scale integration. A microprocessor has a large degree of flexibility built into it. By itself it cannot perform a given task, but must be programmed and connected to a set of additional system devices. These additional system devices usually include memory elements and input/output devices. In general, a set of system devices, including the microprocessor, memory, and input/output elements, interconnected for the purpose of performing some well-defined function, is known as a *microcomputer* or *microprocessor system*.

B. Find the Persian equivalents of the following terms and expressions and write them in the spaces provided.

1. administrate
2. bootstrap
3. console
4. cross-assembler
5. cross-software
6. debugger
7. file erasing
8. invisible cursor
9. line-based editors

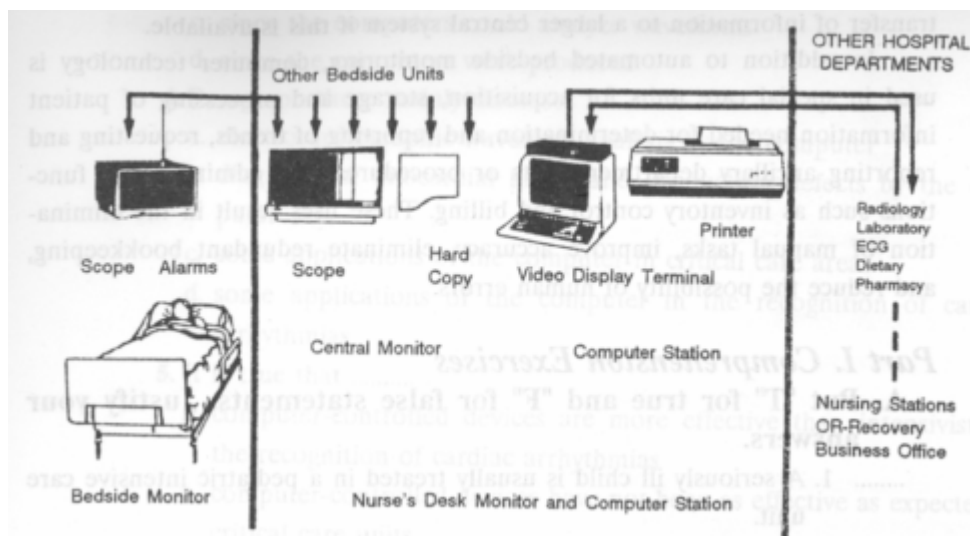
Section One: Reading Comprehension

Computerized Critical Care Areas

A critical care unit is an area in a hospital where highly trained personnel and sophisticated equipment are concentrated to take care of a limited number of actually or potentially severely ill patients. Special units may be known as general intensive care units or labeled according to the type of patients treated. Thus we have medical, surgical, neurological, respiratory, or pediatric intensive care units. Coronary care units, recovery rooms, telemetry monitoring areas, burn and trauma units are also critical care areas. Physicians trained to work in these units are commonly referred to as 'intensivists'. Likewise, nursing personnel permanently assigned to these units undergo specialized training.

Computers are now commonplace in critical care areas in both large and small hospitals. They cover a wide range of applications, from the micro-processor that controls specialized bedside and nurses' desk monitoring equipment to the mini- or mainframe computer that is part of either a dedicated critical care system or an integrated overall hospital-wide information facility. Figure 10-1 represents a commonly seen special care unit arrangement, consisting of (a) microprocessor-controlled patient monitoring hardware with bedside and nursing-desk scopes and controls, as well as hard-copy functions, and (b) remote computer stations, usually video display terminals (VDT) communicating through a central mainframe installation with ancillary departments, other patient care areas, and business and financial offices.

The first attempts at computer-assisted patient monitoring in critical care areas took place in the 1960s. Some of the early applications were based on electrocardiographic waveform analysis and attempted to establish the morphologic diagnosis of myocardial ischemia or injury, conduction defects, or chamber enlargement. Other developments focusing on the automated recognition of cardiac arrhythmias followed but had limited success, and to this day arrhythmia interpretation by computer remains an elusive goal since technology has not yet equalled the human mind in recognizing complex patterns.



Microprocessor-Controlled Bedside and Central-Station Monitors, Complemented by a Computer Station.

Today it is unusual to find drug infusion devices, ECG and blood pressure monitors, intraaortic balloon assist pumps, or other critical care unit devices that are not controlled by microprocessors. Many of these new devices also have built-in communication controllers that allow them to transfer information to, and/or be controlled by, an external computer system. A microprocessor-controlled bedside physiologic monitoring unit may be used in a coronary care unit. It is primarily used to acquire, display, and transmit a patient's heart rate, electrocardiogram, and arterial blood pressure, but additional parameters can be incorporated. Built-in audible and visual alarms alert the staff if preset upper or lower limits are exceeded in any monitoring Channel .

With the advancement of technology, personal computers and even small hand-held computers have become popular because they offer powerful processing tools in small and relatively inexpensive packages. Programs are available that accept as input hemodynamic and blood gas information to calculate and print a patient's hemodynamic profile. Other programs control infusion of drugs or blood, perform cardiac output and drug dosage calculations, help manage the treatment of acid-base disorders, or assist in hyperalimentation therapy. These small dedicated computers allow, in some cases,

transfer of information to a larger central system if this is available.

In addition to automated bedside monitoring, computer technology is used in special care units for acquisition, storage and processing of patient information needed for determination and reporting of trends, requesting and reporting ancillary department tests or procedures, and administrative functions such as inventory control and billing. These uses result in the elimination of manual tasks, improve accuracy, eliminate redundant bookkeeping, and reduce the possibility of human errors.

Part I. Comprehension Exercises

A. Put “T” for true and “F” for false statements. Justify your answers.

- 1. A seriously ill child is usually treated in a pediatric intensive care unit.
- 2. Critical care areas in hospitals are rarely computerized.
- 3. Computers have been used in critical care areas since they first appeared on the market.
- 4. Arrhythmia interpretation has successfully been performed by the computer since 1960.
- 5. Small hand-held computers may also be used in critical care units.
- 6. Small computers in critical care units are used for a wide variety of applications.

B. Choose a, b, c, or d which best completes each item.

- 1. The first paragraph mainly describes
 - a. physicians who work in critical care units
 - b. patients who are treated in critical care units
 - c. the critical care areas in a hospital
 - d. the sophisticated equipment used in a hospital
- 2. Compared to human beings, computers
 - a. are superior in mental ability
 - b. are inferior in mental ability
 - c. can better interpret complicated heart patterns
 - d. can better recognize heart arrhythmias
- 3. As we understand from the text, critical care units have been established in the hospitals.....
 - a. for more than 35 years
 - b. for less than 35 years

- c. since the computers could analyze waveforms
 - d. since the computers were produced
4. Paragraph three mainly discusses..... .
- a. electrocardiograph waveform analysis by the computer
 - b. diagnosis of myocardial injury and conduction defects by the computer
 - c. some applications of the computer in critical care areas
 - d. some applications of the computer in the recognition of cardiac arrhythmias
5. It is true that
- a. computer-controlled devices are more effective than intensivists in the recognition of cardiac arrhythmias
 - b. computer-controlled devices have not been as effective as expected in critical care units
 - c. many of the new care unit devices are controlled by the nursing personnel
 - d. most of critical care unit devices are now computerized

C. Answer the following questions orally.

1. What is a critical care unit?
2. What are some examples of intensive care units?
3. What is the function of built-in communication controllers in a computerized critical care unit?
4. How is a built-in audible alarm helpful?
5. What are some of the applications of small computers used in critical care units?
6. What does the last paragraph mainly discuss?

Part II. Language Practice

A. Choose a, b, c, or d which best completes each item.

1. Patients in critical care units are taken care of by.....
 - a. surgeons
 - b. ophthalmologists
 - c. intensivists
 - d. pediatricians
2. The objective of is to register graphically movements of the heart.
 - a. a cardiographer
 - b. a cardiograph
 - c. a transducer
 - d. a simulator

D. Put the following sentences in the right order to form a paragraph. Write the corresponding letters in the boxes provided.

- a. Less critical files should be backed up at intervals dictated by the frequency of updates to those files.
- b. Thus a fire, flood, or other natural disasters cannot cause total loss of information.
- c. The data contained in a medical database are vital to the institution or the medical office, and their safety and integrity must be preserved.
- d. In addition, a standard operating procedure should be instituted whereby weekly or monthly backup files are created and stored in removable storage media, kept preferably at a location remote from the institution or office.
- e. Critical data files, shared by many users and updated frequently, should be backed up on removable storage media at regular and frequent intervals.

1	2	3	4	5

Section Two: Further Reading

Planning and Designing a Computerized Critical Care Unit

Experience has shown that a computer system can help reduce the length of stay of a patient in a special care area. In this age of increased awareness of the cost of providing high technology health care, this fact results in better utilization of the resources in the unit, and lower costs can be expected. Those individuals involved in the process of planning, designing, and activating a state-of-the-art critical care unit incorporating advances in computer technology may want to follow the steps outlined below. However, every institution presents a different environment, and, therefore, individual designs will probably be significantly personalized.

1. A planning committee including medical and nursing staff members as well as high-level administrative personnel should be established. Delegating the task of planning a project of this magnitude to lower-level management may not produce satisfactory results.

2. A comprehensive evaluation of existing and projected patient care needs should be undertaken to determine specifications for the special care unit in question. Some of these specifications may be determined by existing facilities and/or budgetary constraints. A state-of-the-art unit should incorporate the capability to acquire and process signals representing biological variables as continuous functions of time (physiologic monitoring), as well as acquisition, storage, processing, and recall of discrete patient information. If no centralized computer system is available in the institution, then a small dedicated special care system may be a realistic approach. On the other hand, if a comprehensive hospital information system is currently available, then a major objective should be the integration of the special care unit into the central system.

3. When considering automation in the unit, a decision must be made early in the planning process whether to obtain a commercial turnkey hardware-software package, as opposed to the in-house development of a dedicated microprocessor or minicomputer-based special care unit system. If a mainframe central installation is available, a link to it should be considered in either case. Figure 10-2 lists several approaches to linking a mainframe-based hospital information system and a dedicated local or satellite computer.

The choice of one of these approaches may depend to a large degree on existing facilities, equipment, staff, and experience. In-house developed computer systems have the advantages of being designed and built according to the needs and desires of the staff and afford the ability to make changes as they become necessary, sometimes on short notice. One disadvantage of the in-house approach is that the development time may be long, and, therefore, personnel costs may be high. Commercially available turnkey systems, on the other hand, are usually ready for production work once installation is completed, and development time and costs are substantially lower. However, modifications to the system to meet existing institutional policies or procedures, if needed, may be expensive or not possible. This rigidity of design entails, in many cases, modifications in policies or procedures to conform to a somewhat inflexible, commercial package.

These decisions may not be simple but should be based on the approach

that best suits the present and future needs in the existing hospital environment. In any case, provisions should be included during the planning stages for future implementation or expansion of capabilities for automated entry, communication, archiving, processing, and reporting of information.

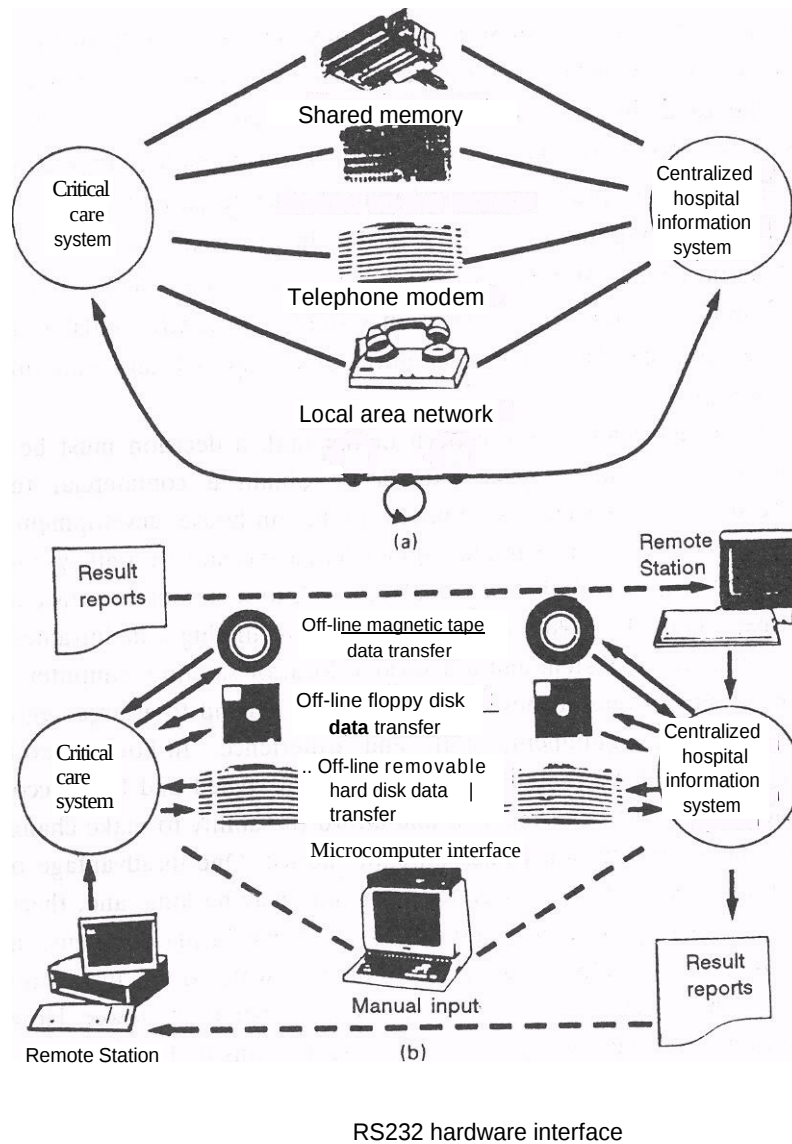


Figure 10-2. Direct (a) and Indirect (b) Approaches to Establishing Communications Between a Dedicated Special Care System and a Central Mainframe-Based Hospital Information System.

4. Whether acquiring a commercial package or designing an in-house system, the administrative aspects of the operation of the special care units should not be neglected. A special care unit package should provide administrative services (or interact with any existing system that already provides them). These should include, but not be limited to, inventory control, patient charges, bed use, and cost analysis information with daily, weekly, and/or yearly reports and appropriate audit trails.

5. Once the new computerized techniques for data acquisition, storage, processing, and reporting become established, usually after a suitable 'paralleling' period, the old manual methods should be discontinued. However, contingency procedures based on the old methods should be established, documented, and tested frequently in the event of an equipment breakdown or other computer system failures.

6. If not already available, an uninterruptible power supply (UPS) should be included as an integral part of the state-of-the-art special care unit. A UPS system should provide power to all computer systems in the unit, if the normal ac service is interrupted. Many computer storage devices (i.e. random-access memory chips) are volatile and do not retain information if the power is interrupted, even for a fraction of a second. Therefore, there is a real potential for losing critical patient information obtained during spontaneous clinical events that cannot be reproduced. UPS systems are available in many configurations and capacities depending on the particular electrical service required. Typically, they include a utility-fed rectifier that supplies dc power to a set of batteries and an inverter that provides clean, transient-free ac power to the equipment. The batteries provide backup during short power failures (minutes) or until the hospital emergency generators take over in case of longer outages.

The actual design of the state-of-the-art critical care unit follows the planning stages and should be a multidisciplinary task. Physicians, nurses, architects, clinical and/or biomedical engineers, data processing personnel, and systems engineers should integrate the design team.

Often, too little thought is given to the practical aspects of the room layout, including computer cabling and connections and the design and placement of computer terminal cabinets. These items are often ignored until after the room is already under construction or completed. Many potential problems can be eliminated by building an actual full-size prototype of the proposed critical care area to ensure optimum placement and accessibility of all monitors, computer-related equipment, and other necessary devices.

B. Write the answers to the following questions.

1. Who should be involved in a planning committee?
2. How do constraints affect the plan specifications?
3. What is the difference between a commercial turnkey system and an in-house developed computer system?
4. What are the advantages and the disadvantages of an in-house computer system?
5. Why is designing a critical care unit considered a multidisciplinary task?



Section Three: Translation Activities

A. Translate the following passage into Persian.

Selection of Monitoring Equipment

There are always questions regarding the number and types of biological variables that should be continuously monitored on special care unit patients. In most institutions, the choice is dictated largely by the capabilities and limitations of the monitoring systems commercially available at any given time

It is desirable for potential users to become familiar with the technical terminology on equipment specification sheets. These specifications usually describe the actual capabilities of the equipment much more clearly and in more detail than do aggressive sales persons or colorful, eye-catching sales literature. If an in-house biomedical engineering department is available, it should evaluate this information and help the medical and nursing staffs interpret it.

Determining what systems are available may require a comprehensive review of the scientific and trade literatures as well as calls and site visits to vendors and users for detailed information. It is recommended that a firm understanding of the capabilities and limitations of any system be established before pricing and contracts are considered. It should be stressed that the selection process should include site visits to institutions that have used or are using equipment similar to that being considered. This should include visits to institutions not in the vendors' reference list.

Finally, it should be noted that the successful implementation of critical care unit systems rests not only on the adequacy of the hardware and software

but also on the human component. Users, including physicians, nurses, and technicians, should not only be involved in the selection process, but should receive comprehensive hands-on training in the use of the equipment before and during the actual implementation phase.

B. Find the Persian equivalents of the following terms and expressions and write them in the spaces provided.

- | | |
|-------------------------------------|-------|
| 1. ancillary | |
| 2. arterial blood pressure | |
| 3. biomeical engineer | |
| 4. cardia | |
| 5. cardiac arrhythmia | |
| 6. coronary care unit | |
| 7. critical care unit | |
| 8. drug infusion device | |
| 9. homodynamic profile | |
| 10. hyperalimentation therapy | |
| 11. intensive care unit | |
| 12. intensivist | |
| 13. intraaortic balloon assist Pump | |
| 14. myocardial ischemia | |
| 15. neurologicall | |
| 16. pediatric intensive care unit | |
| 17. recovery room | |
| 18. respiratory | |
| 19. telemetry monitoring area | |

Unit 11

Section One: Reading Comprehension

Introduction to Control Systems Analysis

Automatic control has played a vital role in the advance of engineering and science. In addition to its extreme importance in space-vehicle systems, missile-guidance systems, aircraft-autopiloting systems, robotic systems, and the like, automatic control has become an important and integral part of modern manufacturing and industrial processes. For example, automatic control is essential in the numerical control of machine tools in the manufacturing industries. It is also essential in such industrial operations as controlling pressure, temperature, humidity, viscosity, and flow in process industries.

In studying control engineering, we need to define those terms that are necessary to describe control systems, such as plants, disturbances, feedback control, and feedback control systems.

Plants. A plant is a piece of equipment, perhaps just a set of machine parts functioning together, the purpose of which is to perform a particular operation. In control systems, any physical object to be controlled such as a heating furnace, or a spacecraft is called a plant.

Disturbances. A disturbance is a signal that tends to adversely affect the value of the output of a system. If a disturbance is generated within the system, it is called *internal*, while an *external* disturbance is generated outside the system and is an input.

Feedback Control. Feedback control refers to an operation that, in the presence of disturbances, tends to reduce the difference between the output of a system and some reference input and that does so on the basis of this difference. Here only unpredictable disturbances are so specified, since predictable or known disturbances can always be compensated for within the system.

Feedback Control Systems. A system that maintains a prescribed relationship between the output and some reference input by comparing them and using the difference as a means of control is called a *feedback control*

system. An example would be a room-temperature control system. By measuring the actual room temperature and comparing it with the reference temperature (desired temperature), the thermostat turns the heating or cooling equipment on or off in such a way as to ensure that the room temperature remains at a comfortable level regardless of outside conditions.

Servo Systems. A servo system (or servomechanism) is a feedback control system in which the output is some mechanical position, velocity, or acceleration. Therefore, the terms servo system and position- (or velocity- or acceleration-) control system are synonymous. Servo systems are extensively used in modern industry. For example, the completely automatic operation of machine tools, together with programmed instruction, may be accomplished by the use of servo systems. It is noted that a control system, whose output (such as the position of an aircraft in space in an automatic landing system) is required to follow a prescribed path in space, is sometimes called a servo system, also. Examples include the robot-hand control system, where the robot hand must follow a prescribed path in space, and the aircraft automatic landing system, where the aircraft must follow a prescribed path in space.

Automatic Regulating Systems. An automatic regulating system is a feedback control system in which the reference input or the desired output is either constant or slowly varying with time and in which the primary task is to maintain the actual output at the desired value in the presence of disturbances. There are many examples of automatic regulating systems, some of which are the Watt's flyball governor, automatic regulation of voltage at an electric power plant in the presence of a varying electrical power load, and automatic control of the pressure and temperature of a chemical process.

Process Control Systems.

An automatic regulating system in which the output is a variable, such as temperature, pressure, flow, liquid level, or pH, is called a *process control system*. Process control is widely applied in industry. Programmed controls such as the temperature control of heating furnaces in which the furnace temperature is controlled according to a preset program are often used in such systems.

Closed-Loop Control Systems. Feedback control systems are often referred to as *closed-loop control systems*. In practice, the terms feedback control and closed-loop control are used interchangeably. In a closed-loop control system the actuating error signal, which is the difference between the input signal and the feedback signal (which may be the output signal itself or a function of the output signal and its derivatives), is fed to the controller so as to reduce the error and bring the output of the system to a desired value.

The term dosed-loop control always implies the use of feedback control action in order to reduce system error.

Open-Loop Control Systems. Those systems in which the output has no effect on the control action are called *open-loop control systems*. In other words, in an open-loop control system the output is neither measured nor fed back for comparison with the input. One practical example is a washing machine. Soaking, washing, and rinsing in the washer operate on a time basis. The machine does not measure the output signal, that is, the cleanliness of the clothes.

Adaptive Control Systems. The dynamic characteristics of most control systems are not constant for several reasons, such as the deterioration of components as time elapses or the changes in parameters and environment. Although the effects of small changes on the dynamic characteristics are attenuated in a feedback control system, if changes in the system parameters and environment are significant, a satisfactory system must have the ability of adaptation. Adaptation implies the ability to self-adjust to self-modify in accordance with unpredictable changes in conditions of environment or structure. The control system having a candid ability of adaptation (that is, the control system itself detects changes in the plant parameters and makes necessary adjustments to the controller parameters in order to maintain an optimal performance) is called the *adaptive control system*.

Learning Control Systems. Many apparently open-loop control systems can be converted into closed-loop control systems if a human operator is considered a controller, comparing the input and output and making the corrective action based on the resulting difference or error.

If we attempt to analyze such human-operated dosed-loop control systems we encounter the difficult problem of writing equations that describe the behavior of a human being. One of the many complicating factors in this case is the learning ability of the human operator. As the operator gains more experience, he or she will become a better controller, and this must be taken into consideration in analyzing such a system. Control systems having an ability to learn are called *learning control systems*.

Part I. Comprehension Exercises

A. Put “T” for true and “F” for false statements. Justify your answers.

..... 1. Automatic control is an essential part of modern manufacturing and industrial processes.

- 2. Any machine which is being controlled is referred to as a plant.
- 3. The mechanism of a servo system is different from that of a position control system.
- 4. A control system whose output is required to follow a prescribed condition may be called a servo system.
- 5. The automatic controller functions more effectively in an open-loop control system.
- 6. An adaptive control system is capable of accommodating unpredictable environmental changes, whether these changes occur within the system or external to it.
- 7. An adaptive control system is designed to modify the control signal as the system environment changes, so that performance is always optimal whereas the human operator recognizes familiar inputs and can use past learned experiences to react in an optimal manner.

B. Choose a, b, c, or d which best completes each item.

- 1. A disturbance of a system.
 - a. has a positive effect on the output
 - b. has an unfavorable effect on the output
 - c. increases the efficiency
 - d. controls the efficiency
- 2. Feedback gives an automatic control system the ability..... .
 - a. to deal with the unexpected disturbances in the plant behavior
 - b. to deal with the predictable disturbances in the plant behavior
 - c. to maintain a steady relationship between the output and some reference input
 - d. to maintain the actual value of a disturbance constant
- 3. It is true that
 - a. the mechanism of the automatic regulating system is based on that of the process control system
 - b. the mechanism of the process control system is different from that of the automatic regulating system
 - c. an automatic regulating system compares the actual value of the plant output with the desired value
 - d. an automatic regulating system maintains the actual value of the plant output at the desired value in the presence of disturbances

4. In a closed-loop control system, the controller
 - a. regulates the internal disturbances of the plant and keeps them under control
 - b. provides information about the actual plant output
 - c. reduces the difference between the input signal and the output signal and brings the output of the system to a desired value
 - d. determines the value of the error and reduces its effect on the system
5. We may infer from the text that open-loop control systems
 - a. should be used for systems in which unpredictable disturbances occur
 - b. should be used for systems in which the inputs are known in advance
 - c. are more complicated than closed-loop control systems
 - d. are more powerful than closed-loop control systems

C. Answer the following questions orally.

1. What part has automatic control played in the advancement of engineering?
2. What is a plant?
3. What are the internal and external disturbances?
4. What is the function of a feedback control system?
5. What is a servo system?
6. What is the mechanism of an open-loop control system based on?
7. What is an adaptive control system?

Part II. Language Practice

A. Choose a, b, c, or d which best completes each Item.

1. The most fascinating developments in adaptive control systems lie in the areas of pattern recognition and systems.

a. learning	b. operating .
c. analyzing	d. reading
2. An automatic compares the actual value of the plant output with the desired value, determines the deviation, and produces a control signal that will reduce the deviation to zero or to a small value.

a. amplifier	b. sensor
c. controller	d. transformer
3. A maintains the plant output constant at the desired value in the presence of external disturbances.

a. capacitor	b. compensator
c. resistor	d. regulator

4. In everyday life, occurs when we are aware of the consequences of our actions.
 - a. adaptation
 - b. regulation
 - c. feedback
 - d. control
5. Control systems without feedback are called.
 - a. closed-loop
 - b. open-loop
 - c. adaptive
 - d. learning

B. Fill in the blanks with the appropriate form of the words given.

1. Heat

- a. For high current levels, an external pass transistor may be required with sinks to reduce the effective thermal resistance.
- b. A heat coil is a protective device that grounds or opens a circuit, or does both, by means of a mechanical element that is allowed to move when the fusible substance that holds it in place is above a predetermined temperature by the current in the circuit.
- c. A heater connector is designed to engage the male terminal pins of a or cooling appliance.
- d. A heater transformer supplies power for electron-tube filaments or of indirectly heated cathodes.

2. Adapt

- a. An System is capable of accommodating unpredictable environmental changes, whether these changes occur within the system or external to it.
- b. The vagueness surrounding most definitions and classifications of adaptive systems is due to the large variety of mechanisms by which may be achieved.
- c. When high is called for most present-day requirements will be met by an identification-decision-modification system.

3. Accomplish

- a. The dynamic characteristics of a plant must be measured and identified continuously. This should be without affecting the normal operation of the system.
- b. When tied in with learning approaches, pattern-recognition techniques will adaptive-learning control.
- c. A business system may consist of many groups. Feedback methods of reporting the of each group must be established in such a system for proper operation.

4. Cool

- a. A heat exchanger or is used in rotating machinery to transfer heat between two fluids without direct contact between them.
- b. The of regulator elements refers to the method used for removing heat generated in the regulating process.
- c. In an electron device, a metallic part or fin extends thearea to facilitate the dissipation of the heat generated in the device.
- d. Air may be used as a..... to remove heat from a machine.

5. Change

- a. Modification refers to the of control signals according to the results of the identification and decision.
- b. If parameters are..... rapidly, a procedure known as alternate biasing is employed.
- c. Adaptive control systems are designed to modify the control signal as the system environment so that performance is always optimal.
- d. Feedback allows us to cope with a..... environment by adjusting our actions in the presence of unforeseen events.

C. Fill in the blanks with the following words.

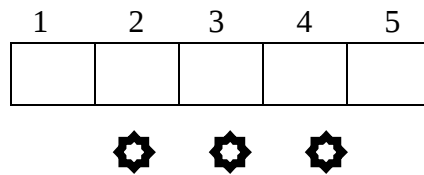
parameters	feedback	control	case
inaccurate	external	given	

An advantage of the closed-loop system is the fact that the use of makes the system response relatively insensitive to..... disturbances and internal variations in system It is thus possible to use relatively and inexpensive components to obtain the accurate control of a/an plant, whereas doing so is impossible in the open-loop

D. Put the following sentences in the right order to form a paragraph. Write the corresponding letters in the boxes provided.

- a. Some systems may have multiple inputs and multiple outputs.
- b. A system may have one input and one output.
- c. Such a system is called a single-input, single-output control system.
- d. An example of such multiple-input, multiple-output systems is a process control system that has two inputs (pressure input and temperature input) and two outputs (pressure output and temperature output).

- e. An example is a position control system, where there is one command input (desired position) and one controlled output (output position).



Section Two: Further Reading

Examples of Control Systems

Speed Control System. The basic principle of a Watt's speed governor for an engine is illustrated in the schematic diagram of Figure 11-1. The amount of fuel admitted to the engine is adjusted according to the difference between the desired and the actual engine speeds.

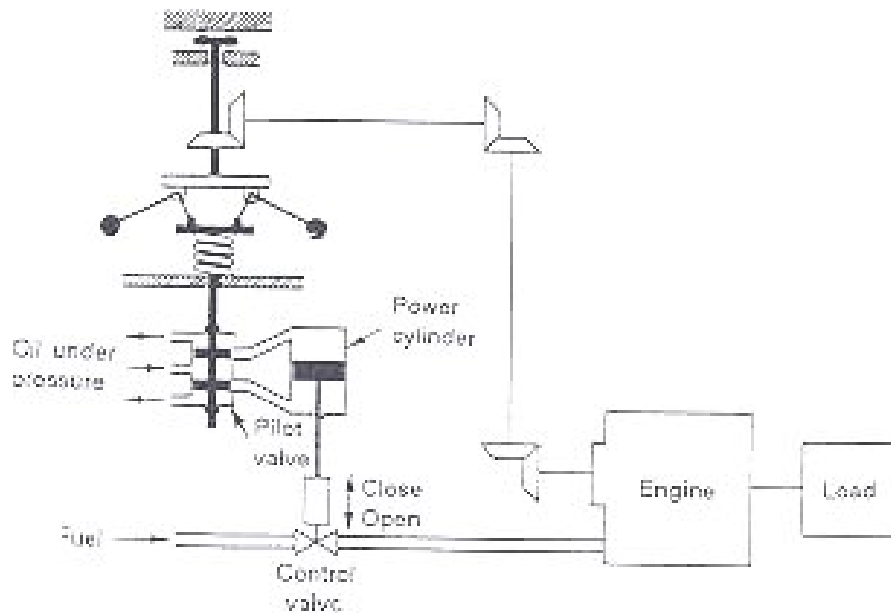


Figure 11-1. Speed Control System

The sequence of actions may be stated as follows: The speed governor is adjusted such that, at the desired speed, no pressured oil will flow into either side of the power cylinder. If the actual speed drops below the desired value due to disturbance, then the decrease in the centrifugal force of the speed

governor causes the control valve to move downward, supplying more fuel, and the speed of the engine increases until the desired value is reached. On the other hand, if the speed of the engine increases above the desired value, then the increase in the centrifugal force of the governor causes the control valve to move upward. This decreases the supply of fuel, and the speed of the engine decreases until the desired value is reached.

Robot Control System. Industrial robots are frequently used in industry to improve productivity. The robot can handle monotonous jobs as well as complex jobs without errors in operation. The robot can work in an environment intolerable to human operators. For example, it can work in extreme temperatures or in a high- or low-pressure environment or under water or in space. There are special robots for fire fighting, underwater exploration, and space exploration, among many others.

The industrial robot must handle mechanical parts that have particular shapes and weights. Hence, it must have at least an arm, a wrist, and a hand. It must have sufficient power to perform the task and the capability for at least limited mobility. In fact, some robots of today are able to move freely by themselves in a limited space in a factory.

The industrial robot must have some sensory devices. In low-level robots, microswitches are installed in the arms as sensory devices. The robot first touches an object and then, through the microswitches, confirms the existence of the object in space and proceeds in the next step to grasp it.

In a high-level robot, an optical means (such as a television system) is used to scan the background of the object. It recognizes the pattern and determines the presence and orientation of the object. A computer is necessary to process signals in the pattern-recognition process (see Figure 11-2).

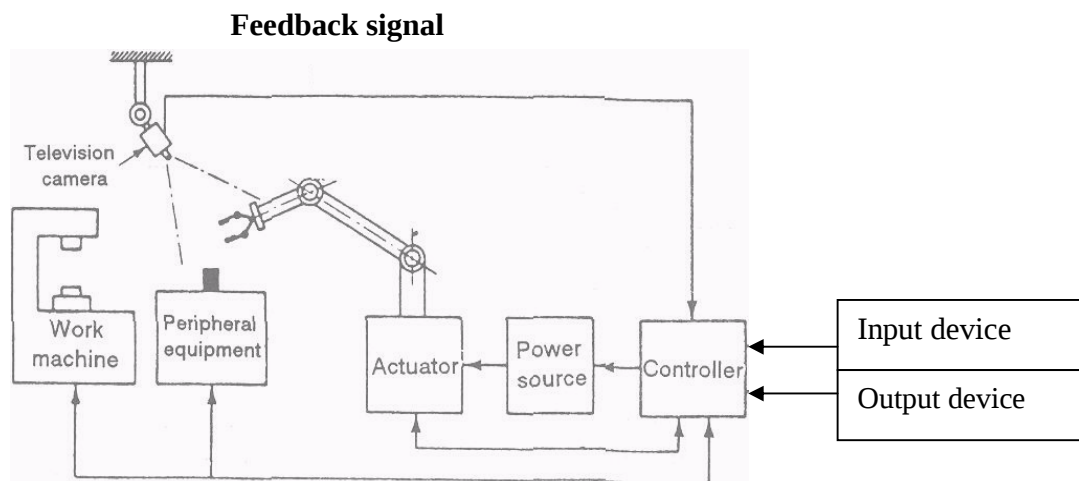


Figure 11-2. Robot Using a Pattern Recording Process.

In some applications, the computerized robot recognizes the presence and orientation of each mechanical part by a pattern recognition process that consists of reading the code numbers attached to it. Then the robot picks up the part and moves it to an appropriate place for assembling, and there it assembles several parts into a component. A well-programmed digital computer acts as a controller.

Temperature Control System. Figure 11-3 shows a schematic diagram of temperature control of an electric furnace. The temperature in the electric furnace is measured by a thermometer, which is an analog device. The analog temperature is converted to a digital temperature by an A/D converter. The digital temperature is fed to a controller through an interface. This digital temperature is compared with the programmed input temperature, and if

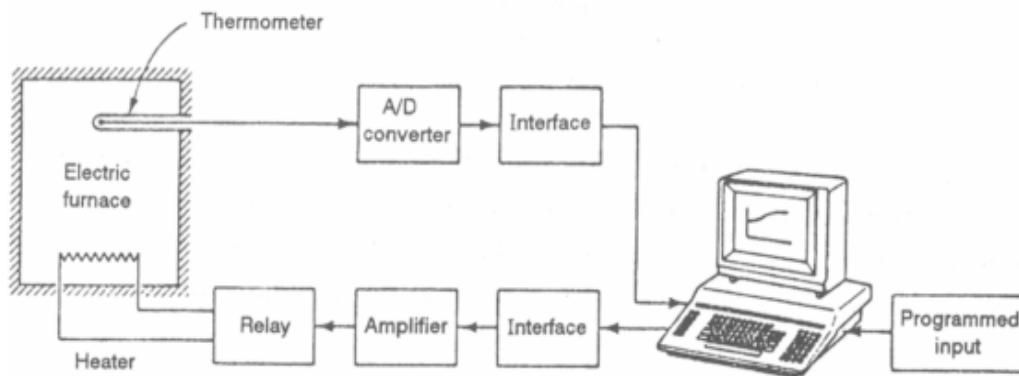


Figure 11-3. Temperature Control System.

there is any discrepancy (error), the controller sends out a signal to the heater, through an interface, amplifier, and relay, to bring the furnace temperature to a desired value.

Temperature Control of the Passenger Compartment of a Car. Figure 11-4 shows a functional diagram of temperature control of the passenger compartment of a car. The desired temperature, converted to a voltage, is the input to the controller. The actual temperature of the passenger compartment is converted to voltage through a sensor and is fed back to the controller for comparison with the input. The ambient temperature and radiation heat transfer from the sun, which are not constant while the car is driven, act as disturbances. This system employs both feedback control and feed forward control. (Feed forward control gives corrective action before the disturbances affect the output.)

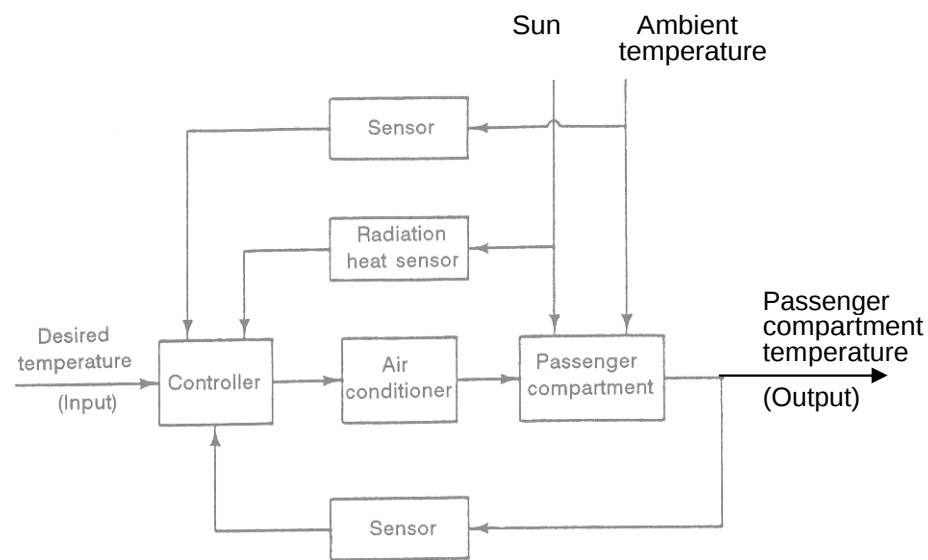


Figure 11-4. Temperature Control of Passenger Compartment of a Car.

The controller receives the input signal, output signal, and signals from sensors from disturbance sources. The controller sends out an optimal control signal to the air conditioner to control the amount of cooling air so that the passenger compartment temperature is equal to the desired temperature.

Comprehension Exercises

A. Choose a, b, c, or d which best completes each item.

1. The second paragraph mainly discusses
 - a. the factors affecting the speed of an engine
 - b. the disturbances created due to fluctuations in the speed of an engine
 - c. the mechanism of the speed governor to adjust the speed of an engine
 - d. the rate of oil flow in a speed governor to reduce external disturbances
2. In the speed control system just described, the amount of fuel to be applied to the engine is known as

a. the force	b. the disturbance
c. the feedback signal	d. the actuating signal
3. According to the text, robots
 - a. cannot work at very low temperatures

- b. cannot handle jobs without error
- c. are equipped with proper devices to be able to perform the tasks required
- d. are powered to handle various jobs performed by man in industry

4. A high-level robot
- a. performs pattern recognition to determine how to pick up, move, and assemble the parts into a component
 - b. determines the presence and orientation of parts to change their shapes and weights if needed
 - c. does not have enough control over the process it goes through
 - d. does not have the proper program for pattern recognition in its computer memory
5. The examples given in this text are of control systems.
- a. closed-loop
 - b. open-loop
 - c. hydraulic
 - d. pneumatic

- B. Write the answers to the following questions.**
1. What is the function of a speed governor?
 2. How freely is an industrial robot able to move?
 3. What is the use of a microswitch installed in the arm of a robot?
 4. How does a temperature control system work?
 5. What are considered disturbances in the temperature control of the passenger compartment of a car?
 6. What constitutes the temperature control of the passenger compartment of a car?



Section Three: Translation Activities

- A. Translate the following passage into Persian.**

Historical Review

The first significant work in automatic control was James Watt’s centrifugal governor for the speed control of a steam engine in the eighteenth century. Other significant works in the early stages of development of control theory were due to Minorsky, Hazen, and Nyquist, among many others. In 1922,

Minorsky worked on automatic controllers for steering ships and showed how stability could be determined from the differential equations describing the system. In 1932, Nyquist developed a relatively simple procedure for determining the stability of closed-loop systems on the basis of open-loop response to steady-state sinusoidal inputs. In 1934 Hazen, who introduced the term servomechanisms for position control systems, discussed the design of relay servomechanisms capable of closely following a changing input.

During the decade of the 1940s, frequency-response methods made it possible for engineers to design linear closed-loop control systems that satisfied performance requirements. From the end of the 1940s to early 1950s, the root-locus method due to Evans was fully developed.

The frequency-response and root-locus methods, which are the core of classical control theory, lead to systems that are stable and satisfy a set of more or less arbitrary performance requirements. Such systems are, in general, acceptable but not optimal in any meaningful sense. Since the late 1950s, the emphasis in control design problems has been shifted from the design of one of many systems that work to the design of one optimal system in some meaningful sense.

As modern plants with many inputs and outputs become more and more complex, the description of a modern control system requires a large number of equations. Classical control theory, which deals only with single-input, single-output systems, becomes powerless for multiple-input, multiple-output systems. Since about 1960, because the availability of digital computers made possible time-domain analysis of complex systems, modern control theory, based on time-domain analysis and synthesis using state variables, has been developed to cope with the increased complexity of modern plants and the stringent requirements on accuracy, weight, and cost in military, space, and industrial applications.

B. Find the Persian equivalents of the following terms and expressions and write them in the spaces provided.

1. actuate
2. adaptive control system
3. aircraft-autopiloting system
4. alternate biasing
5. automatic regulating system
6. classical control theory
7. closed-loop control system

8.compensator	
9. complex system	
10. electromagnetic valve	
11. external disturbance	
12. feedback control	
13. internal disturbance	
14. linear control system	
15. microswitch	
16. missile-guidance system	
17. multiple-input	
18. multiple-output	
19. numerical control	
20. open-loop system	
21. pneumatic control	
22. position control system	
23. preset	
24. process control system	
25. robot locus method	
26. robotic system	
27. sensory device	
28. servomechanism	
29. servo system	
30. space-vehicle system	
31. steady-state	
32. time-domain analysis	
33. Watt's centrifugal governor	
34. Watt's (Watt's flyball governor) governor	
35. Watt's speed governor	

Magnetic Line Voltage Starters

Magnetic control means the use of electromagnetic energy to close switches. Line voltage (across the line) magnetic starters are electromechanical devices that provide a safe, convenient, and economic means of full voltage starting and stopping motors. In addition, these devices can be controlled remotely. They are used when a full-voltage starting torque may be applied safely to the driven machinery and when the current inrush resulting from across-the-line starting is not objectionable to the power system. Control for these starters is usually provided by pilot devices such as push buttons, float switches, timing relays, etc. Automatic control is obtained from the use of some of these pilot devices.

Magnetic vs Manual Starters

Using manual control, the starter must be mounted so that it is easily within reach of the machine operator. With magnetic control, push-button stations are mounted nearby, but automatic control pilot devices can be mounted almost anywhere on the machine. The push buttons and automatic pilot devices can be connected by control wiring into the coil circuit of a remotely mounted starter, possibly closer to the motor to shorten the power circuit.

In the construction of a magnetic controller, the armature is mechanically connected to a set of contacts so that, when the armature moves to its closed position, the contacts also close. There are different variations and positions, but the operating principle is the same.

The simple up-and-down motion of a solenoid-operated, three-pole magnetic switch is shown in Figure 12-1. Not shown are the motor overload relays and maintaining and auxiliary electrical contacts. Double break contacts are used on this type of starter to cut the voltage in half on each contact, thus providing high arc rupturing capacity and longer contact life.

The operating principle that makes a magnetic starter different from a manual starter is the use of an electromagnet. Electrical control equipment makes extensive use of a device called a solenoid. This electromechanical device is used to operate motor starters, contactors, relays, and valves. By

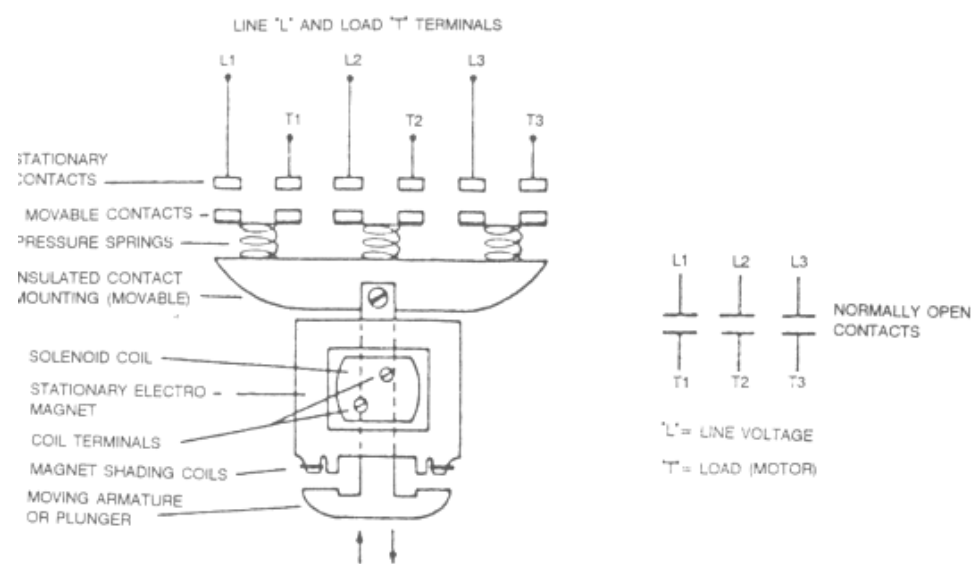


Figure 12-1. Three-Pole, Solenoid-Operated Magnetic Switch (Contactor) and Electrical Wiring Symbols.

placing a coil of many turns of wire around a soft iron core, the magnetic flux set up by the energized coil tends to be concentrated; therefore, the magnetic field effect is strengthened. Since the iron core is the path of least resistance to the magnetic lines of force, magnetic attraction concentrates according to the shape of the magnet core.

There are several different variations in design of the basic solenoid magnetic core and coil. Figure 12-2 shows a few examples. As shown in the solenoid design of Figure 12-2C, linkage to the movable contacts assembly is obtained through a hole in the movable plunger. The plunger is shown in the open de-energized position.

The center leg of each of the E-shaped magnet cores in Figures 12-2B and C is ground shorter than the outside legs to prevent the magnetic switch from accidentally staying closed (due to residual magnetism) when power is disconnected.

Figure 12-3 shows a manufactured magnet structure and how the starter contacts are mounted on the armature.

When a magnetic motor starter coil is energized and the armature has been sealed in, it is held tightly against the magnet assembly. A small air gap is always deliberately placed in the center leg, iron circuit. When the coil is de-energized, a small amount of magnetism remains. If it were not for this gap

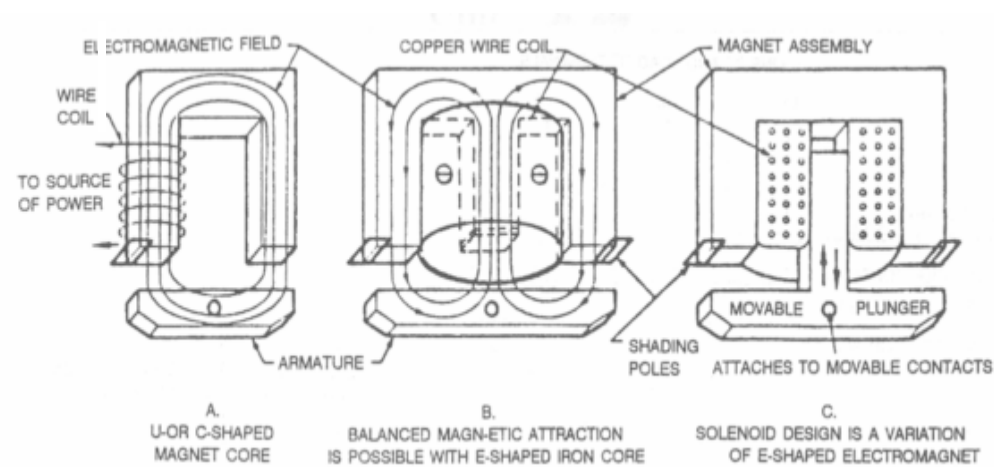


Figure 12-2. Some Variations of Basic Magnet Core and Coil Configurations of Electromagnets.

in the iron circuit, the residual magnetism might be enough to hold the movable armature in the sealed-in position. This knowledge can be important to the electrician when troubleshooting a motor that will not stop.

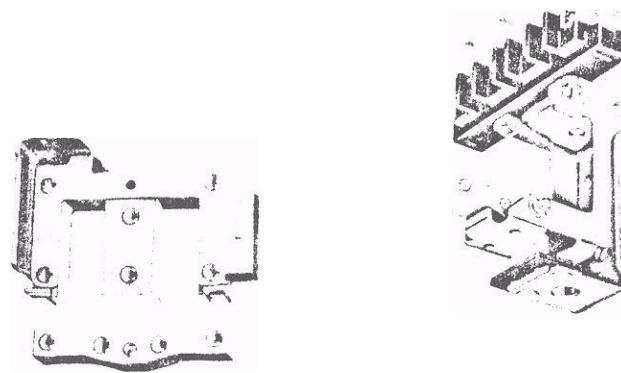


Figure 12-3. Magnet Structure (Left) and Movable Contacts and Armature Guide Assembly (Right) of a Four-Pole Magnetic Switch (Courtesy Square D Co.).

The OFF or OPEN position is obtained by de-energizing the coil and allowing the force of gravity or spring tension to release the plunger from the magnet body, thereby opening the electrical contacts. The actual contact surfaces of the plunger and core body are machine ground to insure a high degree of flatness on the contact surfaces so that operation on alternating current is quieter. Improper alignment of the contacting surfaces and foreign matter between the surfaces may cause a noisy hum on alternating-current magnets.

Another source of noise is loose laminations. The magnet body and plunger (armature) are made up of thin sheets of iron laminated and riveted together to reduce *eddy currents* and hysteresis, iron losses showing up as heat

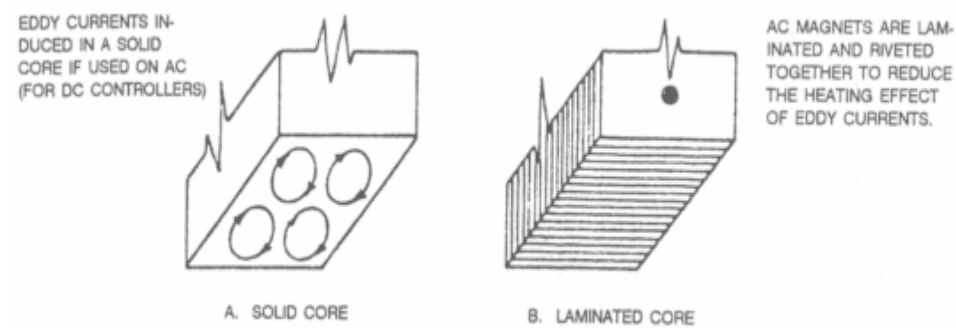


Figure 12-4. Types of Magnet Cores.

(see Figure 12-4). Eddy currents are shorted currents induced in the metal by the transformer action of an ac coil. Although these currents are small, they heat up the metal, create an iron loss, and contribute to inefficiency. At one time, laminations in magnets were insulated from each other by a thin, nonmagnetic coating; however, it was found that the normal oxidation of the metallic laminations reduces the effects of eddy currents to a satisfactory degree, thus eliminating the need for a coating.

Part I. Comprehension Exercises

A. Put "T" for true and "F" for false statements. Justify your answers.

- 1. Line voltage magnetic starters provide electric motors with safe current.
- 2. Magnetic controllers can be mounted anywhere on the machine.
- 3. The armature plays the major part in the construction of a magnetic controller.
- 4. A solenoid is an electromechanical device.
- 5. Alternating-current magnets are inclined to create noise.
- 6. Eddy currents do not affect the efficiency of alternating-current magnets.

B. Choose a, b, c, or d which best completes each item.

- 1. The first paragraph mainly discusses
 - a. the application of full voltage starting torque to the driven mach

- b. the application of safe current to power systems
 - c. the use of electromagnetic energy to control electrical machinery
 - d. the use of pilot devices to control magnetic starters
2. In solenoid, magnetic concentration
- a. directly depends on the design of the solenoid magnetic core and coil
 - b. directly depends on the number of turns of wire around the solenoid core
 - c. reduces the amount of core resistance to the magnetic lines of force
 - d. reduces the magnetic field effect caused by the magnetic coil
3. The shorter length of the center leg of an E-shaped magnet core when power is disconnected.
- a. concentrates more energy than the outside legs
 - b. concentrates less energy than the outside legs
 - c. causes the magnetic switch to stay closed
 - d. causes the magnetic switch to function properly
4. As we understand from the text,
- a. the shape of the magnet core of a solenoid determines the concentration of the magnetic attraction
 - b. the magnetic flux set up by the coil concentrates according to the energy applied to it
 - c. manual starters are practically more useful than magnetic starters
 - d. solenoids are not widely used in electrical control equipment
5. If a small air gap were not placed in the center leg of the magnet core, when the magnetic motor starter coil was de-energized.
- a. the residual magnetism would increase
 - b. the plunger might release the magnet body
 - c. a motor might not stop working
 - d. a motor might not continue working

C. Answer the following questions orally.

1. How are magnetic starters usually controlled?
2. What is the function of double break contacts on a solenoid-operated, three-pole magnetic switch?
3. What is the operating principle of a magnetic controller?
4. What is the use of the hole in the movable plunger?
5. How do the electrical contacts of a magnetic starter open?
6. What are eddy currents?

Part II. Language Practice

A. Choose a, b, c, or d which best completes each item.

1. An electric controller , , is used to start and stop a motor.
 - a. a switch
 - a starter
 - a push button
 - a timing relay
2. An starter connects the motor to the supply without the use of a resistance or autotransformer to reduce the voltage.
 - across-the-line
 - automatic
 - autotransformer
 - increment
3. Due to.....,ferromagnetic bodies retain a certain magnetization after the magnetizing force has been removed.
 - solenoid conductivity
 - solenoid resistivity
 - residual modulation
 - residual magnetism
4. Voltage induced in the body of a conducting mass by a variation of magnetic flux results in
 - eddy currents
 - electromagnetic energy
 - hysteresis
 - magnetism
5. Alternating-current magnets are laminated and riveted together to reduce
 - magnetic flux
 - magnetic charge
 - induction
 - hysteresis

B. Fill in the blanks with the appropriate form of the words given.

1. Energize

- The available is the amount of work that a system is capable of doing.
- When excessive current is drawn, the relay de- the starter and stops the motor.

2. Devise

- The human operator may be replaced by a mechanical ,electrical ,or similar
- Circuit controllers are to close and open electric circuits.

3. Maintain

- The controller functions to the furnace temperature close to the varying set point.
- Quick tripping may be by a time delay overload relay.

4. Vary

- a. Most physical systems are nonlinear to extents.
- b. If the range of.....of the system variables is not wide, the system may be linearized within a relatively small range of variation of variables.
- c. In a continuous time control system, all system are functions of a continuous time t.
- d. A time-invariant control system is one whose parameters do notwith time.

5. Overload

- a. To provide or running protection to keep a motor from overheating, overload relays are used on starters to limit the amount of current drawn to a predetermined value.
- b. Unless the cause of the has been removed, the overload relay will trip again.
- c. A magnetic overload relay is used to stop an electrical machine when it is

C. Fill in the blanks with the following words.

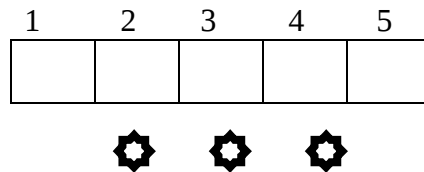
reached	used	material
overload	line	immediately
cut off	high	blockage

Instantaneous trip current relays areto take a motor off the as soon as a predetermined load condition is For example, when a blockage ofon a woodworking machine causes a sudden..... current, an instantaneous trip relay can the motor quickly. After the cause of the is removed, the motor can be restartedbecause the relay resets itself as soon as the is removed.

D. Put the following sentences in the right order to form a paragraph. Write the corresponding letters in the boxes provided.

- a. As a result, the coil of the magnetic relay must be wound with wire large enough in size to pass the motor current.
- b. In some cases, the relay may also be used so that it is actuated when the current falls to a certain value.
- c. The magnetic overload relay coil is connected in series directly with the motor or is indirectly connected by current transformers (as in circuits with large motors).

- d. They are used when an electrical contact must be opened or closed as the actuating current rises to a certain value.
- e. These overload relays operate by current intensity and not heat.



Section Two: Further Reading

Motor Overheat

An electric motor does not know enough to quit when the load gets too much for it. It keeps going until it burns out. If a motor is subjected, over a period of time, to internal or external heat levels that are high enough to destroy the insulation on the motor windings, it will fail-burn out.

A solution to this problem *might* be to install a larger motor whose capacity is in excess of the normal horsepower required. This is not too practical since there are other reasons for a motor to overheat besides excess loads. A motor will run cooler in the winter cold than in the summer heat of a tropical climate. A high, surrounding air temperature (*ambient temperature*) has the same effect as higher-than-normal current flow through a motor-it tends to deteriorate the insulation on the motor windings.

High ambient temperature is also created by *poor ventilation* of the motor. Motors must get rid of their heat, so any obstructions to this process must be avoided. High inrush currents of *excessive starting* create heat within the motor. The same is true with *starting heavy loads*. There are several other related causes that generate heat within a motor such as *voltage unbalance*, *low voltage*, and single phasing. In addition, when the rotating member of the motor will not turn (a condition called *locked rotor*), heat is generated.

The ideal overload protection for a motor is an element with current sensing properties very similar to the heating curve of the motor. This would act to open the motor circuit when full load is exceeded. The operation of the protective device is ideal if the motor is allowed to carry small, short and harmless overloads, but is quickly disconnected from the line when an overload has persisted too long. Dual element, or time-delay, fuses may

provide motor overload protection, but they have the disadvantage of being nonrenewable and must be replaced.

An overload relay is added to the magnetic switch that was shown in Figure 12-1. Now it is called a motor starter. The overload relay assembly is the heart of motor protection. The motor can do no more work than the overload relay permits. Like the dual element fuse, the overload relay has characteristics permitting it to hold in during the motor accelerating period when the inrush current is drawn. Nevertheless, it still provides protection on small overloads above full-load current when the motor is running. Unlike the fuse, the overload relay can be reset. It can withstand repeated trip and reset cycles without need of replacement. It is emphasized that the overload relay does *not* provide short circuit protection. This is the function of overcurrent protective equipment like fuses and circuit breakers, generally located in the disconnecting switch enclosure.

Current drawn by a motor is a convenient and accurate measure of the motor load and motor heating. Therefore, the device used for overload protection, the overload relay, is usually connected with the motor current. It is provided as part of the starter or controller. As the relay carries the motor current, it is affected by that current. If a dangerous over-current condition occurs, it operates or trips the relay to open the control circuit of the magnetic starter and disconnect the motor from the line; this helps insure the maximum operating life of the motor. In a manual starter, an overload trips a mechanical latch causing the starter contacts to open and disconnect the motor from the line.

The controller is normally installed in the same room or area as the motor. This makes it subject to the same ambient temperature as the motor. The tripping characteristic of the proper thermal overload relay will then be affected by room temperature exactly as the motor is affected. This is done by selecting a thermal relay element (from a chart provided by the manufacturer) that trips at the danger temperature for the motor windings. When excessive current is drawn, the relay de-energizes the starter and stops the motor.

Overload relays can be classified as being either *thermal* or *magnetic*. Magnetic overload relays react only to current excesses, and are not affected by temperature. As the name implies, *thermal overload relays* depend on the rising ambient temperature and temperatures caused by the overload current to trip the overload mechanism.

Comprehension Exercises

A. Choose a, b, c, or d which best completes each item.

1. Overloading causes a motor to overheat which in turn results in
 - a. starter breakdown
 - b. its failure
 - c. low coil currents
 - d. low magnetic pull
2. As we understand from the text, a motor can best be protected against overheating by.....
 - a. an overload relay
 - b. reducing the ambient temperature
 - c. a time-delay fuse
 - d. reducing high inrush currents
3. Paragraph Jive mainly describes.....
 - a. a motor accelerating period
 - b. a motor running period
 - c. the function of a dual element fuse
 - d. the function of an overload relay
4. It may be inferred from the text that
 - a. it is impossible to design a motor that will adjust itself to all the various changes of heat
 - b. it is possible to design a motor that will adjust itself to all the various changes of heat
 - c. poor ventilation does not always create high ambient temperature
 - d. starting heavy loads do not always create heat within a motor
5. It is true that
 - a. the overload relay has nothing to do with the motor current
 - b. the controller is usually installed in a different room away from the motor
 - c. current drawn by a motor is directly proportional to the motor load
 - d. current drawn by a motor de-energizes the starter and stops the motor

B. Write the answers to the following questions.

1. Why are time-delay fuses not considered the best overload protective devices?
2. How does an overload protective device function?
3. What is a dual element fuse?
4. What is a motor starter?
5. What factors may cause a motor to be overheated?
6. Why is an overload relay connected with the motor current?
7. How are overload relays classified?



Section Three: Translation Activities

A. Translate the following passage into Persian.

Time Limit Overload Relays

Time delay overload relays make use of the oil dash pot principle. Motor current passing through the coil of the relay exerts a magnetic pull on a plunger. The magnetic flux set up inside the coil tends to raise the plunger which is attached to a piston immersed in oil. As the current increases in the relay coil, so does the magnetic flux. The force of gravity is overcome and the plunger and piston move upward. During this upward movement, oil is forced through bypass holes in the piston. As a result, the operation of the contacts is delayed. A valve disc is turned to open or close bypass holes of various sizes in the piston. This action changes the rate of oil flow and so adjusts the time delay factor. The rate of upward travel of the core and piston depends directly upon the degree of overload. The greater the current load, the faster the upward movement. As the rate of upward movement increases, the relay tripping time decreases.

This inverse time characteristic prevents the relay from tripping on the normal starting current or on harmless momentary overloads. In these cases, the line current drops to its normal value before the operating coil is able to lift the core and piston far enough to operate the overload control contacts. However, if the overcurrent continues for a prolonged period, the core is pulled far enough to operate the contacts. As the line current increases, the relay tripping time decreases. Tripping current adjustment is achieved by adjusting the plunger core with respect to the overload relay coil. Quick tripping is obtained through the use of a light grade dashpot oil and by adjustment of the oil bypass holes.

A valve in the piston allows almost instantaneous resetting of the circuit to restart the motor. The current must then be reduced to a very low value before the relay will reset. This action is accomplished automatically when the tripping of the relay disconnects the motor from the line. Magnetic overload relays are available with either automatic reset contacts or hand reset contacts.

B. Find the Persian equivalents of the following terms and expressions and write them in the spaces provided.

1. armature

2. autotransformer	
3. auxiliary	
4. blockage	
5. contactor	
6. contribute	
7. eddy current	
8. electromechanical float switch	
9. excessive current	
10. increment	
11. intensity	
12. laminated core	
13. lock-rotor	
14. machinery	
15. magnetic control	
16. magnetic flux	
17. manual	
18. movable plunger	
19. nonlinear	
20. poor ventilation	
21. residual magnetism	
22. residual modulation	
23. solid core	
24. time eddy fuse	
25. time-invariant control system	
26. timing relay	

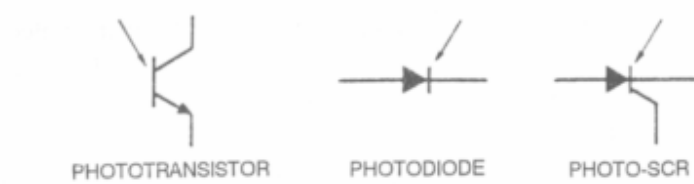


Figure 13-3. Schematic Symbols for the Phototransistor, the Photodiode, and the Photo-SCR.

When the photo transistor is in darkness, no electrons are emitted by the base junction, and the transistor is turned off. When the photo transistor is in the presence of light, it turns on and permits current to flow through the relay coil. The diode connected parallel to the relay coil is known as a kickback or freewheeling diode. Its function is to prevent induced voltage spikes from occurring when the current suddenly stops flowing through the coil and the magnetic field collapses.

In the circuit shown in Figure 13-4, the relay coil will turn on when the photo transistor is in the presence of light, and turn off when the phototransistor is in darkness. Some circuits may require the reverse operation. This can be accomplished by adding a resistor and a junction transistor to the circuit, Figure 13-5. In this circuit a common junction transistor is used to

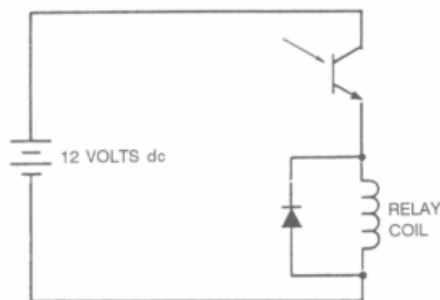


Figure 13-4. Photo transistor Controls Relay Coil.

control the current flow through the relay coil. Resistor R 1 limits the current flow through the base of the junction transistor. When the photo transistor is in darkness, it has a very high resistance. This permits current to flow to the base of the junction transistor and turn it on. When the photo transistor is in the presence of light, it turns on and connects the base of the junction

transistor to the negative side of the battery, This causes the junction transistor to turn off. The photo transistor in the circuit is used as a *stealer* transistor. A stealer transistor steals the base current away from some other transistor to keep it turned off.

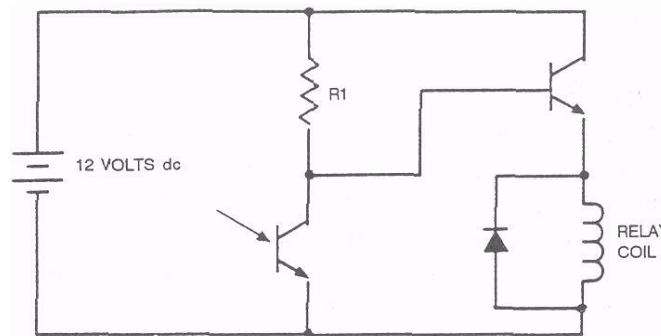


Figure 13-5. The Relay Turns on When the Photo transistor Is in Darkness.

Some circuits may require the photo transistor to have a higher gain than it has under normal conditions. This can be accomplished by using the photo transistor as the driver for a Darlington amplifier circuit, Figure 13-6. A Darlington amplifier circuit generally has a gain of over 10,000.

Photo diodes and photo-SCRs are used in circuits similar to those shown for the photo transistor. The photo diode will permit current to flow through it in the presence of light. The photo-SCR has the same operating characteristics as a common junction SCR. The only difference is that light is used to trigger the gate when using a photo-SCR.

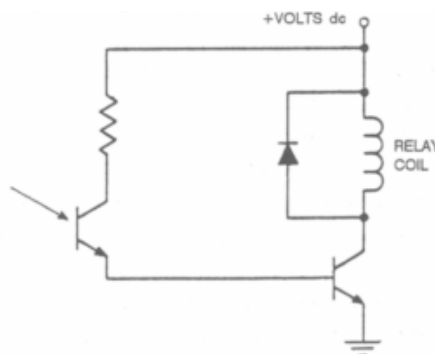


Figure 13-6. The Phototransistor Is Used as the Driver for a Darlington Amplifier.

Regardless of the type of photoemissive device used, or the type circuit it is used in, the greatest advantage of the photoemissive device is speed. A photoemissive device can turn on or off in a few microseconds. Photovoltaic or photoconductive devices generally require several milliseconds to turn on or off. This makes the use of photoemissive devices imperative in high speed switching circuits.

Photoconductive Devices

Photoconductive devices exhibit a change of resistance due to the presence or absence of light. The most common photoconductive device is the cadmium sulfide cell or cad cell. The cad cell has a resistance of about 50 ohms in direct sunlight and several hundred thousand ohms in darkness. It is generally used as a light sensitive switch. The schematic symbol for a cad cell is shown in Figure 13-7, Figure 13-8 shows a typical cad cell.

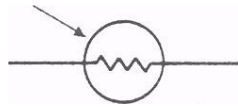


Figure 13-9. Schematic Symbol for a cad cell



Figure 13-8. Cad Cell (Courtesy EG & G Vactec , Inc.)

Figure 13-9 shows a basic circuit of a cad cell being used to control a relay. When the cad cell is in darkness, its resistance is high. This prevents the amount of current needed to turn the relay on from flowing through the circuit. When the cad cell is in the presence of light, its resistance is low. The amount of current needed to operate the relay can now flow through the circuit.

Although this circuit will work if the cad cell is large enough to handle the current, it has a couple of problems.

1. There is no way to adjust the sensitivity of the circuit. Photo-operated switches are generally located in many different areas of a plant. The surrounding light intensity can vary from one area to another. It is, therefore, necessary to be able to adjust the sensor for the amount of light needed to operate it.
2. The sense of operation of the circuit cannot be changed. The circuit shown in Figure 13-9 permits the relay to turn on when the cad cell is in the presence of light. There may be conditions that would make it desirable to turn the relay on when the cad cell is in darkness.

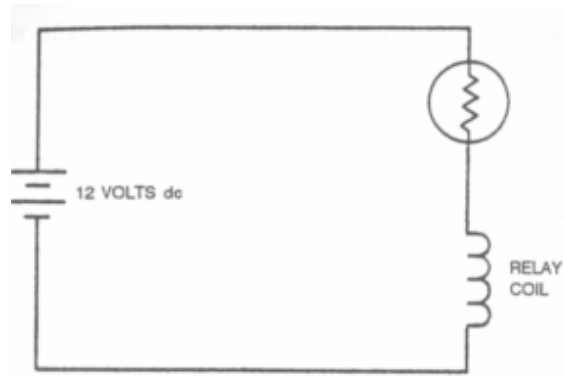


Figure 13-9. Cad Cell Controls Relay Coil.

Part I. Comprehension Exercises

A. Put “T” for true and “F” for false statements . Justify your answers.

- 1. Photodetectors sense the presence of an object by physical contact with the object.
- 2. Photovoltaic devices release electrons in the presence of light.
- 3. The amount of voltage produced by a solar cell depends on the material it is made of.
- 4. A large and a small solar cell made of the same material produce the same amount of voltage.
- 5. Photoemissive devices are of three types.
- 6. A phototransistor may not emit electrons in darkness.
- 7. Photovoltaic devices turn on and off faster than photoemissive devices.

B. Choose a, b, c, or d which best completes each item.

- 1. As we understand from the text, the increasing application of photodetectors in industry is due to their
 - a. speed and their ability to sense the presence or absence of almost any object
 - b. flexibility of being used in almost every type of industry
 - c. ability to sense objects without making any contact with the objects
 - d. all of the above
- 2. Paragraph three mainly discusses solar cells.
 - a. the amount of voltage produced by
 - b. the amount of current produced by

- c. the characteristics of
- d. the characteristics of the material used in
- 3. In a circuit consisting of a resistor, a junction transistor, a diode, a photo transistor, and a relay
 - a. the relay coil will turn on when the photo transistor is in the presence of light
 - b. the relay coil will turn on when the photo transistor is in darkness
 - c. the photo transistor exhibits a high resistance in the presence of light
 - d. the photo transistor emits more electrons in the presence of light
- 4. The junction transistor used in the circuit controls the current flow through.....
 - a. the relay coil
 - b. the resistor
 - c. the diode
 - d. the phototransistor
- 5. It is true that a cad cell
 - a. is not usually used as a light sensitive switch
 - b. is usually used to produce high currents
 - c. acts much faster than photovoltaic cell
 - d. exhibits lower resistance in the presence of sunlight

C. Answer the following questions orally.

1. What does the amount of current produced by a solar cell depend on?
2. How can the amount of voltage produced by solar cells be increased?
3. What are the three categories of photo-operated devices?
4. What is the difference between photo voltaic and photoemissive devices?
5. What is the function of the kickback diode in the circuit?
6. What are the problems of a cad cell circuit?
7. What is the advantage of photovoltaic cells over other photo-operated devices?

Part II. Language Practice

A. Choose a, b, c, or d which best completes each item.

1. Photodiode current is very low. If this current is to be used effectively in control applications, it must be amplified by an external current amplifier or by a device called

a. a photodetector	b. a phototransistor
c. a photodiode	d. a photoconductor
2. The collector-base junction of the phototransistor acts as

a. a transistor	b. a resistor
c. a photodetector	d. a photodiode

3. Solar cells often called devices are made of silicon and have the ability to produce a voltage in the presence of light.
 - a. photovoltaic
 - b. photoemissive
 - c. photodiode
 - d. photoconductive
4. Current developed by a phototransistor is dependent mainly on of light and very little on the applied voltage.
 - a. the frequency
 - b. the wavelength
 - c. the intensity
 - d. both a and b
5. The ability of the low-current gate circuit of to control large amounts of power in its anode-cathode circuit makes this device particularly useful in industrial electronics.
 - a. photo-SCR
 - b. photodetector
 - c. photodiode
 - d. phototransistor

B. Fill in the blanks with the appropriate form of the words given.

1. Industry

- a. Photodetectors designed for use are made to be mounted and used in different ways.
- b. Optoelectronics deals with light-sensitive semiconductor devices used in

2. Dope

- a. Photovoltaic cells consist of a single semiconductor crystal which has been with both N- and P-type materials.
- b. Photodiodes and phototransistors are each sensitive to a specific range of light frequencies. What these frequencies are, is determined by the materials from which they are constructed and by the extent of of the junctions.
- c. When light falls on the PN junction, which is the boundary of these, a voltage appears across the junction.

3. Exhibit

- a. Photodiodes and phototransistors..... unique spectral characteristics and have a variety of uses in industry.
- b. Selective properties within the range of the visible spectrum areby the average human eye.
- c. Photoconductive cellsthe particular property that their

resistance decreases in the presence of light and increases in the absence of light.

4. Contact

- a. The switch of relays may be normally open or normally closed.
- b. The relay coil and terminals can usually be located by inspection.

5. Sense

- a. The current that a photovoltaic cell can deliver to a load depends on the area of the light- material which makes up the cell.
- b. The of a photocell to a particular color depends on the nature of the material from which the cell is constructed, and on the manner of its construction,
- c. Manufacturers, specify the spectral of their devices in the form of a frequency response curve.

C. Fill in the blanks with the following words.

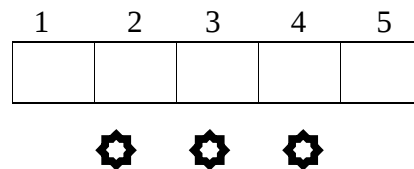
wavelengths	frequency	violet
associated	human	color
nanometer	range	eye

Light behaves like an electromagnetic radiation, and with light are the characteristics of wavelength and The wavelength of light determines the of that light. White light consists of many which may be separated by a prism. The eye responds to the prismatic colors ranging from to red . The wavelengths of light which the can 'see' are in the nanometer of 400 nm to 700 nm. A is a 10^{-9} part of a meter.

D. Put the following sentences in the right order to form a paragraph. Write the corresponding letters in the boxes provided.

- a. A glass window in the housing permits light to fall on the active material of the cell.
- b. Photoconductive cells are made of a thin layer of semiconductor material such as cadmium sulfide, cadmium selenide, or lead sulfide.
- c. The cell simply acts as a conductor whose resistance changes when illuminated.

- d. The semiconductor layer is enclosed in a sealed housing.
- e. Photoconductive cells exhibit the peculiar property that their resistance decreases in the presence of light and increases in the absence of light.



Section Two: Further Reading

Optoelectronics-The Phototransistor

There is a variety of semiconductor-junction light-sensitive devices which fall under the heading of *optoelectronic* devices. These devices are being used, in conjunction with other semiconductors, in such diverse control applications as automatic light level controls in photocopy machines to computer control of machine tools.

Included in the optoelectronics' family are: light-emitting diodes (LEDs), photodiodes, phototransistors, photo-SCRs [also called *light-activated* SCRs (LASCRs)], optocouplers or optoisolators, and solid-state relays (SSRs).

Photodiodes

Silicon photodiodes are light-sensitive devices, also called *photodetectors*, which convert light signals into electrical signals. A window or lens permits light to fall on the junction (Figure 13-10). When light shines on the *reversebiased* PN photodiode junction, hole-electron pairs are created. The

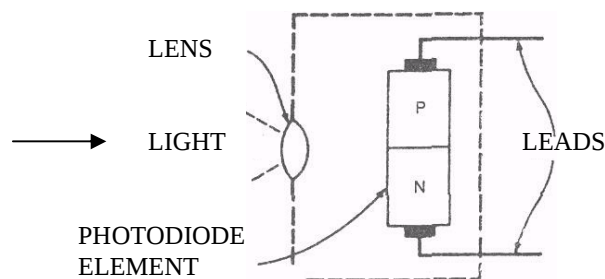


Figure 13-10. Light Falls on the PN Junction of a Photodiode.

movement of these hole-electron pairs in a properly connected circuit results in current flow. The current is proportional to the intensity of light and is also affected by the frequency of the light falling on the photojunction,

The response of the human eye is not uniform in the range of the visible spectrum. As Figure 13-11 shows, the eye is most sensitive to light whose wavelength is 550 nm and falls off to 400 and 700 nm. The spectral response of the eye, then, is 400 to 700 nm, peaking at 550 nm.

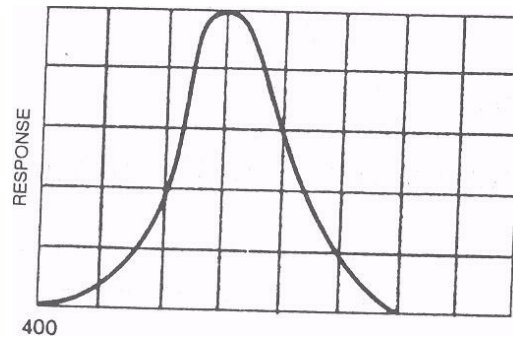


Figure 13-11. Spectral Response of the Human Eye. (To convert wavelength to angstrom units, multiply above values by 10.)

The spectral response of a silicon photodiode is shown in Figure 13-12. We see that maximum sensitivity is to radiation at 900 nm, and that the total response is in the range from 400 to 1100 nm. This includes response both *in*

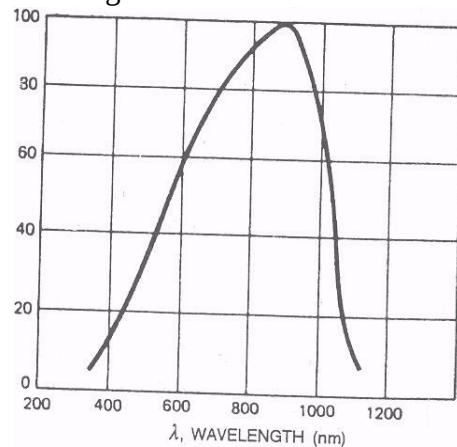


Figure 13-12. Spectral Response of a Silicon Photodiode (Motorola).

the visual range and *outside* it. We call the total range 'light' even though light normally refers to the frequencies within the visual(human)spectrum. The spectral response curve of a particular silicon photodiode depends on the geometry and the extent of doping of the junction. It is apparent, then, for maximum efficiency, that the spectral characteristics of the light source (light emitter) used with a photodiode must match the characteristics of the photodiode. So in the case of the photodiode in Figure 13-12, the light source for it must have a wavelength close to 900 nm.

Phototransistors

The current developed by a photodiode is very low. This current cannot be used directly in control applications but must be amplified. After amplification, the photocurrent may be high enough to be used in a control system for example, to set a relay. The phototransistor is a light detector which combines a photodiode and a transistor amplifier. Figure 13-13 shows an NPN phototransistor. Here a lens concentrates the light on a very thin P-type wafer, sandwiched in between an N-type collector and an emitter. Although the phototransistor has three sections, only two leads may issue from the housing, the emitter and collector leads. In this device base current is supplied

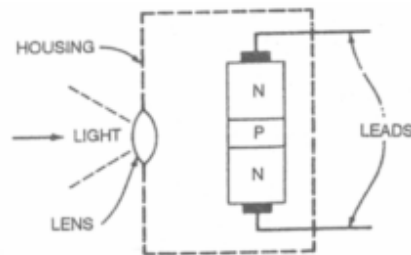


Figure 13-13. NPN Phototransistor

by the current created by the light falling on the base-collector photodiode junction. Some phototransistors have a third lead issuing from the housing. In such a phototransistor, base bias is provided from an external circuit, on which the photodiode current is superimposed.

Current in a phototransistor is dependent mainly on the intensity of light entering the transistor window and is little affected by the voltage applied to the external circuit. Figure 13-14 is a graph of collector current I_c , as a function of collector-emitter voltage V_{CE} and as a function of illumination H . It is evident that the phototransistor acts as a constant-current source, and

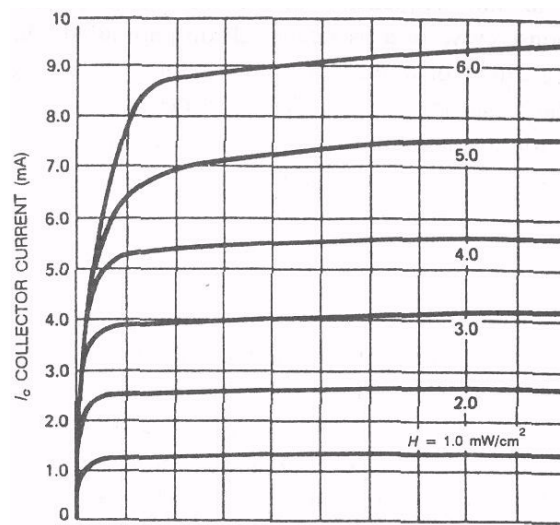


Figure 13-14. Collector Characteristic for the MRD300 (Motorola).

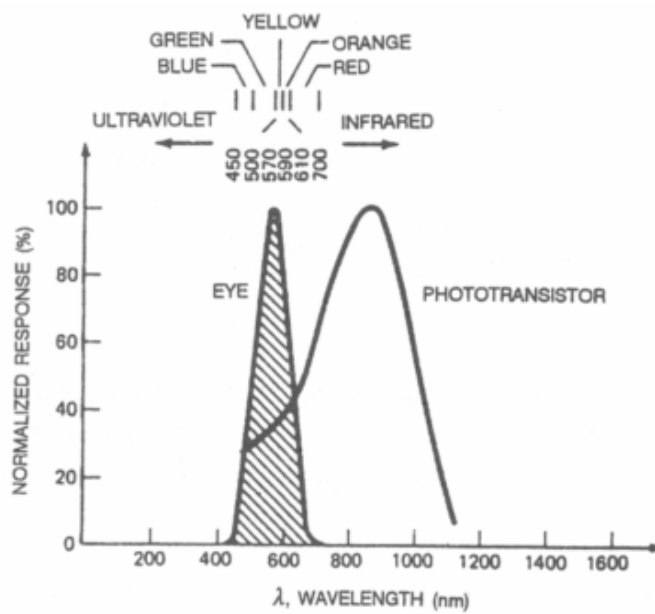


Figure 13-15. Spectral Responses of a Phototransistor and the Human Eye (General Electric).

Figure 13-15 is the frequency response curve of a phototransistor for comparison with the response of the human eye. Peak sensitivity of the

phototransistor is at about 900 nm.

The angular alignment of a phototransistor and its source of illumination are important considerations. The reason is that the illumination of the photojunction is proportional to the cosine of the angle between the direction of radiation of the light source and the perpendicular to the surface of the photojunction. In addition, the optical lens, or window, and its size further affect the response of the phototransistor.

Phototransistor Relay

Figure 13-16 is the circuit diagram of a one-stage relay employing an NPN

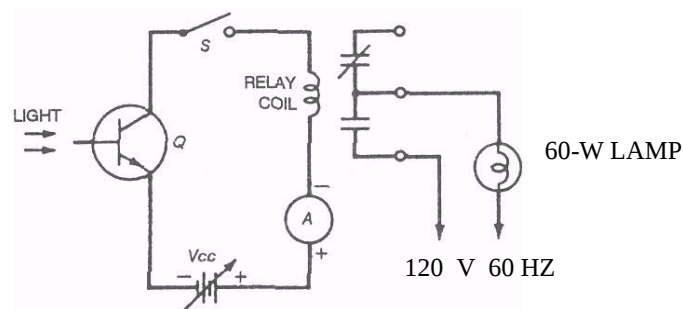


Figure 13-16. One-Stage Phototransistor-Operated Relay.

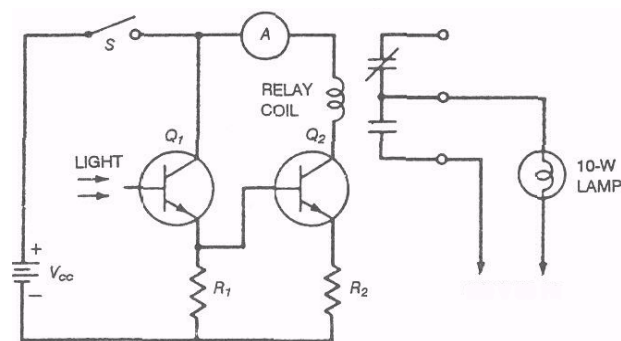


Figure 13-17. Two-Stage Phototransistor and Amplifier.

phototransistor Q_1 . A very sensitive relay is connected in the collector circuit and is actuated when a strong light is focused directly on the transistor lens. The resistance and the pickup current of the relay must be relatively low so that the transistor may be operated within its rated characteristics. A is a milliammeter used for measuring collector current. If an ordinary 60-W light bulb is used as the light source, the light must be held fairly close to the

photowindow to actuate the relay.

Less light is required to actuate the same relay in the circuit shown in Figure 13-17. Hence the light source need not be held so close to the light window. In this circuit, the output of phototransistor Q1 is amplified by transistor Q2. The relay coil is connected in the collector of Q2. Photocurrent flow in Q1 through *R1*, biases Q2. In the absence of photocurrent (that is, in the absence of light on Q1), Q2 is biased to cutoff, since the emitter and the base of Q2 are at the same potential. When light shines on Q1 emitter current flows through *R1* and forward-biases the emitter to base of Q2, thus causing current to flow in Q2. If the light is strong enough, there will be sufficient current in the collector of Q2 to actuate the relay. As in the preceding circuit, a sensitive relay is required.

Comprehension Exercises

A. Choose a, b, c, or d which best completes each item.

1. The third paragraph mainly describes
 - a. how the intensity of light affects current
 - b. how the frequency of light affects current
 - c. the movement of hole-electron pairs
 - d. the mechanism of photodiodes
2. As we understand from the text,
 - a. the spectral characteristics of various photodiodes are different
 - b. the spectral characteristics of all photodiodes are similar
 - c. the visual spectrum extends outside the spectral response of a photodiode
 - d. the efficiency of a photodiode does not relate to the light source
3. The higher the doping of a particular photodiode junction
 - a. the higher the number of hole-electron pairs created
 - b. the lower the number of hole-electron pairs created
 - c. the higher the efficiency of its spectral response
 - d. the lower the efficiency of its spectral response
4. It is true that
 - a. the photodiode current is amplified to be used in a control system
 - b. the primary source of the photodiode current is an external circuit
 - c. a photo transistor consists of two transistor amplifiers
 - d. a photo transistor consists of a collector and an emitter
5. The last two paragraphs mainly discuss
 - a. relay circuits with one or two transistors

- b. milliammeters used in different relay circuits
 - c. the kinds of relays used in a circuit
 - d. the number of transistors employed in a relay circuit
6. We may infer from the text that
- a. an optoelectronic device is an electronic device containing optic and electric ports
 - b. an optoelectronic device cannot be used as a control element in complicated machines
 - c. the current developed by a photodiode is high enough to be used directly in control applications
 - d. the spectral response of the eye can be extended to the spectral response of a silicon photodiode

B. Write the answers to the following questions.

1. What are some applications of optoelectronic devices?
2. What is a photodetector?
3. How does a photodiode create current?
4. What is the current developed by a photodiode proportional to?
5. How does the spectral response of the eye differ from that of a photodiode?
6. What is known as light?
7. What brings about the maximum efficiency of the spectral response of a particular photodiode?
8. What does a phototransistor consist of?
9. What is the illumination of the photojunction proportional to?
10. How do you describe the mechanism of a relay circuit having two transistors?



Section Three: Translation Activities

A. Translate the following passage into Persian.

Photovoltaic Cells

These are also light-sensitive semiconductor devices, but they produce a voltage when illuminated, which increases as the intensity of light falling on

the semiconductor junction of this two-element cell increases. The usual basic material from which these cells are made today is silicon or selenium.

Photovoltaic cells convert light into electric energy, which may be used directly to supply small amounts of electric power for electrically powered devices. Because of the low levels of power which photovoltaic cells generate, they have been used in the past in low-power devices such as light meters and photographic exposure meters. However, with an improvement in the efficiency of these cells, more power has been produced, as in solar cells, which are photovoltaic devices. Solar cells appear destined to play a substantial role in the development of new sources of energy.

Photovoltaic cells consist of a single semiconductor crystal which has been doped with both N- and P-type materials. When light falls on the PN junction, which is the boundary of these dopants, a voltage appears across the junction. About 0.6 V is developed by the photovoltaic cell in bright sunlight. The amount of power the cell can deliver depends on the extent of its active surface. An average cell will produce about 30 milliwatts per square inch (30 m W/in^2) of surface, operating into a load of 4Ω . To increase the power output, large banks of cells are used in series and parallel combinations. An example is the use of solar cells to power the experimental circuits on lunar and space modules.

B. Find the Persian equivalents of the following terms and expressions and write them in the spaces provided.

- 1 . angular alignment
- 2 . collapse
- 3 . Darlington amplifier circuit
- 4 . dope
- 5 . evident
- 6 . exhibit
- 7 . lexibility
- 8 . illumination
- 9 . light activated SCR (LASCR)
- 10 . light emitting diode (LED)
- 11 . optocoupler
- 12 . optoelectronic
- 13 . optoisolator
- 14 . perpendicular
- 15 . photocell

16. photodetector	
17. photodiode	
18. photoemissive	
19. photojunction	
20. photo-operated device	
21. photo-operated switch	
22. photovoltaic	
23. photowindow	
24. reverse biased	
25. semiconductor	
26. solid-state relay (SSR)	
27. spectrum	
28. stealer transistor	
29. trigger	
30. wafer	
31. wavelength	

Unit 14

Learning Comprehension

Instrument Classification and Characteristics

Classification

Instruments are subdivided into separate classes according to several characteristics. These classifications are useful in broadly establishing several characteristics of particular instruments such as accuracy, cost, and general intended applications.

Instrument Types

Instruments are divided into active or passive ones according to whether the energy for measurement is entirely produced by the quantity being measured or the quantity being measured simply modulates the magnitude of an external energy source. This is illustrated by examples.

A passive instrument is the pressure-measuring device shown in Figure 14-1. The pressure of the fluid is translated into a movement of a pointer. The energy expended in moving the pointer is derived from an external source. The change in pressure measured: there are no other energy

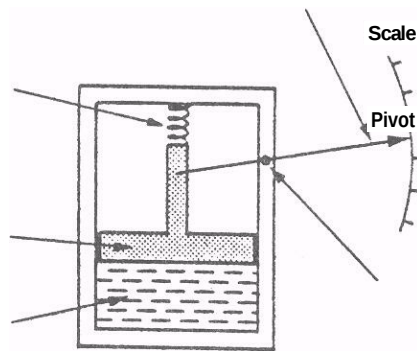


Figure 14-1: Passive Pressure Gauge.

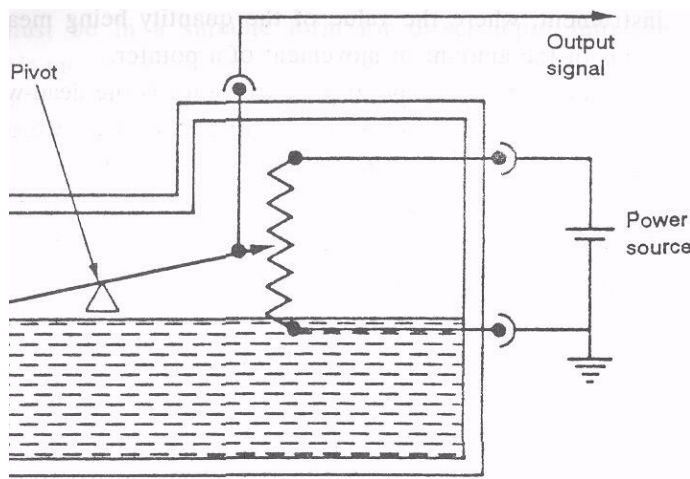


Figure 14-2. Petrol-Tank Level Indicator.

An active instrument is a float-type petrol-tank level indicator shown in Figure 14-2. Here, the change in petrol level moves a float, and the output signal consists of a proportion of the voltage applied across the two ends of the potentiometer. The output signal comes from the external power source; the float system is merely modulating the value of the voltage from the power source.

Thus, the external power source is usually in electrical form, but it can be other forms of energy such as a pneumatic or hydraulic pressure.

The main difference between active and passive instruments is the measurement resolution which can be obtained. With the simple deflection type instrument, the amount of movement made by the pointer for a change is closely defined by the nature of the instrument. To increase measurement resolution by making the pointer travel a longer arc, the pointer tip moves through a longer arc, the scope for error is clearly limited by the practical limit of how long the instrument can be. In an active instrument, however, adjustment of the external energy input allows much greater control over measurement resolution.

Active Instruments

The petrol-tank level indicator just mentioned is a good example of a deflection type of instrument.

instrument, where the value of the quantity being measured is displayed in terms of the amount of movement of a pointer.

An alternative type of pressure gauge is the dead-weight gauge shown in Figure 14-3 which is a null-type instrument. Here, weights are put on top of the piston until the downward force balances the fluid pressure. Weights are added until the piston reaches a datum level, known as the null point. Pressure measurement is made in terms of the value of the weights needed to reach this null position.

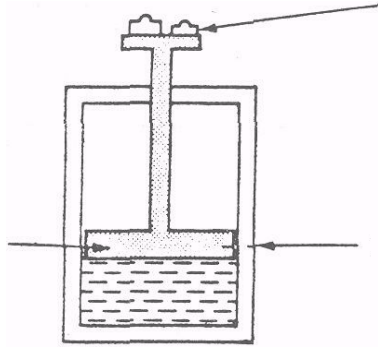


Figure 14-3. Dead-Weight Pressure Gauge.

The accuracy of these two instruments depends on different things. For the first one, it depends on the linearity and calibration of the spring, while for the second, it relies on the calibration of the weights. As calibration of weights is much easier than careful choice and calibration of a linear-characteristic spring, this means that the second type of instrument will normally be the more accurate. This is in accordance with the general rule that null-type instruments are more accurate than deflection types.

Monitoring/Control Instruments

An important distinction between different instruments is whether they are suitable only for monitoring functions or whether their output is in a form that can be directly included as part of an automatic control system. Instruments which only give an audio or visual indication of the magnitude of the physical quantity measured, such as a liquid-in-glass thermometer, are only suitable for monitoring purposes. This class normally includes all null-type instruments and most passive transducers.

For an instrument to be suitable for inclusion in an automatic control

system, its output must be in a suitable form for direct input into the controller. Usually, this means that an instrument with an electrical output is required, although other forms of output such as optical or pneumatic signals are used in some systems.

Analog/Digital Instruments

An analog instrument gives an output which varies continuously as the quantity being measured changes. The output can have an infinite number of values within the range that the instrument is designed to measure. The deflection type of pressure gauge is a good example of an analog instrument. As the input value changes, the pointer moves with a smooth continuous motion. While the pointer can therefore be in an infinite number of positions within its range of movement, the number of different positions which the eye can discriminate between is strictly limited, this discrimination being dependent upon how large the scale is and how finely it is divided.

A digital instrument has an output which varies in discrete steps and so can only have a finite number of values. The rev-counter sketched in Figure 14-4 is an example of a digital instrument. A cam is attached to the revolving body whose motion is being measured, and on each revolution the cam opens and closes a switch. The switching operations are counted by an electronic counter. This system can only count whole revolutions and cannot discriminate any motion which is less than a full revolution.

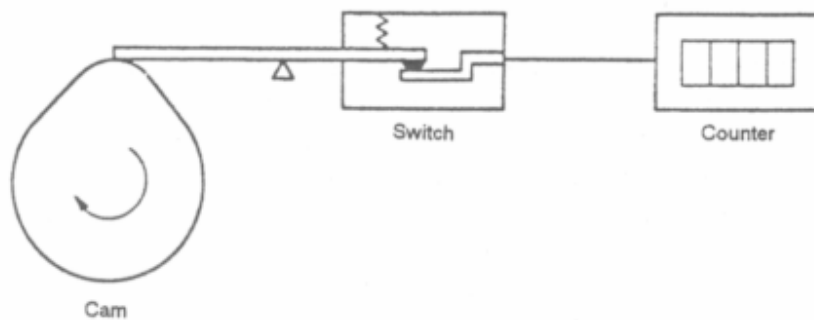


Figure 14-4. Rev-Counter.

The distinction between analog and digital instruments has become particularly important with the rapid growth in the application of micro-

computers to automatic control systems. Any digital computer system, of which the microcomputer is but one example, performs its computations in digital form. An instrument whose output is in digital form is therefore particularly advantageous in such applications, as it can be interfaced directly to the control computer. Analog instruments must be interfaced to the microcomputer by an analog-to-digital (*ND*) converter, which converts the analog output signal from the instrument into an equivalent digital quantity which can be read into the computer.

Part I. Comprehension Exercises

A. Put "T" for true and "F" for false statements. Justify your answers.

- 1. The output of a passive instrument is directly produced by the quantity being measured.
- 2. The output of an active instrument is determined by an external power source.
- 3. In passive instruments, the external power source can be various forms of energy.
- 4. In a passive pressure gauge, the energy used to move the pointer is derived from the change in pressure of the fluid measured.
- 5. A deflection-type instrument is more accurate than a null-type instrument.
- 6. An analog computer does its calculations one step at a time whereas a digital computer continuously works out its calculation

B. Choose a, b, c, or d which best completes each item.

- 1. The function of the float system in a petrol-tank level indicator is
 - a. to modulate the value of the voltage from the external power source
 - b. to evaluate the amount of energy from the external power source
 - c. to calculate the amount of the external voltage applied across the two ends of the potentiometer
 - d. to adjust the energy produced by the quantity being measured
- 2. It is understood from the text that the level of measurement resolution obtained by
 - a. passive instruments can be highly under control
 - b. active instruments can be highly under control
 - c. a simple pressure gauge cannot be increased
 - d. a simple pressure gauge can be highly increased

3. It may be inferred from the text that
 - a. the scope for improving measurement resolution is infinite
 - b. the scope for improving measurement resolution is not infinite
 - c. active and passive instruments can both be used for accurate measurements
 - d. active and passive instruments do not vary in their basic structures
4. The mechanism of a null-type instrument is based on
 - a. the force applied by weights
 - b. the linearity of a spring
 - c. the movement of a pointer
 - d. the state of equilibrium
5. It is true that active instruments
 - a. are only suitable for monitoring purposes
 - b. are not as effective for measuring purposes as passive ones
 - c. can be used for control purposes
 - d. cannot be used for control purposes
6. It can be inferred *from* the last paragraph that
 - a. the time involved in the process of converting an analog signal to a digital quantity can be critical in the control of fast processes
 - b. the time involved in the process of converting an analog signal to a digital quantity is too small to be considered a disadvantage of the analog instrument
 - c. the analog to digital converter does not affect the speed of operation of the control computer
 - d. the analog to digital converter does not affect the accuracy of the control computer

C. Answer the following questions orally.

1. How do you describe the mechanism of a passive pressure gauge?
2. What constitutes the potentiometer in the petrol-tank level indicator?
3. How are null-type instruments different from deflection types?
4. Why is an analog instrument interfaced to a microcomputer by an A/D converter?
5. How does a rev-counter work?

Part II. Language Practice

A. Choose a, b, c, or d which best completes each item.

1. A passive transducer.....
 - a. can adjust the magnitude of some external power source
 - b. can transform the energy obtained from an external source

- c. has no source of power other than the input signals
- d. may have either internal or external sources of power
- 2. A three-terminal rheostat, or a resistor with one or more adjustable sliding contacts, that functions as an adjustable voltage divider is called
 - a. a potentiometer
 - b. a converter
 - c. a gauge
 - d. a counter
- 3. The energy in the output signal of is entirely produced by the quantity being measured.
 - a. a passive pressure gauge
 - b. a petrol-tank level indicator
 - c. a deflection type instrument
 - d. a digital instrument
- 4. A instrument can be directly included as part of an automatic control system.
 - a. null-type
 - b. passive
 - c. control
 - d. monitoring
- 5. The energy in the output signal of instruments comes from an external source of power.
 - a. passive
 - b. active
 - c. monitoring
 - d. null-type

B. Fill in the blanks with the appropriate form of the words given.

1. Require

- a. Choice between active and passive instruments for a particular application involves carefully balancing the measurement-resolution against cost.
- b. The higher the gate current, the lower the anode-cathode voltage to turn the SCR on.
- c. The current passing through the coil a definite time interval to reach its maximum or steady-state value.
- d. A zener diode may be used as a voltage regulator for a load a voltage equal to the zener voltage.

2. Obtain

- a. Angular velocity measurements can be by differentiating the output signal from angular displacement transducers.
- b. In measurement systems which contain an angular acceleration transducer, such as a gyro accelerometer, it is possible to a velocity measurement by integrating the acceleration measurement signal.

3. Consist

- a. The drag-cup tachometer has a central spindle carrying a permanent magnet which rotates inside a non-magnetic drag-cup of a cylindrical sleeve of electrically conductive material.
- b. A bistable multivibrator of two direct cross-coupled dc amplifiers.

4. Accurate

- a. The drag-cup tachometer has a typical measurementof $\pm 0.5\%$ and is commonly used in the speedometers of motor vehicles and as a speed indicator for aero engines.
- b. If the input data entering the computer are correct and if the program of instructions is reliable, then we can expect that the computer generally will produce output.

5. Use

- a. A silicon rectifier, when to convert ac to dc, acts as a closed switch when its anode is positive relative to its cathode and as an open switch when its anode is negative relative to its cathode.
- b. The vertical amplifiers of an oscilloscope must be calibrated if the scope is to be for measuring the amplitude of waveforms.
- c. The ability of the low-current gate circuit of an SCR to control large amounts of power in its anode-cathode circuit makes this device particularly in industrial electronics.

C. Fill in the blanks with the following words.

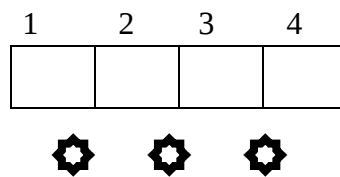
squirrel-cage	drag-cup	speed	tachometer
accuracy	with	measures	analog

Probably the most common form of output device used is the dcThis is a relatively simple device which speeds up to about 5000 rpm an accuracy of $\pm 1\%$. Where better is required within a similar range of..... measurement, ac tachometers are used. The rotor type has an accuracy of $\pm 0.25\%$ and rotor types have accuracies up to $\pm 0.05\%$.

D. Put the following sentences in the right order to form a paragraph. Write the corresponding letters in the boxes provided.

- a. It consists of a pair of spherical balls pivoted on a rotating shaft.

- b. The pointer can be arranged to give a visual indication of speed by causing it to move in front of a calibrated scale, or its motion can be converted by a translational displacement transducer into an electrical signal.
- c. These balls move outward under the influence of centrifugal forces as the rotational velocity of the shaft increases and lift a pointer against the resistance of a spring.
- d. The mechanical flyball is a velocity-measuring instrument which was first developed many years ago and is still used extensively in speed-governing systems for engines, turbines, etc.



Section Two: Further Reading

Rotational Velocity Measurement

The main application of rotational velocity transducers is in speed control systems. They also provide the usual means of measuring translational velocities, which are transformed into rotational motions for measurement purposes by suitable gearing. Many different instruments and techniques are available for measuring rotational velocity as presented below.

DC Tachometric Generator

The dc tachometric generator, or dc tachometer as it is generally known, has an output which is approximately proportional to its speed of rotation. Its basic structure is identical to that found in a standard dc generator used for producing power, and is shown in Figure 14-5. Both permanent-magnet types and separately excited field types are used. However, certain aspects of the design are optimized to improve its accuracy as a speed-measuring instrument. One significant design modification is to reduce the weight of the rotor by constructing the windings on a hollow fiberglass shell. The effect of this is to minimize any loading effect of the instrument on the system being measured.

The dc output voltage from the instrument is of a relatively high

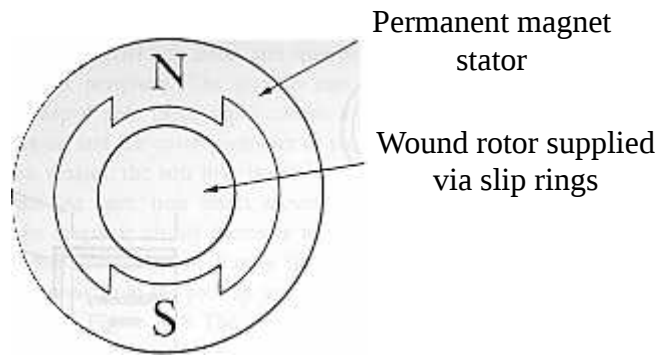


Figure 14-5. The DC Tachometer.

magnitude, giving a high measurement sensitivity which is typically 5 V per 1000 rpm. The direction of rotation is determined by the polarity of the output voltage. A common range of measurement is 0-5000 rpm. Maximum non-linearity is usually about $\pm 1\%$ of the full-scale reading.

One problem with these devices which can cause problems under some circumstances is the presence of an ac ripple in the output signal. The magnitude of this can be up to 2% of the output dc level.

AC Tachometric Generator

The ac tachometric generator, or ac tachometer as it is generally known, has an output approximately proportional to rotational speed, as in the dc tachogenerator. Its mechanical structure takes the form of a two-phase induction motor, with two stator windings and (usually) a drag-cup rotor, as shown in Figure 14-6. One of the stator windings is excited with an ac voltage and the measurement signal is taken from the output voltage induced in the second winding. The magnitude of this output voltage is zero when the rotor is stationary, and otherwise proportional to the angular velocity of the rotor. The direction of rotation is determined by the phase of the output voltage, which switches by 180° as the direction reverses. Therefore, both the phase and magnitude of the output voltage have to be measured. A typical range of measurement is 0-4000 rpm with an accuracy of $\pm 0.05\%$ of full-scale reading.

While the form of ac tachometer described above is the commonest one, a second form also exists. This has a squirrel-cage rather than a drag-cup rotor and so is cheaper. Its structure and mode of operation are otherwise identical, but the measurement accuracy is reduced.

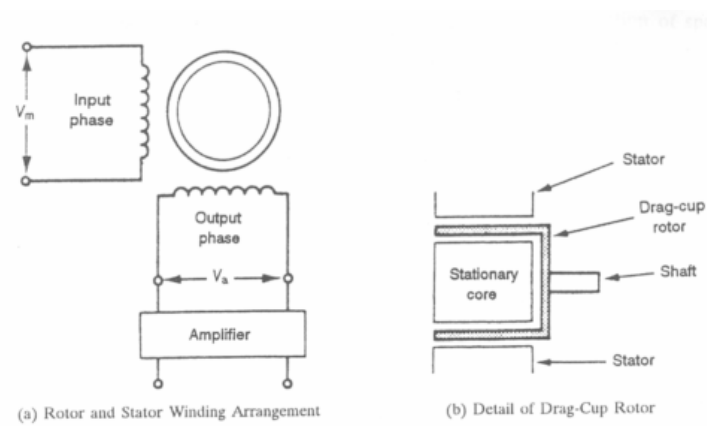


Figure 14-6. The A C Tachometer.

Variable-Reluctance Velocity Transducer

The form of a variable-reluctance transducer is shown in Figure 14-7. It can be seen that this consists of two parts, a rotating disk connected to the moving

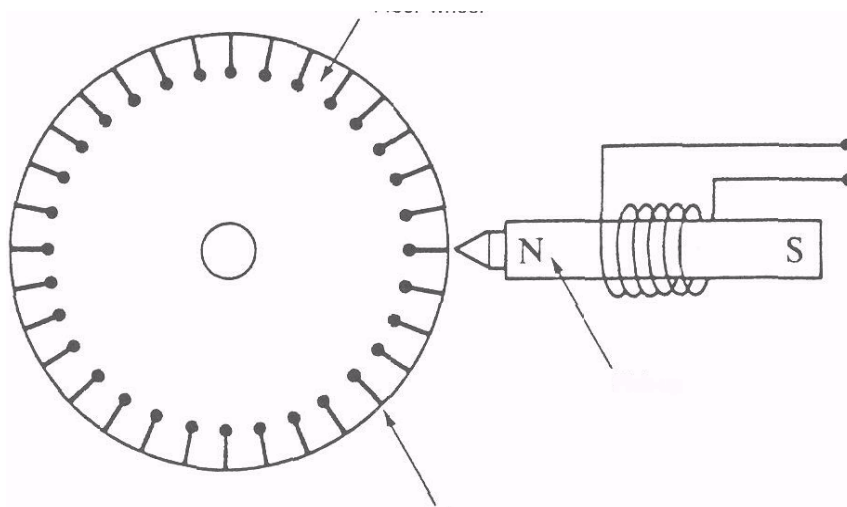


Figure 14-7. Variable-Reluctance Transducer.

body being measured and a pick-up unit. The rotating disk is constructed from a bonded-fiber material into which soft iron poles are inserted at regular intervals around its periphery. The pick-up unit consists of a permanent magnet with a shaped pole piece which carries a wound coil. The distance between the pick-up and the outer perimeter of the disk is around 0.5 mm.

As the disk rotates, the soft iron inserts on the disk move in turn past the pick-up unit. As each iron insert moves toward the pole piece, the reluctance of the magnetic circuit increases and hence the flux in the pole piece also increases. Similarly, the flux in the pole piece decreases as each iron insert moves away from the pick-up unit, and the pattern of flux changes with time as shown in Figure 14-8. The changing magnetic flux inside the pick-up coil causes a voltage to be induced in the coil whose magnitude is proportional to the rate of change of flux. This voltage is positive while the flux is increasing and negative while it is decreasing and therefore its variation with time is as shown in Figure 14-9.

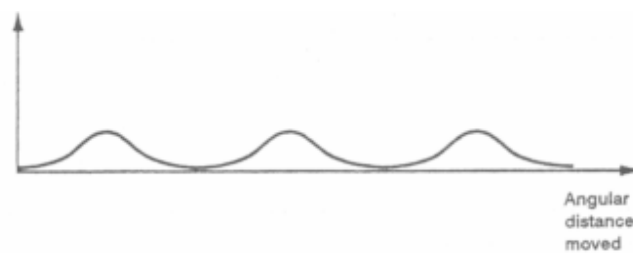


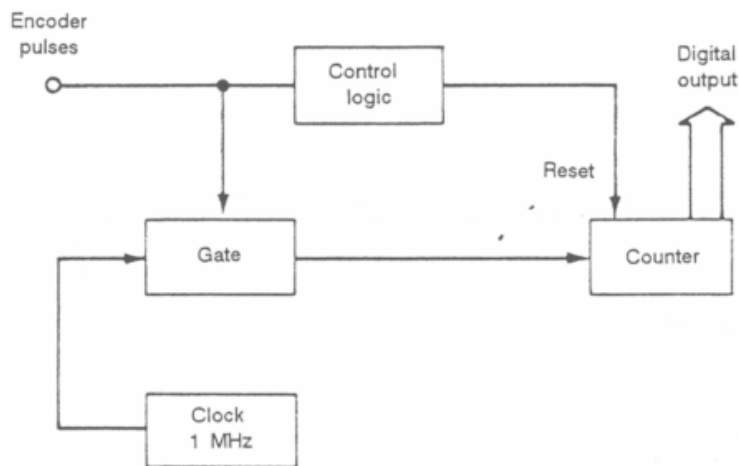
Figure 14-8. Pattern of Flux Change With Rotation of a Variable-Reluctance Transducer.



Figure 14-9. Pattern of Induced Voltage Change With Rotation of a Variable-Reluctance Transducer.

The form of this output can be regarded as a sequence of positive and negative pulses whose frequency is proportional to the rotational velocity of the disk. This can be converted into an analog, varying-amplitude, dc voltage output by means of a frequency-to-voltage converter circuit connected to the output terminals of the pick-up. However, greater measurement accuracy can be obtained by converting the output waveform into sharp pulses which are counted by an electronic counter. It is normal procedure to produce the pulses at each instant that the induced voltage in the coil changes sign as it passes through zero. This is achieved by electronic means.

The maximum angular velocity which the instrument can measure is limited because of the finite width of the induced pulses. As the velocity increases, the distance between the pulses is reduced, and at a certain velocity the pulses start to overlap. At this point, the pulse counter ceases to be able to distinguish the separate pulses.



The total pulse count measured over a certain length of time only gives information about the average velocity over that period. Measurement of the actual velocities at the instants of time that each output pulse occurs can be achieved by the scheme shown in Figure 14-10. In this circuit, the pulses from the transducer gate the train of pulses from a 1 MHz clock into a counter, Control logic resets the counter and updates the digital output value after receipt of each pulse from the transducer. The measurement resolution of this system is highest when the speed of rotation is low.

Photoelectric Pulse-Counting Methods

An alternative to the variable-reluctance transducer, but which uses very similar principles to it, is the method where pulses are produced by photoelectric techniques and counted. These pulses are generated by one of the two alternate methods illustrated in Figure 14-11. In Figure 14-11(a), the pulses are produced as the windows in a slotted disk pass in sequence between a light source and a detector. The alternate form, Figure 14-11(b), has both light source and detector mounted on the same side of a reflective disk which has black sectors painted onto it at regular angular intervals. In either case, the pulses are counted by an electronic counter. The frequency of the pulses is proportional to the angular velocity of the body connected to the rotating disk. Pulses generated in this manner are narrower than those generated by magnetic means and so the instrument is capable of measuring higher velocities.

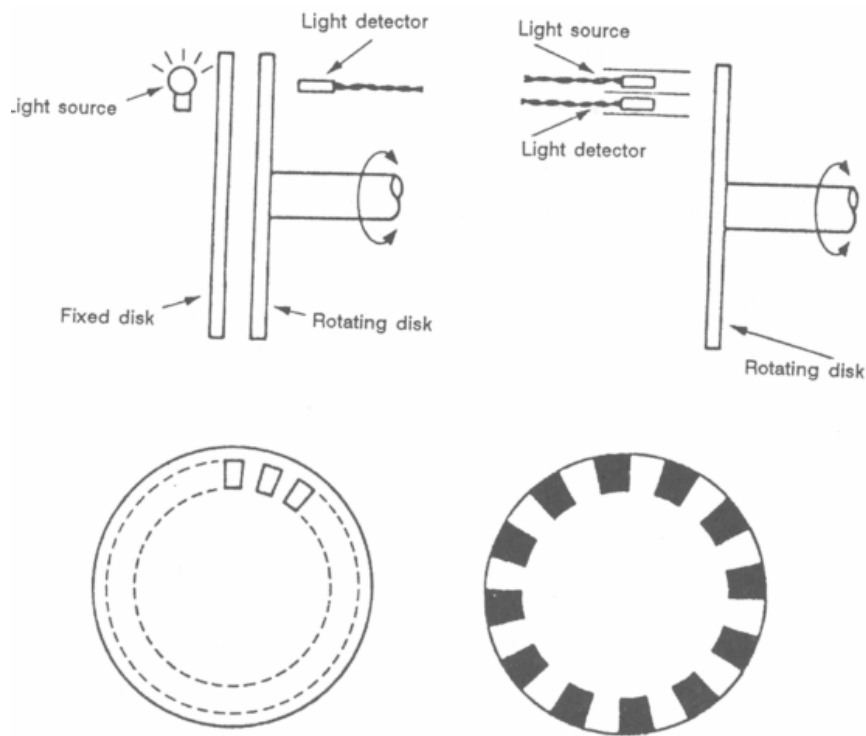


Figure 14-10. Scheme to Measure Instantaneous Angular Velocities.

5. How can an analog, varying amplitude, dc voltage be obtained from a variable-reluctance transducer?
6. What does the finite width of the induced pulses in a variable-reluctance transducer lead to?
7. How can actual velocities at the instants of time that each output pulse occurs be measured?
8. How do you describe the photoelectric pulse generation techniques?

Section Three: Translation Activities

A. Translate the following passage into Persian.

Stroboscopic Methods

The stroboscopic technique of rotational velocity measurement operates on a similar physical principle to the variable-reluctance and photoelectric pulse-counting methods, except that the pluses involved consist of flashes of light generated electronically and whose frequency is adjustable so that it can be matched with the frequency of occurrence of some feature on the rotating body being measured. This feature can either be some naturally occurring one such as the spokes of a wheel or gear teeth, or it can be an artificially created pattern of black and white stripes. In either case, the rotating body appears stationary when the frequencies of the light pulses and body features are in synchronism. Flashing rates up to 25,000 per minute are available from commercial stroboscopes, according to the range of velocity measurement required, and the typical measurement accuracy obtained is $\pm 1\%$ of the reading.

Measurement of the flashing rate at which the rotating body appears stationary does not automatically indicate the rotational velocity, because synchronism also occurs when the flashing rate is some integral submultiple of the rotational speed. The practical procedure followed is therefore to adjust the flashing rate until synchronism is obtained at the largest flashing rate possible, R_1 . The flashing rate is then carefully decreased until synchronism is again achieved at the next lower flashing rate, R_2 . The rotational velocity is then given by:

$$v = \frac{R_1 R_2}{R_1 - R_2}$$

B. Find the Persian equivalents of the following terms and expressions and write them in the spaces provided.

1. beam splitter
2. calibration
3. cross-coupled
4. deflection type
5. discriminate
6. drag-cup
7. emerge
8. feature
9. fiber optic
10. gyroscope
11. interferometer
12. monitoring control instrument
13. null type
14. photoelectric pulse-counting method
15. pilot
16. potentiometer
17. rotational velocity
18. sleeve
19. spherical
20. spindle
21. squirrel-cage
22. stationary core
23. stroboscopic method
24. tachogenerator
25. tachometer
26. variable-reluctance transducer
27. zener diode

Unit 16

Section One: Reading Comprehension

Television Broadcasting

The term 'broadcasting' means to send out in all directions. As illustrated in Figure 16-1, the transmitting antenna radiates electromagnetic radio waves that can be picked up by the receiving antenna. For commercial television broadcast stations, the service area is about 25 to 75 mi in all directions from the transmitter. The radiation is in the form of two rf carrier waves, modulated by the desired information. Amplitude modulation(AM) is used for the picture signal. However, frequency modulation (FM) is used for the sound signal.

Referring to Figure 16-1, we see that the desired sound for the televised program is converted by the microphone to an audio signal, which is amplified for the sound-signal transmitter. For transmission of the picture, the camera tube converts the visual information into electrical signal variations. A camera tube is a cathode-ray tube with a photoelectric image plate.

The electrical variations from the camera tube become the *video signal*, which contains the desired picture information. The video signal is amplified and coupled to the picture-signal transmitter for broadcasting to receivers in the service area.

Separate carrier waves are used for the picture signal and sound signal, but they are radiated by one transmitting antenna. Furthermore, the picture and sound signals are included in the broadcast channel for each station. A television channel for a commercial broadcast station is made 6 MHz wide to include both the picture and sound. At the receiver also, one antenna is used for the picture and sound signals.

The receiving antenna intercepts the radiated picture and sound carrier signals, which are then amplified and detected in the receiver. The detector output includes the desired video signal containing the information needed to reproduce the picture. Then the recovered video signal is amplified and coupled to a picture tube that converts the electric signal back into light.

Reproducing the Picture

The picture tube is very similar to the cathode-ray tube used in the oscilloscope. The glass envelope contains an electron-gun structure that

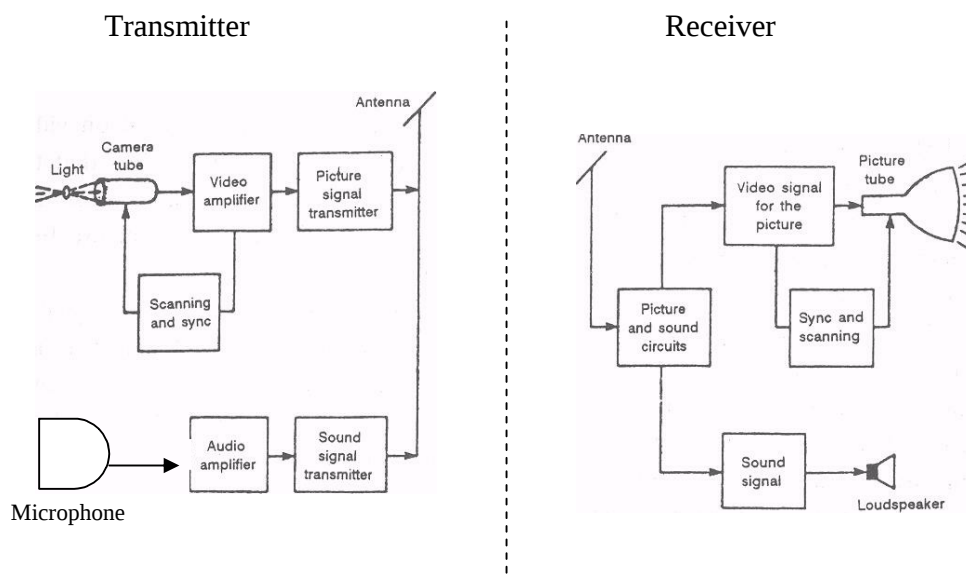


Figure 16-1. Block Diagram of the Television Broadcasting System.

produces a beam of electrons aimed at the fluorescent screen. When the electron beam strikes the screen, it emits light.

When the signal voltage makes the control grid less negative, the beam current is increased, making the spot of light on the screen brighter. More negative grid voltage reduces the brightness. If the grid voltage is negative enough to cut off the electron-beam current at the picture tube, there will be no light. This state corresponds to black. A color picture tube has three electron guns for the tricolor screen.

The picture tube is also called a *kinescope* or a *CRT*. Its function is to convert the video signal into a picture.

Scanning and Synchronizing

In order for the camera tube to convert the picture information into video signal, the image is dissected into a series of horizontal lines. Similarly, the picture tube reassembles the image line by line. These horizontal lines are produced by making the electron beam scan across the screen. There are 525 lines per picture frame. In addition to this horizontal scanning, vertical scanning is necessary to spread the lines from top to bottom of the screen. There are 30 complete picture frames per second.

Furthermore, the scanning at the camera tube and picture tube must be synchronized, or timed, with respect to the video signal. The synchronization

is necessary to reassemble the picture information on the correct lines. These functions are provided by the block of scanning and synchronizing circuits shown in Figure 16-1 for the transmitter and receiver. The term 'synchronizing' is usually abbreviated *sync*.

Most programs are produced live in the studio but recorded on video tape at a convenient time to be shown later. The quality is so good that the picture looks practically the same as a live program. The studio also has projectors to use 35-mm still slides, opaque slides, and motion-picture film, either 16 or 35 mm, as the program source.

For remote pickups, as in broadcasting a sports event, the signal is relayed to the studio for broadcasting in the assigned channel. When there is a national hookup for important programs, each station in the network receives video signal by intercity relay links, usually leased from the telephone company. A system for satellite relay stations covering the country is being developed for this nationwide television service.

Part I. Comprehension Exercises

A. Put "T" for true and "F" for false statements. Justify your answers

-1. Radio waves are used to carry the picture and the sound signals.
-2. Carrier waves are modulated by the desired information after being picked up by the receiving antenna.
-3. Sound signals are converted to audio signals prior to transmission.
-4. The radiated picture and sound carrier signals are amplified and detected in the receiver.
-5. The electron beam produced by the electron gun strikes the fluorescent screen causing the screen to emit light.
-6. The electron beam scans the screen horizontally to produce picture frames.

B. Choose a, b, c, or d which best completes each item.

- 1. As we understand from the first paragraph,
 - a. amplitude modulation may also be used for audio frequency signals
 - b. audio waves are modulated and then broadcast to the receiver
 - c. carrier waves are used to produce sound and picture signals
 - d. sound and picture signals cannot be transmitted directly
- 2. According to the text,
 - a. visual information is converted into electrical signal variations

- b. electrical variations become video signals containing the desired picture information
 - c. video signals are amplified and broadcast to the receiver
 - d. all of the above
3. Picture and sound signals are
- a. radiated by different transmitting antennas
 - b. received by two receiving antennas.
 - c. transmitted through the same television channel
 - d. broadcast by the same carrier wave
4. At the receiver,
- a. the detector produces the desired video signal
 - b. the video signal produced by the detector is amplified and coupled to a picture tube
 - c. the picture tube converts the video signal into a picture
 - d. all of the above
5. As we understand from the text,
- a. the picture tube used in television sets has a different mechanism from that used in an oscilloscope
 - b. the sound carrier signal is detected in a receiver different from that detecting the picture carrier signal
 - c. by varying the negative potential on the grid in the electron gun, the intensity of the beam varies
 - d. by increasing the negative potential on the grid the brightness of the spot of light on the screen increases

B. Answer the following questions orally

- 1. How are picture and sound signals transmitted to the receiving antenna?
- 2. What is the function of the camera tube?
- 3. What is the function of the picture tube?
- 4. How many picture frames are produced on the screen *per* second?
- 5. How does synchronization affect a picture frame?
- 6. How are remote pickups done?
- 7. What are intercity relay links?

Part II. Language Practice

A. Choose a, b, c, or d which best completes each item.

- 1. In amplitude modulation, of a carrier signal is varied by the

modulating voltage, whose frequency is invariably lower than that of the carrier.

- a. the frequency
- b. the amplitude
- c. the phase
- d. all of the above

2. A cathode-ray is an electron-beam tube in which the beam can be focused on to a small cross section on a luminescent screen and varied in position and intensity to produce a visible pattern.

- a. instrument
- b. oscillograph
- c. tube
- d. oscilloscope

3. An is an electrode structure that produces one or more electron beams.

- a. electron gun
- b. electron tube
- c. electronic controller
- d. electronic converter

4. A cathode-ray tube used to produce an image by variation of the beam intensity as the beam scans a raster is called

- a. an electron gun
- b. an oscillator
- c. a picture element
- d. a picture tube

5. A is an electron tube used to provide an image in color by the scanning of a raster and by varying the intensity of excitation of phosphors to produce light of the chosen primary colors.

- a. color-purity magnet
- b. color-picture tube
- c. color-selecting-electrode system
- d. color-field collector

B. Fill in the blanks with the appropriate form of the words given.

1. Amplify

- a. Diodes may be combined with electronic DC to form an electronic voltmeter or other electronic instruments.
- b. The horizontal amplifier in an oscilloscope the time-base voltage from the sweep oscillator, providing a control for the width of the resulting pattern.
- c. The operational amplifier serves as the heart of the analog computer, because it possesses the widely useful ability to provide a high value of precisely controlled

2. Transmit

- a. In an AM, amplitude modulation can be generated at any point after the radio frequency source.
- b. There are several different systems for TVand reception.

- c. The video chain at the station begins with a transducer which converts light into electric signals.

3. Radiate

- a. Any power escaping into free space is governed by the characteristics of free space. If such power escapes on purpose, it is said to have been
- b. Free space is the space that does not interfere with the normal and propagation of radio waves.
- c. Antennas electromagnetic waves, or, putting it differently, radiation will result from the flow of high-frequency current in a suitable circuit.

4. Detect

- a. The diode is used for AM demodulation or
- b. A detects the presence of electric waves.
- c. A radar the presence of objects and their distance by the transmission and return of electromagnetic energy.

5. Synchronize

- a. The task of the circuits in a television receiver is to process received information, in such a way as to ensure that the vertical and horizontal oscillators in the receiver work at the correct frequencies.
- b. In television, the synchronizing signal is employed for the of scanning.
- c. In a computer, each event, or the performance of each operation, starts as a result of a signal generated by a clock.

C. Fill in the blanks with the following words.

frequencies	between	useful	band
bandwidth	distance	long	low
developed	stations		

When practical radio transmission started in the year 1901, the radio frequencies of about 100 kHz were used for.....distances of hundreds to thousands of miles. As radio, higher frequencies were used for services requiring more..... . Now we have television broadcasting in the VHF of 30 to 300 MHz and the UHF band of 300 to 3,000 MHz. However, the..... for wireless transmission becomes much shorter at these highBroadcasting is practically limited to the line-of-sight distance

..... the transmitting and receiving antennas in the VHF and UHF bands. The service range is up to 75 mi for VHFand 25 to 35 mi for UHF stations.

D. Put the following sentences in the right order to form a paragraph. Write the corresponding letters in the boxes provided.

- a. Among the more important factors are the modulation system used, the operating frequency, the operating range, and the type of display required which in turn depends on the destination of the intelligence received .
- b. Receivers range widely in complexity, from a very simple crystal receiver with headphones, to a far more complex radar receiver with its involved antenna arrangements and visual display system.
- c. A great variety of receivers are used in communication systems, because the exact form of a suitable receiver for a particular usage is influenced by a number of different factors .
- d. Both these processes are the reverse of the corresponding transmitter processes .
- e. Whatever the receiver, its most important function is demodulation (and sometimes also decoding).

1	2	3	4	5
				

Section Two: Further Reading

The Cathode Ray Tube

The cathode ray tube operates as follows. First, electrons are emitted from a heated cathode. Then these electrons are accelerated to give them a high velocity. Next they are formed into a beam which can be deflected vertically and horizontally. Finally they are made to strike a screen coated on its inner surface with a phosphor.

The CRT comprises an electron gun and a deflection system enclosed in a glass tube with a phosphor coated screen. The electron gun forms the electrons into a beam. It contains a cathode which is heated to produce a stream of electrons. On the same axis as the cathode is a cylinder known as the grid. By varying the negative potential on the grid, the intensity of the beam can be varied. A system of three anodes follows. These accelerate the beam and also operate as a lens to focus the beam on the screen as a small dot. Varying the potential on the central anode, a_2 , allows the focus to be adjusted.

On leaving the electron gun, the beam passes through the deflection system. There are two systems in use today for moving the electron beam around on the face of the CRT. They are the **electrostatic** and the **electromagnetic** deflection systems. The type of deflection will depend on the ultimate use of the CRT display. Electrostatic deflection is accomplished by placing four metal plates inside the neck of the CRT with connecting wires to the outside for voltage application. Two of the plates control the vertical movement of the beam, and two plates control the horizontal movement.

The two vertical plates are placed above and below the path of the electron beam, as shown in Figure 16-2, while the two horizontal plates are on either side of the electron beam path. When one plate is grounded and a positive AC voltage is applied to the other plate, the path of the electron beam will bend toward the positive plate. The beam is moving too fast to strike the deflection plate, so only its direction is changed. The advantage of electrostatic deflection is that it performs equally well at all frequencies, from DC to the limits of the CRT's ability to produce a trace. The disadvantage is the small deflection angles of about $\pm 15^\circ$. A 21 in. picture tube would be about 48 in. long.

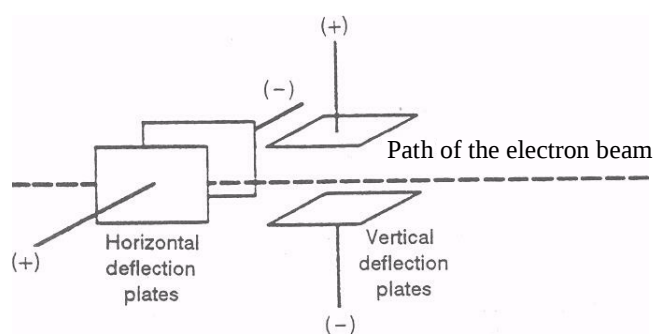


Figure 16-2. Electrostatic Deflection Plates.

Electromagnetic deflection uses two sets of coils in place of the deflection plates. A sawtooth current is required through the coils in order to achieve linear trace and retrace. The coils, called the **deflection yoke**, fit around the outside of the neck of the CRT and have an inductive component as well as the wire resistance considered to be in series with the inductance. A sawtooth voltage applied to the resistive element will produce a sawtooth current through the resistance. However, a square wave voltage must be applied to the inductive portion of the yoke to produce a sawtooth current through the inductance. The sum of the square wave voltage and the sawtooth voltage is called a **trapezoidal** voltage waveshape. The advantage of the electromagnetic deflection system is the wide deflection angles of about $\pm 60^\circ$. The disadvantage is that coils work best at only one frequency. However, this is not really a disadvantage because in television, each half of the yoke is required to work at only one frequency—60 Hz for the vertical and 15,750 Hz for the horizontal.

The vertical circuit consists of an oscillator that will operate near 60 cycles to generate the voltage wave needed for deflection. The vertical frequency adjustment (called the *vertical hold* control) is part of the oscillator circuit. This is a free-running oscillator that will be synchronized by the station signal at the exact frequency of that station. The vertical size control (or height control) adjusts the DC voltage to the oscillator and controls the strength of the output signal, which determines the size of the vertical scan of the CRT (see Figure 16-3).

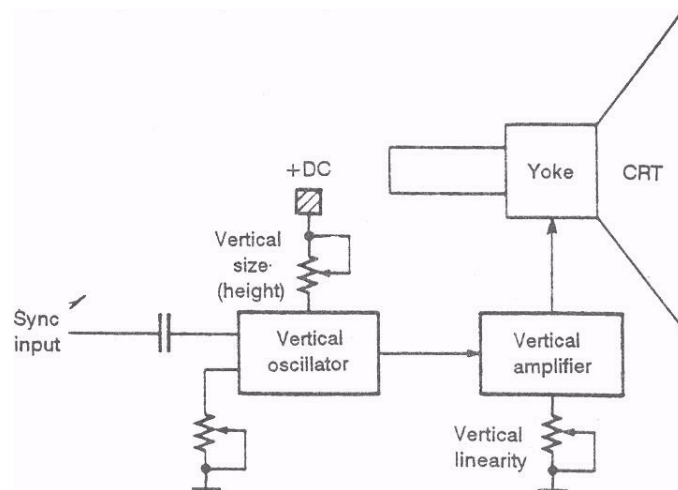


Figure 16-3. Vertical Deflection Block Diagram.

The oscillator signal is amplified by the vertical output power amplifier to convert the voltage signal to a current signal to drive the deflection yoke. The vertical **linearity** control is generally a part of this circuit. The linearity control is a gain control that varies the operating Q point of the amplifier and distorts the waveshape by compressing the top or bottom of the wave as required to achieve the best distribution of the vertical scan from top to bottom of the CRT screen. The size control and the linearity controls interact for best picture results.

Vertical yoke currents of 1 A are common. The vertical oscillator and the vertical amplifier with the deflection yoke comprise the entire vertical circuit.

The horizontal deflection section has the same functions as the vertical system, but is developed in a slightly different manner. The horizontal oscillator used for television was the first practical use of the phase-locked loop. The oscillator has a free-running frequency of 15,750 cycles adjustable by a control setting (*horizontal hold* control), as shown in Figure 16-4. That is, with no input signal, the circuit components are selected to control the frequency. The oscillator is a voltage-controlled circuit that depends on the DC voltage at the input to establish the frequency of operation.

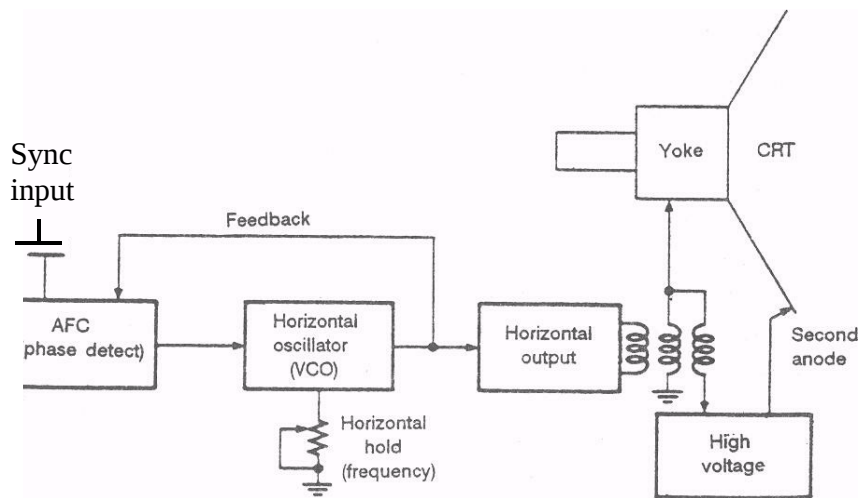


Figure 16-4. Horizontal Deflection and High-Voltage Block Diagram.

A feedback signal from the horizontal circuit is applied to a phase detector, which compares the frequency of an incoming pulse from the station

(when a station is selected) and develops the DC voltage to control the oscillator. The oscillator output is amplified by a power amplifier to develop the current signal used to drive the horizontal section of the deflection yoke. This exchange from voltage to current is accomplished through a horizontal output transformer. Because of the higher horizontal frequency (compared to the vertical), the horizontal yoke current need not be as great as the current in the vertical deflection system.

The final element is the phosphor coated screen. When the electron beam strikes the screen, the phosphor coating fluoresces. Various colors of light are produced depending on the phosphor used.

Comprehension Exercises

A. Choose a, b, c, or d which best completes each item.

1. The intensity of the beam is controlled by varying
 - a. the number of electrons emitted for the cathode
 - b. the potential on the central anode
 - c. the negative potential on the grid
 - d. the form of the electron beam
2. The second paragraph mainly discusses the mechanism of
 - a. the electron gun
 - b. the cathode ray tube
 - c. the deflection system
 - d. the central anode
3. The electrostatic deflection system consists of
 - a. four plates and a screen
 - b. two sets of deflection plates
 - c. an electron gun and a screen
 - d. four plates and an electron gun
4. It is true that
 - a. electrostatic deflection system uses two sets of coils
 - b. electromagnetic deflection system uses two sets of plates
 - c. electromagnetic deflection system works well at all frequencies
 - d. electrostatic deflection system performs well at all frequencies
5. According to the text,
 - a. the electromagnetic deflection system is made up of the deflection yoke which contains inductive and resistive elements
 - b. the sawtooth voltage applied to the resistive element produces a sawtooth current through the inductance
 - c. either a sawtooth voltage or a square wave voltage is required to produce a trapezoidal voltage waveshape
 - d. both a sawtooth current and a square wave voltage are needed to produce a trapezoidal waveshape

6. The voltage wave required for the vertical deflection is provided by
- | | |
|------------------|----------------|
| a. an oscillator | b. an inductor |
| c. a resistor | d. a converter |
7. The best picture results from the interaction between
- the vertical hold control and the oscillator
 - the oscillator and the vertical size control
 - the size control and the linearity controls
 - the amplifier and the oscillator

B. Write the answers to the following questions.

- What is the source of electrons for the electron beam?
- What is the function of the electron gun?
- In what way is the system of anodes like a lens?
- What are the two deflection systems called?
- How do the plates change the direction of the beam?
- What is a trapezoidal voltage waveshape?
- What does the vertical circuit consist of?
- How does the phase detector in the horizontal deflection system control the oscillator?



Section Three: Translation Activities

A. Translate the following passage into Persian.

Color Television

The block diagram in Figure 16-1 illustrates the television broadcasting system for monochrome. In color television, a color camera is necessary at the transmitter and a color picture tube at the receiver. The color camera provides video signal for the red, green, and blue picture information. A color picture tube has red, green, and blue phosphors on the viewing screen to reproduce the picture in color. A Typical color picture tube has three electron guns for the tricolor screen. The phosphors can be dot trios of red, green, and blue, or vertical stripes of color. Then each gun produces an electron beam to illuminate the red, green, or blue phosphor dots on the fluorescent screen.

Although the camera and picture tube operate with red, green, and blue all other colors including white can be reproduced by combinations of these three colors. Furthermore, in commercial television, the red, green, and blue signals are combined for broadcasting. The purpose is to transmit only a chrominance signal for color and a luminance signal that contains the monochrome information. It is necessary to transmit the luminance signal so that monochrome receivers can reproduce the picture in black and white. The chrominance signal or *chroma signal* has all the information needed to reproduce the picture in color.

The luminance signal is called the Y video signal. The chrominance signal can be called the C signal. Actually, the C signal is a modulated subcarrier of 3.58 MHz. This 3.58-MHz C signal modulates the assigned picture carrier in the standard 6-MHz television broadcast channel. Furthermore, the 3.58-MHz chrominance signal itself is modulated by two color video signals. The process of interweaving the Y signal for luminance with the 3.58-MHz color subcarrier signal for color is called *multiplexing*. In terms of the modulated chrominance signal, 3.58 MHz is the frequency for color in the television broadcast system.

B. Find the Persian equivalents of the following terms and expressions and write them in the spaces provided.

- | | |
|----------------------------|-------|
| 1.coil | |
| 2. deflection yoke | |
| 3. detect | |
| 4. distort | |
| 5. free-running oscillator | |
| 6. hookup | |
| 7. illuminate | |
| 8. intercept | |
| 9. kinescope | |
| 10. modulate | |
| 11.opaque | |
| 12. sawtooth current | |
| 13. subcarrier | |
| 14. trapezoidal | |
| 15. vertical hold control | |

Section One: Reading Comprehension

Transmission Lines

Transmission lines are a means of conveying signals or power from one point to another. From such a broad definition, any system of wires can be considered as forming one or more transmission lines. However, if the properties of these lines must be taken into account, the lines might as well be arranged in some simple, constant pattern. This will make the properties much easier to calculate, and it will also make them constant for any type of transmission line. Thus all practical transmission lines are arranged in some uniform pattern: this simplifies calculations, reduces costs, and increases convenience. There are two types of commonly used transmission lines. The parallel-wire (balanced) line is shown in Figure 17-lb, and the coaxial (unbalanced) line in Figure 17-la.

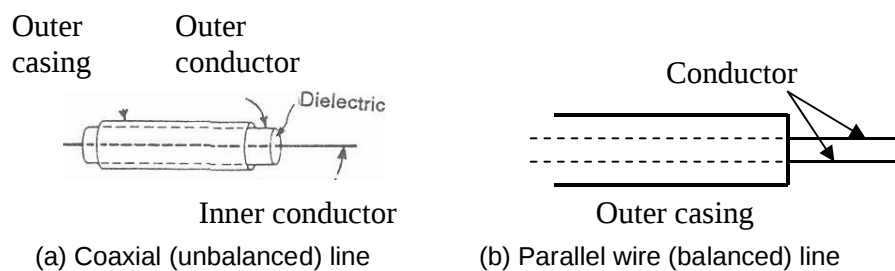


Figure 17-1. Transmission Lines.

The parallel-wire line is employed where balanced properties are required: for instance, in connecting a *folded-dipole* antenna to a TV receiver or a *rhombic* antenna to an HF transmitter. On the other hand, the coaxial line is used when unbalanced properties are needed, as in the interconnection of a broadcast transmitter to its grounded antenna. It is also employed at UHF and microwave frequencies, to avoid the risk of radiation from the transmission line itself.

Any system of conductors is likely to radiate if the conductor separation approaches a half-wavelength at the operating frequency. This is far more likely to occur in a parallel-wire line than in a coaxial line, whose outer conductor surrounds the inner one and is invariably grounded. Accordingly,

parallel-wire lines are never used for microwaves, whereas coaxial lines may be employed for frequencies up to at least 18 GHz. However, from the general point of view the limit is on the *lowest* usable frequency; below about 1 GHz, waveguide cross-sectional dimensions become inconveniently large. Within each broad grouping or type of transmission line there is an astonishing variety of different kinds, dictated by various applications. Lines may be rigid or flexible, air-spaced or filled with different dielectrics, with smooth or corrugated conductors as the circumstances warrant. Different diameters and properties are also available. Flexible lines are naturally more convenient than rigid ones, since they may be bent to follow any physical layout and are much easier to stow and transport. On the other hand, rigid cables can generally carry much higher powers, and it is easier to make them air-dielectric rather than filled with a solid dielectric. This consideration is important, especially for high powers, since all solid dielectrics have significantly higher losses than air, particularly as frequencies are increased.

Rigid coaxial air-dielectric lines consist of an inner and outer conductor with spacers of low-loss dielectric separating the two every few centimeters. There may be a sheath around the outer conductor to prevent corrosion, but this is not always the case. A flexible air-dielectric cable generally has corrugations in both the inner and the outer conductor, running at right angles to its length, and a spiral of dielectric material between the two.

The power-handling ability of a transmission line is limited by flashover between the conductors due to a high-voltage gradient breaking down the dielectric. It depends on the type of dielectric material used, as well as the distance between the conductors. Thus, for the high-power cables employed in transmitters, nitrogen under pressure may be used to fill the cable and reduce flashover. Since nitrogen is less reactive than the oxygen component of air, corrosion is reduced as well. Dry air under pressure is also used as a means of keeping out moisture. Clearly, as the power transmitted is increased, so must be the cross-sectional dimensions of the cable.

Since each conductor has a certain length and diameter, it must have resistance and inductance; since there are two wires close to each other, there must be capacitance between them. Finally, the wires are separated by a medium called the *dielectric*, which cannot be perfect in its insulation; the current leakage through it can be represented by a shunt conductance. The resulting equivalent circuit is as shown in Figure 17-2.

At radio frequencies, the inductive reactance is much larger than the resistance. The capacitive susceptance is also much larger than the shunt

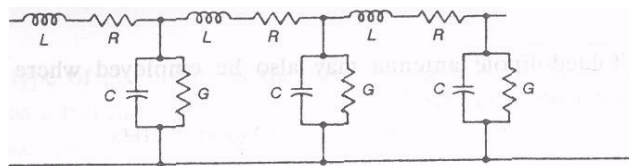


Figure 17-2. General Equivalent Circuit of Transmission Line.

conductance. Thus both R and G may be ignored, resulting in a line that is considered lossless (as a very good approximation for RF calculations). The equivalent circuit is simplified as shown in Figure 17-3.

It is to be noted that the quantities L , R , C , and G , shown in Figures 17-2 and 17-3, are all measured per unit length, e.g., per meter, because they occur continuously along the line.

They are thus distributed throughout the length of the line. Under no circumstances can they be assumed to be lumped at any one point.

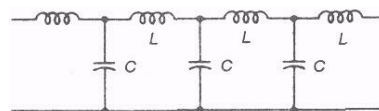


Figure 17-3. Transmission-Line RF Equivalent Circuit.

Part I. Comprehension Exercises

A. Put "T" for true and "F" for false statements. Justify your answers.

- 1. The parallel-wire line may be used to connect a broadcast transmitter to its grounded antenna.
- 2. The parallel-wire line is usually used at HF and UHF frequencies.
- 3. A parallel-wire line is more liable to radiation than a coaxial line.
- 4. The higher the frequency, the higher the power loss a solid dielectric will have.
- 5. The sheath around the outer conductor of a rigid coaxial air-dielectric cable is not of much use.

B. Choose a, b, c, or d which best completes each item.

1. The first paragraph mainly discusses
 - a. the basic principles of transmission lines
 - b. the basic calculations for transmission lines
 - c. how signals are conveyed from one point to another
 - d. how transmission lines are arranged
2. As we understand from the text,
 - a. a rhombic antenna is the most popular antenna used in TV systems

- b. a folded-dipole antenna may also be employed where unbalanced properties are needed
 - c. waveguides are not normally used below 1 GHz
 - d. coaxial lines are not normally used between 1 and 18 GHz
3. Paragraphs 2, 3, and 4 mainly describe
- a. the balanced and unbalanced transmission lines
 - b. the fundamentals of transmission lines
 - c. practical transmission lines for use in audio-frequency applications
 - d. practical transmission lines manufactured in different forms
4. It is true that
- a. flashover due to a high-voltage gradient has no effect on high-power cables
 - b. flashover may be reduced due to the high reactive property of nitrogen
 - c. a high-power cable of small cross-sectional dimension can withstand serious flashover
 - d. a high-power cable must be made so as not to give up under flashover conditions
5. As we understand from Figure 17-2,
- a. all the quantities shown cause equal problems throughout the length of the line
 - b. all the quantities shown are proportional to the length of the line
 - c. resistance along the line occurs between the two wires in the cable
 - d. shunt conductance along the line is due to high resistivity of wires in the cable

C. Answer the following questions orally.

1. What are the two types of transmission lines commonly used?
2. What is the use of parallel-wire line?
3. What are the advantages of rigid cables over the flexible one?
4. What does a rigid air-dielectric line consist of?
5. What is a spacer?
6. What comprises a flexible air-dielectric cable?
7. What causes the capacitance along the line?
8. How are the quantities L, R, C, and G, considered at radio frequencies?

Part II. Language Practice

A. Choose a, b, c, or d which best completes each item.

1. What is formed by two coaxial conductors is
 - a. a parallel-wire line b. a directional-power relay
 - c. a coaxial line d. a signal carrier

2. One type ofline is the two-wire open line which is sometimes used as a transmission line between antenna and transmitter or antenna and receiver.
 - a. rigid air-dielectric
 - b. flexible air-dielectric
 - c. parallel
 - d. coaxial
3. The electric and magnetic fields in the two-wire parallel line extend into space for relatively great distances, andlosses occur.
 - a. transmission
 - b. power
 - c. reflection
 - d. radiation
4. Any one of a class of antennas producing the radiation pattern approximating that of an elementary electric dipole is known as antenna.
 - a. rhombic
 - b. grounded
 - c. dipole
 - d. quarter-wave
5. The property of a system of conductors and dielectrics that permits the storage of electrically separated charges when potential differences exist between the conductors is referred to as..... .
 - a. resistance
 - b. capacitance
 - c. inductance
 - d. conductance

B. Fill in the blanks with the appropriate form of the words given.

1. Flexible

- a. Concentric cables may be made, with the inner conductor consisting of wire insulated from the outer conductor by a solid, continuous insulating material.
- b. Early attempts at obtaining employed the use of rubbed insulators between the two conductors.

2. Shield

- a. The pair consists of parallel conductors separated from each other and surrounded by a solid dielectric.
- b. The conductors are contained within a copper braid tubing that acts as a
- c. The fields are confined to the space between the two conductors; thus, the coaxial line is a perfectlyline.

3. Ground

- a. Theparts may be connected to ground without affecting operation of the device.

- b. No electric or magnetic fields extend outside of the..... conductor in a coaxial line.
- c. A ground bus is a bus to which the from individual pieces of equipment are connected, and that, in turn, is connected to ground at one or more points.
- d. A ground cable band is used for the armor or sheaths of cables or both.

4. Space

- a. Space charge is the electric charge in a region of, due to the presence of electrons and/or ions.
- b. The direct wave is basically limited to so-called line-of-sight transmission distances.
- c. Ashaft is a separate shaft connecting the shaft ends of two machines.
- d. Apulse or space is the signal pulse that, in direct-current neutral operation, corresponds to a circuit open, or no current condition.

5. Break

- a. The length of a multipleis the sum of two or more breaks.
- b. A motor develops the breakaway torque to..... away its load from rest to rotation.
- c. Breaking capacity is the current that the device is capable of at a stated recovery voltage under prescribed conditions of use and behavior.

C. Fill in the blanks with the following words.

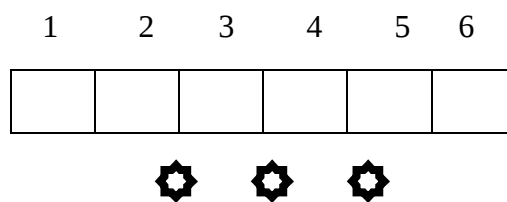
Characteristic	operating	line
frequency	below	be

The impedance of the cable, termed theor **surge impedance** Z_0 , is considered to.....,independent of the cable length and the..... Frequency. This consideration is valid when the..... is properly terminated and when the operating..... is above a few tens of kilohertz, but..... a few gigahertz

D. Put the following sentences in the right order to form paragraph. Write the corresponding letters in the boxes provided.

- a. A quantitative indication of the nature of a particular standing wave is given by the *standing-wave ratio* (SWR).

- b. When voltage and current waves are reflected on a line due to a discontinuity, standing waves are produced.
- c. It is of the nature of this pattern that there are points of maximum and minimum values.
- d. Standing waves are the result of the summing of instantaneous values of incident and reflected waves at every point along a line.
- e. The standing-wave ratio is defined as the ratio of the maximum value of a wave to its minimum value.
- f. The summing process produces a pattern of variation (the standing wave) along the line.



Section Two: Further Reading

Antennas

A source has no way of knowing whether a line is infinite or finite when it begins to supply current and voltage waves to the line. If the line is terminated (connected to a load) in a resistance whose value is equal to Z_0 , the voltage and current waves will 'enter' that resistance and be dissipated. The energy that the waves represent will be taken off the line by the terminating device (the resistance); none of the energy will be returned to the line.

On the other hand, if the line is simply an open line of finite length, something must happen to the waves when they reach the end of the line. Since there is nothing connected to the line to absorb them, they will be reflected back from the end of the line and will travel along the line toward the source. On the line there will now be voltage and current waves coming from the source, and voltage and current waves traveling back from the end of the line. The waves from the source are called *incident waves*, those reflected from the end are called *reflected waves*. As with any ac voltage or current, the two sets of waves will combine phasorally at each point along the line—the incident voltage wave with the reflected voltage wave, incident current wave

with reflected current wave. As a result of the waves combining, there will be established on the line, patterns of voltage and current variations. It happens that these patterns do not move or travel. These are, therefore, *standing waves* and are known by that name.

The standing waves of voltage and current of an open line are depicted in Figure 17-4. The points in a standing-wave pattern where a voltage or current is a maximum are called *loops*, the points where values are minimums are called *nodes*.

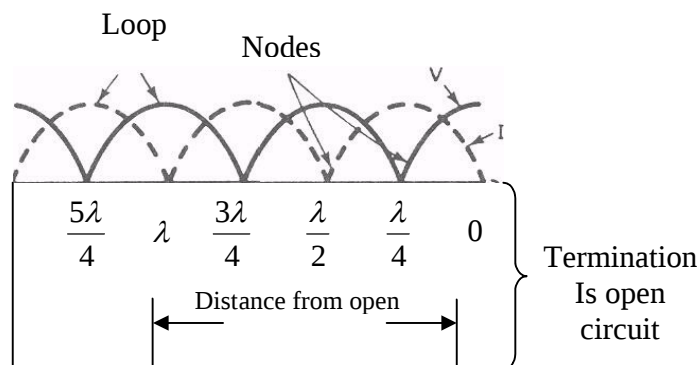


Figure 17-4. Standing Waves on an Open Line.

You will observe that on an open line, the voltage standing wave has a loop at the open end, and the current standing wave has a node at that end. This is as we would expect from the theory of ordinary circuits—the voltage is maximum across an open, the current is zero at the open.

Let us imagine that we have all the facilities for exciting a transmission line with RF energy and for carefully measuring transmitted and reflected power on the line. The line, a parallel-wire line, is open and is slightly longer than one wavelength at the exciting frequency. Excited, it develops standing waves because it is neither infinite in length nor terminated in a load that will remove all of the RF energy transmitted. If not removed, energy reaching the open end of the line will be reflected. The reflected energy, measured at the sending end, will be equal to the energy fed to the line from the source minus the energy lost during its trip from the generator, along the line to the open end, and back.

We expect some energy to be lost—the usual I^2R loss due to current flowing in the resistance of the conductors, and a much smaller loss in the dielectric between the conductors. We can predict these losses quite accurately, however, using facts about conductor size, measured current, and so on.

We are puzzled, therefore, when we discover that the reflected energy, as

measured at the sending end of the line, is significantly less than the transmitted energy minus the predicted losses. What can explain the greater-than-expected loss of energy? The answer is *radiation*! A small but significant amount of RF energy has simply left the transmission line and is traveling away from it. This phenomenon is the one that makes possible all wireless communication. An antenna is a device that enhances the process of radiation of RF energy from a system.

The radiation of electrical energy is very common. The science of measuring and predicting radiation is highly developed. However, because of its nature-it is silent, invisible, and odorless-it does not lend itself to an explanation in simple physical terms. The explanation that follows is a common one. It does not tell the entire story, but is useful in providing a working understanding of radiation independent of the very elegant but highly sophisticated mathematical treatments.

To continue with the example of the transmission line of the preceding paragraph, let us imagine the conductors of the last $\lambda/4$ of the line being spread apart slightly, as in Figure 17-5a. The standing waves of voltage and current produce an electric field between the conductors and a magnetic field around each conductor. These are represented by the lines of force of Figure 17-5b. Notice *the fringing* (bowing out) of the electric field lines at the end of the line. (Fringing of electric field lines is a common phenomenon at the boundaries of an electric field between two conductors.) Fringing occurs because field lines running in the same direction exert a repulsive force on each other. At a field boundary, this force produces the spreading out and bowing of the lines.

When the electric field lines are the product of an ac voltage, the lines must be produced, collapsed, and reproduced in a reverse direction at a rate equal to the frequency of the voltage. We can imagine the lines being sent out from, and withdrawn to, the conductors. It is useful to theorize that, when the frequency exceeds approximately 20,000 cycles (40,000 reversals) per second, the outermost line of the field simply cannot keep up with the process of reversal. Being unable to return to its conductor, it closes upon itself, forming a closed loop. This loop is repulsed by the outermost line of force produced by the next alternation of the ac voltage. That line of force subsequently fails to make it back to the conductor and forms another closed loop, etc. The process repeats itself during each cycle. The closed loops are driven farther and farther away from the conductors. The result is a continuous wave train of energy being discharged (radiated) and repulsed from the conductors.

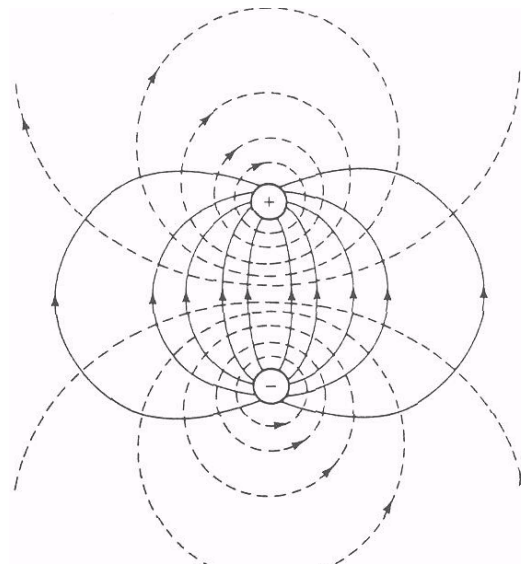
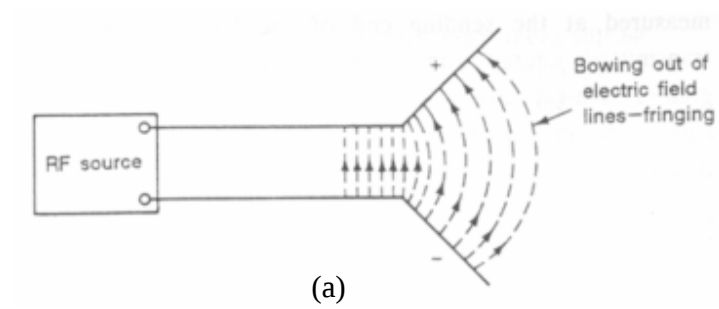


Figure 17-5. (a) Fringing of Electric Field Lines at the End of Transmission Line; (b) Cross-Sectional View of Twin-Lead Transmission Line With Electric and Magnetic Field Lines.

A Basic Antenna: The Half-Wave Dipole

Let us return to the picture of energy being radiated from the slightly spread end of a parallel-wire transmission line. Refer to Figure 17-5 again. What might we do to maximize radiation, since that is what is desired in an antenna? The answer is to bend further the final quarter-wavelength of each conductor until each is at right angles to the line (see Figure 17-6). Since

the length of each conductor bent is $\lambda/4$, the overall length of the perpendicular portion of the line is now $\lambda/2$. The device thus produced is called a *half-wave dipole*. The half-wave dipole is a simple, basic antenna. It is commonly used as a basis for comparison for more complex antennas.

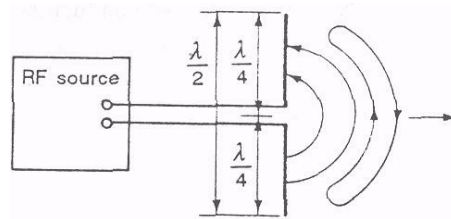


Figure 17-6. Half-Wave Dipole Antenna Made by Spreading Apart Conductors of $\lambda/4$ of Transmission

Line

Information about standing waves on transmission lines can be transferred to the half-wave dipole antenna. It is like an open line. The distance from the tip of each pole to the center feed point is $\lambda/4$. Hence there will be voltage standing-wave loops at the tips of the dipole (like the loop at the end of an open line). There will be a voltage wave node $\lambda/4$ away, at the center feed point. Similarly, there will be current standing-wave nodes at the tips, and a current loop $\lambda/4$ away, at the center feed point (see Figure 17-7). The standing-wave pattern just described is exactly the one needed to maximize radiation.

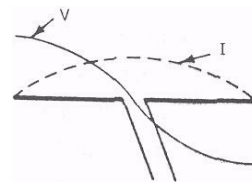


Figure 17-7. Voltage and Current Standing Waves on Half-Wave Dipole Antenna.

Comprehension Exercises

A. Choose a, b, c, or d which best completes each item.

- The second paragraph mainly describes along an open line of finite length.
 - the patterns of voltage and current variations
 - the voltage and current waves traveling
 - incident waves, reflected waves, and standing waves
 - incident waves variations at each point
- The mechanisms of an antenna is based on the process of..... .
 - radiation
 - reflection
 - power dissipation
 - line excitation
- Paragraph 5, 6, and 7 mainly discuss
 - transmission lines excited by RF energy
 - parallel-wire lines longer than one wavelength

- c. basic properties of conductors
- d. basic principles of antennas

4. It is true that

- a. the fringing of electric field lines formed at the end of the line is compatible with the fringing of magnetic field lines around each conductor
- b. the electric field lines produced by an ac voltage lead to a continuous wave train of energy radiated and repulsed from the conductors
- c. the bowing of electric field lines may occur at any point along the line
- d. the bowing of electric lines occurs because of the field lines running in opposite directions

5. We may conclude from the text that

- a. antennas are employed for only the generation of electromagnetic energy
- b. antennas receive electromagnetic waves and convert them into RF currents
- c. an antenna is a passive device; that is, it cannot add any energy to a signal that has been fed to it for processing
- d. an antenna is identical to a circuit containing a transistor; that is, it adds energy to the signal it is processing

6. It is true that

- a. if we have maximum radiation from an antenna, all energy applied to it will be converted to electromagnetic energy and radiated
- b. If an antenna is made up of a parallel-wire line with two quarter-wave sections, the electromagnetic energy radiated will be maximized
- c. the distance between the conductors of a parallel-wire line causes the formation of loops at the tips of the dipole
- d. the distance between the conductors of a parallel-wire line causes the radiated energy to be maximized

B. Write the answers to the following questions.

1. How do you explain the loops and nodes on an open line?
2. What are the values of voltage and current at the end of an open line?
3. What causes energy losses of a transmission line?
4. What kind of energy travels away from an open transmission line?
5. What is wireless communication based on?

6. What are the characteristics of an open-ended line?



Section Three: Translation Activities

A. Translate the following passage into Persian.

Traveling Voltage and Current Waves

Let us imagine that a parallel-wire transmission line is connected through a switch to a source of RF voltage, as in Figure 17.8a. Let us imagine, further, that the line is infinite in length. At the moment the switch is closed the effect of an electrical disturbance begins to be felt on the line. The effect is that of a

Radio-frequency sinusoidal voltage. This effect travels down (or along) the line at approximately the speed of light. (The speed, may be somewhat less than the speed of light depending on the exact nature of the line.) At any given instant, because of the sinusoidal nature of the voltage, at some points along the line the voltage will be zero volts, at other points it will be maximum positive volts. At still other points the voltage will be equal to the peak negative amplitude, or it will be equal to anything in between these values. In other words, at any given instant there is a pattern of sinusoidal voltage variation along the line. And this pattern is traveling away from the source. We say that a voltage wave is traveling down the line from the source. The idea of a voltage wave is shown in Figure 17-8b.

Since the line is imagined to be one of infinite length, its input impedance will be equal to its Z_o . The source will supply a current to the line with a value given by $I = V_s / Z_o$. Because Z_o is resistive in its nature, this current will be sinusoidal and in phase with the source voltage. The current effect will travel down the line just at the voltage effect did. We say that there is a current wave traveling down the line. The current wave is depicted in Figure 17-8c.

In summary, when a line of infinite length is connected to a source of ac voltage there is produced on the line a traveling voltage wave and a traveling current wave. These waves travel away from the source toward the opposite end of the line. The current wave is in phase with the voltage wave at every point along the line; its amplitude is determined by $I = V_s / Z_o$.

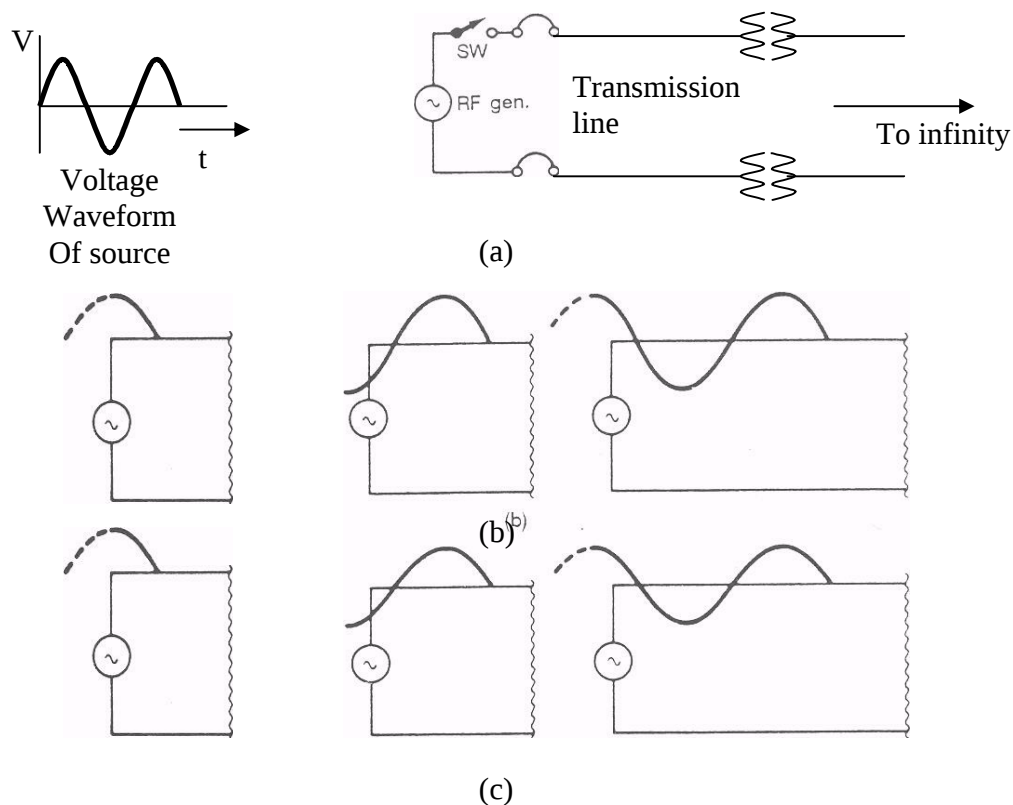


Figure 17-8. (a) Transmission Line and RF Source; (b) Traveling Voltage Wave on Transmission Line; (c) Traveling Current Wave on Transmission Line

B. Find the Persian equivalents of the following terms and expressions and write them in the spaces provided.

1. circumstance
2. concentric cable
3. conductance
4. convenience
5. corrugate
6. dielectric
7. dimension
8. flashover
9. flexible air antenna
10. folded dipole antenna
11. fringing

12. gradient.	
13. grounded antenna	
14. half-wave dipole	
15. incident wave	
16. leakage	
17. node	
18. odorless	
19. open-ended line	
20. parallel-wire line	
21. quarter-wave antenna	
22. reflected wave	
23. rhombic antenna	
24. rigid cables	
25. sheath	
26. standard-wave ratio (SWR)	
27. standing wave	
28. susceptance	

Unit 18

Section One: Reading Comprehension

Waveguides

Any system of wires may be used as a transmission line, but the simplest arrangements are invariably preferred in practice. Thus parallel-wire and coaxial lines are by far the most common. In a similar way, a pipe with *any* sort of cross section could be used as a waveguide, but the simplest cross sections are preferred. Accordingly, waveguides with constant rectangular or circular cross sections are normally employed, although other shapes may be used from time to time for special purposes. As with regular transmission lines, so in waveguides, the simplest shapes are the ones easiest to manufacture, and the ones whose properties are simplest to evaluate. A rectangular waveguide is shown in Figure 18-1, as is a circular waveguide for comparison. In a typical setup, there may be an antenna at one end of a waveguide and some form of load at the other end. The antenna generates electromagnetic waves, which travel down the waveguide to be eventually received by the load. It is seen that the waves are truly *guided*.

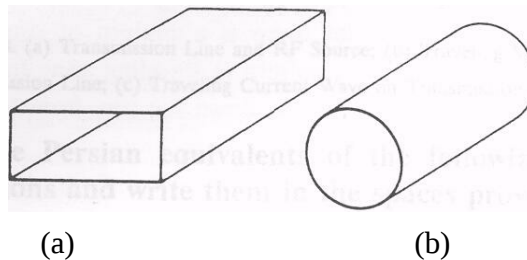


Figure 18-1. Waveguides, (a) Rectangular; (b) Circular.

The walls of the guide are conductors, and therefore reflections from them take place. It is of the utmost importance to realize that *conduction of energy takes place not through the walls*, whose function is only to confine this energy, *but through the dielectric filling the waveguide*, which is usually air. *In discussing the behavior and properties of waveguides, it is necessary to speak of electric and magnetic fields, as in wave propagation, instead of voltages and currents, as in transmission lines.* This is the only possible approach, but it does make the behavior of waveguides more complex to grasp.

Because the cross-sectional dimensions of a waveguide must be of the same order as those of a wavelength, use at frequencies below about 1 GHz is not normally considered, unless special circumstances warrant it.

Both waveguides and transmission lines can pass several signals simultaneously, but in waveguides it is sufficient for them to be propagated in different *modes* to be separated. They do not have to be of different frequencies. Again, a number of waveguide components are similar if not identical to their coaxial counterparts. These components include stubs, quarter-wave transformers, directional couplers and taper sections. Indeed, the operation of a very large number of waveguide components may best be understood by first looking at their transmission-line equivalents.

A major problem with twin-lead transmission lines at higher frequencies is that the amount of direct radiation from such lines increases with the frequency of the signal being transmitted. The result is that a twin-lead transmission line radiates virtually all of the energy it is carrying, and transmits little, if any, to a load at frequencies above several hundred megahertz. The problem of energy loss due to radiation is almost totally eliminated with coaxial lines and waveguides because these forms of transmission lines ‘enclose’ the signal and prevent its radiation.

A second source of energy loss for both parallel-lead and coaxial transmission lines is in the dielectric that supports the separation of the conductors. This is called *dielectric loss*. Although, theoretically, there is no current flow in an insulator, a dielectric, there is some current flow in actual, practical dielectrics, and there is dissipation. Of course, this dissipation is extremely small. However, it increases with frequency. Again, at very high frequencies, an energy loss becomes consequential. Because waveguides are completely hollow and in most cases filled with air, dielectric loss is virtually nil.

A third form of energy loss in transmission lines is in the I^2R heating of the conductors of the line. Heating or ‘copper’ loss is directly proportional to the resistance of a conductor, for a given current. And the resistance of conductors of RF energy increases with frequency! This is the result of the phenomenon called ‘skin effect’. As the frequency of a current increases, it ‘travels’ more and more on the surface of a conductor. The penetration of the disturbance of electron movement becomes shallower. This means that a smaller cross section of a conductor is utilized for current flow. And the consequence of that, in turn, is an increase in the resistance of the conductor since resistance is inversely proportional to cross-sectional area. Coaxial lines

represent some improvement over parallel-wire lines in the matter of heating loss since the conduction area of the outer conductor is significantly larger than that of the inner conductor. However, a waveguide has a major advantage in this regard: the inner (or one) conductor is completely eliminated; and the conduction area of the inner surface of the guide is significantly larger than that of the coaxial line.

The use of waveguides is not all gravy. In comparison with other forms of transmission lines, they are difficult and expensive to install. The skills required for installation are more like those of a plumber than of an electronics technician, or even of an electrician. Waveguides, in most instances, are rigid devices. Their routing must be carefully planned. Joints or connection points must be carefully made to avoid discontinuities in the inner, reflecting surfaces and the consequent creation of standing waves. Other forms of transmission lines are relatively flexible and can simply be unrolled and positioned to conform with almost any surface contour.

Waveguides are more expensive to manufacture. They must be precision made. Inner surfaces must conform to precise dimensions and be free of burrs, unevenness, etc., which could disturb the reflection patterns of the guided waves. Several waveguide sections of various shapes used to accommodate a variety of routing situations are shown in Figure 18-2.

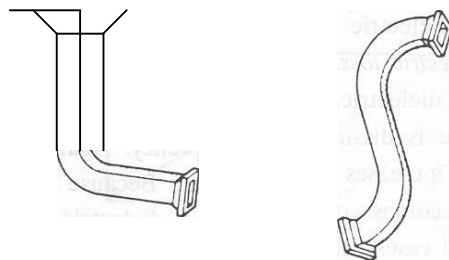


Figure 18-2. Miscellaneous Waveguide Sections.

Part I. Comprehension Exercises

A. Put “T” for true and “F” for false statements. Justify your answers.

- 1. Compared with other shapes, rectangular waveguides are the most common.
- 2. Propagation in rectangular waveguides is too difficult to evaluate compared with other shapes.
- 3. In a waveguide, conduction of energy takes place through the air filling it.

- 4. In order to evaluate the behavior of waveguides, their electric and magnetic properties must be considered.
- 5. Energy loss due to radiation is increased with coaxial lines.

B. Choose a, b, c, or d which best completes each item.

1. Electromagnetic signals propagated through waveguides
 - a. may be of different modes
 - b. must be of different frequencies
 - c. may not have frequencies below 1 GHz
 - d. cannot be evaluated properly
2. it is true that
 - a. waveguides stop radiation loss
 - b. waveguides have the problem of radiation loss
 - c. twin-lead and coaxial lines decrease energy loss
 - d. twin-lead line transmits the energy it carries
3. Waveguides are preferred to twin-lead lines because their loss is almost nil.

a. radiation	b. dielectric
c. heating	d. all of the above
4. According to the text, dielectric loss in coaxial transmission lines

a. decreases with frequency	b. increases with frequency
c. is practically high	d. is virtually nil
5. As the frequency of a current along a conductor increases,
 - a. the resistance of the conductor decreases since resistance is directly proportional to frequency
 - b. the penetration of the disturbance of electron movement becomes shallower
 - c. the efficiency of the conductor increases, too
 - d. the conduction area of the conductor increases, too

C. Answer the following questions orally.

1. Why are waveguides not normally used at frequencies below 1 GHz?
2. What similarities are there between transmission lines and waveguides?
3. What are the disadvantages of waveguides over other transmission lines?
4. What are the skills required for waveguide installation similar to?

5. What difficulty might arise from rigid-waveguide installation?
6. Why are waveguides expensive to manufacture?

Part II. Language Practice

A. Choose a, b, c, or d which best completes each item.

1. A state of a vibrating system to which corresponds one of the possible propagation constants is known as
 a. mode voltage b. mode purity
 c. mode d. frequency
2. At high frequencies, the $I^2 R$ loss is mainly due to
 a. the skin effect b. the skin depth
 c. the self-inductance d. the self-impedance
3. For any component of a field, the ratio of the instantaneous value of at one point to that of any other point does not vary with time.
 a. a waveform b. a wave envelope
 c. a wavefront d. a standing wave
4. Transmission lines used for transmitting microwaves often take the form of completely hollow cylindrical or rectangular tubes called
 a. parallel lines b. coaxial lines
 c. waveguides d. waveforms
5. The time rate at which electric energy is transformed into heat in a dielectric when it is subjected to a changing electric field is referred to as
 a. dielectric loss b. dielectric loss angle
 c. dielectric loss factor d. dielectric loss index

B. Fill in the blanks with the appropriate form of the words given.

1. Reflect

- a. When a mismatch occurs, there is an interaction between the incident and waves.
- b. When a line is terminated with a short circuit, open circuit, or purely reactive load, no energy can be absorbed by the load so that total takes place.
- c. The voltage coefficient is defined as the ratio of the complex electric field strength of the reflected wave to that of the incident wave.

2. Guide

- a. The energy of a wave is concentrated within or near boundaries between materials of different properties.
- b. The wavelength in a waveguide, measured in the longitudinal direction is known as wavelength.

3. Dissipate

- a. Theoretically, there is no current flow in insulators, however, some current always flows in practical dielectrics and there is
- b. If a line is terminated in a resistance, the voltage and current waves will enter that resistance and they will be

4. Contour

- a. Flexible transmission lines can be unrolled and positioned to conform with any surface.....
- b. In a control system, the controlled path can result from the coordinated simultaneous motion of two or more axes.

5. Enclose

- a. Waveguides the signal and prevent its radiation.
- b. An relay has both coil and contacts protected from the surrounding area.
- c. A protective housing used to contain equipment and prevent personnel from accidentally contacting live parts is called an

C. Fill in the blanks with the following words.

reflection	wave	waveguide	equal
called	like	dissipate	frequency
simply	from	severely	wavelength

When the width, w , of a waveguide is exactly to half of the free-space wavelength of λ (i.e., when $A=2w$), the wave will bounceside to side in the waveguide and will make no progress down the guide. The angle of incidence or is zero. The frequency that makes this condition true is the *cutoff frequency* for the guide. The associated..... is called the *cutoff wavelength*. Furthermore, the waveguide will..... attenuate all frequencies lower than the cutoff (or all wavelengths longer than the cutoff wavelength). A/An..... is like a high-pass filter. Its attenuation is that of a resonant wave trap: it does not the energy in the heating of a conductor, it blocks the passage of the energy.

D. Put the following sentences in the right order to form a paragraph. Write the corresponding letters in the boxes provided.

- a. Waveguides must be excited by a generator in such a way that waves capable of being propagated will be produced at the point of excitation.
- b. Waveguides can be excited by connecting a generator to either a probe or loop inserted in the guide, or through a window usually called an iris or Slot.
- c. Unlike methods of exciting twin-lead and coaxial lines, it is not enough simply to connect a generator to two points on guide.
- d. As we have seen, waveguides are literally 'guides for electromagnetic waves'.

1	2	3	4
			

Section Two: Further Reading

Theory of Operation

In many respects, waveguides can be dealt with in ways not unlike those used for other transmission lines. Matters of characteristic impedance, the need for impedance matching, etc., are not significantly different for waveguides. However, learning a few new ideas is required if one wishes to gain an understanding of how a waveguide transmits energy. This understanding is useful as a basis for a working knowledge of some of the operating peculiarities and limitations of waveguides. Conventional explanations of waveguide theory utilize the concepts of electric and magnetic fields extensively.

It is useful to think of a waveguide as a special environment for the propagation of electromagnetic (EM) waves. The ideas involved are not significantly different from those we examined in connection with the propagation of such waves from antennas. In fact, waveguides are energized or excited by a probe which acts very much like an antenna. Energy is removed from a waveguide by an antenna-like probe (see Figure 18-3).

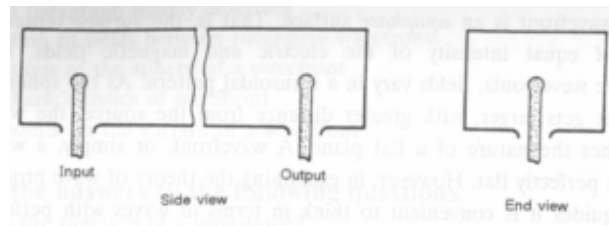


Figure 18-3. Input-Output Coupling for Waveguide Operation.

For electromagnetic waves, the ‘special environment’ which is a waveguide is like a tunnel. However, the waves are not able to simply streak straight down the tunnel like a train going through a tunnel. Rather, the waves are guided in their motion through the tunnel by bouncing from side to side of the tunnel. Indeed, although the propagation velocity of the waves along their zigzag journey is the same as that in free space (the speed of light), the velocity along the axis of the waveguide is less than the speed of light. This, of course, is the result of the actual distance traveled being greater than the length of the tunnel.

You will recall that electromagnetic waves consist of two inseparable components: an electric field component designated E , and a magnetic field component designated H . These components are vector quantities since they have both magnitude and direction in space. The vectors are always at right angles to each other. Together, the two vectors define a plane. The direction of travel (of propagation) of an EM wave is always perpendicular to the plane of the vectors.

When radiation is emitted from a point source—a small, simple antenna—the electromagnetic energy travels (is propagated) away from that source in all directions. Since the source of energy varies in amplitude at a radio-frequency rate, the intensity of the energy being propagated varies at the same rate. The result is that in the space surrounding the antenna, the energy intensity varies in a wave-like pattern. The pattern travels out from the source. The leading edge of this energy disturbance is called, appropriately, a *wavefront*. In the first instant of emission, the wavefront is like a small sphere surrounding the source. With time the wavefront travels away from the source, expanding the sphere. The action is like that of a spherical balloon being inflated—the wavefront corresponds to the surface of the balloon. With each succeeding alternation of the source, a new wave front is generated, and so on.

A wavefront is an *equiphase* surface. That is, the surface represents all points of equal intensity of the electric and magnetic fields. Between successive wavefronts, fields vary in a sinusoidal pattern. As the sphere of the wavefront gets larger, with greater distance from the source, the wavefront approaches the nature of a flat plane. A wavefront, or simply, a wave, can never be perfectly flat. However, in examining the theory of wave propagation in waveguides it is convenient to think in terms of waves with perfectly flat fronts. Such a wave is called a *uniform plane wave*.

Comprehension Exercises

A. Choose a, b, c, or d which best completes each item.

1. As we understand from the text, waveguides
 - a. are based on operational theories completely different from those of other transmission lines
 - b. are energized by a probe whose mechanism is similar to an antenna
 - c. cannot propagate electromagnetic waves as efficiently as other transmission lines
 - d. cannot be excited by antenna-like probes
2. It is true that
 - a. waves travel along a waveguide in a straight line
 - b. waves bounce back and forth as they move along a waveguide
 - c. the actual distance traveled by waves along a waveguide is greater than the length of the waveguide
 - d. the actual distance traveled by waves along a waveguide is equal to the length of the waveguide
3. We can infer from the text that
 - a. electromagnetic waves can be sent straight down a waveguide
 - b. electromagnetic waves are not influenced by the waveguide walls
 - c. the velocity of propagation in a waveguide can be the same as the speed of light
 - d. the electric field, the magnetic field, and the direction of propagation of an EM wave are mutually perpendicular
4. Coordination of variations in amplitude and intensity of the energy propagated result in
 - a. the energy traveling out from the source in a wave-like pattern
 - b. the energy traveling away from the source in all directions
 - c. the uniformity of the waves in free space
 - d. the concentration of the waves in a direct line

5. The last paragraph mainly describes
- a. variation of fields between successive wavefronts
 - b. variations of the sphere of a wavefront
 - c. the characteristics of wavefront
 - d. the points on the surface of a wavefront

B. Write the answers to the following questions.

1. What is the function of a waveguide?
2. Why is the propagation velocity of the waves along the axis of the waveguide less than the speed of light?
3. Why are the electric and magnetic field components of electromagnetic waves vector quantities?
4. What is a wavefront?
5. How does the wavefront vary traveling away from the source?
6. What is a wavefront compared with?
7. How are new wavefronts generated?



Section Three: Translation Activities

A. Translate the following passage into Persian.

Reflection of Waves From a Conducting Plane

As already discussed, an electromagnetic plane wave in space is transverse-electromagnetic, or TEM; the electric field, the magnetic field and the direction of propagation are mutually perpendicular. If such a wave were sent straight down a waveguide, it would not, despite appearances, propagate in it. This is because the electric field (no matter what its direction) would be short-circuited by the walls, since the walls are assumed to be perfect conductors, and thus a potential cannot exist across them. What must be found is some method of propagation which does not require an electric field to exist near a wall and simultaneously be parallel to it. This is achieved by sending the wave down the waveguide in a zigzag fashion, bouncing it off the walls and setting up a field that is maximum at or near the center of the guide, and zero at the walls. In this case the walls have nothing to short-circuit, and therefore they do not interfere with the wave pattern set up between them; thus propagation is not hindered.

Two major consequences of the zigzag propagation are apparent. The first is that the velocity of propagation in a waveguide must be less than in free space, and the second is that waves can no longer be TEM. The second situation arises because propagation by reflection requires not only a normal component but also a component in the direction of propagation for either the electric or the magnetic field, depending on the way in which waves are set up in the waveguide. This extra component in the direction of propagation means that waves are no longer transverse-electromagnetic, because there is now either an electric or a magnetic additional component in the direction of propagation.

B. Find the Persian equivalents of the following terms and expressions and write them in the spaces provided.

1. accommodate
2. boundary
3. circular
4. coaxial line
5. coefficient
6. cutoff frequency
7. cutoff wavelength
8. dielectric loss
9. dissipation
10. equiphase
11. iris
12. mismatch
13. modepurity
14. penetration
15. perpendicular
16. probe
17. rectangular
18. rigid waveguide
19. skineffect
20. slot
21. sphere
22. transmission
23. twin-lead line
24. uniform plane wave
25. wave front
26. waveguide

Unit 19

Section One: Reading Comprehension

Optical Communication Systems

An optical communications system is primarily a 'conventional' telecommunication system that utilizes light waves as a carrier in one or more transmission links. Although lightwaves can carry analog signals, optical systems are invariably also digital communications systems.

An optical communications system, then, is one in which the transmission link is an optical transmission line instead of a terrestrial metallic conductor transmission line, or microwave link or satellite link, etc. Such a system has terminal facilities incorporating typical digital communications functions: analog-to-digital and digital-to-analog converters, multiplexers and demultiplexers, carrier generators (light sources, in this case) and modulators, receivers and demodulators, and so on. What is really new and different to be learned about an optical communications system is the operation of the optical transmission line.

A special tube made of glass or plastic for the purpose of guiding light is called an *optical fiber*. An optical fiber is a waveguide for light waves. The term 'fiber' is appropriate because this tube or guide is a slender, thread-like structure. 'Optical' means that it has to do with light. An optical fiber is able to guide or conduct light along a path that is not a straight line. It can accomplish this feat with only a minimum of attenuation of the light. Fibers with attenuation characteristics of the order of 0.2 dB/km (decibels per kilometer) have been demonstrated in the laboratory. By comparison, the attenuation of 19-gauge twisted-wire-pair transmission line (in a multi pair cable) at voice frequencies is about 0.6 dB/km. Because attenuation on twisted-wire-pair line increases rapidly with frequency, operation is limited to approximately 1 MHz.

It is not difficult to conceive of a perfectly straight tube functioning as a light guide; light could simply shoot down the tube. Optical fibers, however, are seldom perfectly straight. By what principle are they able to conduct light around curves? The answer is *refraction*.

Refraction means bending of a wave-like entity such as light. The direction of a light wave—a light ray—is bent when the ray passes between two

media in which the velocities of propagation (of light) are different. You have experienced this phenomenon if you have ever been puzzled when trying to locate something under water while looking at it from above the surface of the water. You will recall that the object (e.g., a bar of soap in a bathtub) was not where you 'saw' it to be. The light rays from the object were refracted as they left the surface of the water. Light travels faster in air than in water.

Optical fibers guide light by refraction. In simplest form, a light 'conductor' could consist of a solid glass or plastic rod surrounded by air (see Figure 19-1). Since the solid and air have different propagation characteristics, light rays in the rod would be refracted from the sides of the rod and thus be guided through it.

For various reasons, optical fibers for communications applications have a form somewhat more complex than the simple 'light pipe' of Figure 19-1. As shown in Figure 19-2, a typical optical fiber consists of three basic elements: a central *core* (a solid rod of glass or plastic) surrounded by a protective coating of a different material called a *cladding*, which, in turn, is covered with a protective *sheath*. Before proceeding further with the specifics of optical fibers, let us examine briefly the basic 'rules' of refraction and learn some terminology commonly used in discussions involving optical communication.



Figure 19-1. Concept of Light Travel

Through a Light Pipe.

Figure 19-2. Three Basic Parts of

an Optical Fiber.

Snell's Law

The performance of light rays in refraction at the boundary of two light-conducting media is predictable from a principle known as Snell's law. Before looking at Snell's law, however, let us learn the meaning of basic terms associated with refraction. Refer to Figure 19-3. Observe that Figure 19-3 depicts the boundary between two media. Each medium is characterized by a property called its *refraction index*, n . The media in the diagram have indexes of n_1 and n_2 ; n_1 is greater than n_2 . The index of refraction of a material is

inversely proportional to the velocity of propagation of light in the material. For example, air has an index of refraction of approximately 1 (1.0002914); a typical n for glass is 1.5: light propagates more slowly through glass than through air. In Figure 19-3, light travels faster in medium 2 than in medium 1.

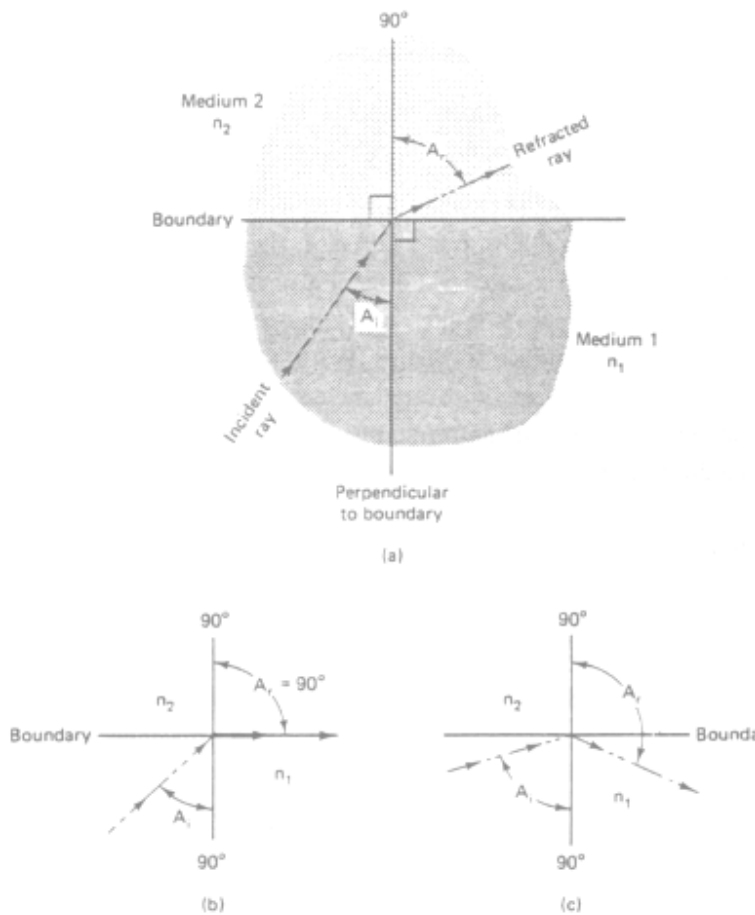


Figure 19-3. Refraction and Reflection: (a) Angles of Incidence and Refraction at Boundary Between Media of Different Indexes of Refraction; (b) $A_i = \text{Critical Angle}$; (c) $A_i > \text{Critical Angle}$.

Examine Figure 19-3 further. Observe that three conditions of a light ray encountering a media boundary are shown. The conditions relate to different *angles of incidence*. The angle of incidence, A_i , of an arriving (incident) light ray is the angle the ray makes with a line perpendicular to the media

boundary at the point where the ray meets the boundary. The *angle of refraction*, A_r , is the angle between the perpendicular to the boundary and the ray as it continues on its way. There is a consistent relationship between A_i and A_r . Snell's law:

The ratio of the sine of the angle of incidence in medium 1 (of two media) to the sine of the angle of refraction in medium 2 is a constant, K , and equal to the ratio of the index of refraction n_2 of the second medium to that, n_1 , of the first:

$$\frac{\sin A_i}{\sin A_r} = \frac{n_2}{n_1} = k$$

Note from Figure 19-3 that when A_i is relatively small, as in Figure 19-3 (a), the ray is bent but is able to exit medium 1; that is, it is able to cross the boundary and continue in the general direction of its original path. In Figure 19-3(c), however, where A_i is quite large, the ray is reflected back into medium 1; it is not able to escape. Figure 19-3(b) illustrates what is called the *critical angle*. The critical angle is the incidence angle which produces a refraction angle of 90° . When the refraction angle is 90° , the ray neither exits the first medium nor is reflected back into it. Its direction is along the boundary.

When optical fibers are used as transmission lines for light in a communications system, it is important that they be operated so that most of the light that is introduced to the fiber remains in it until the destination is reached. That is, the condition of interest is when the angle of incidence is greater than the critical angle, the condition of Figure 19-3(c).

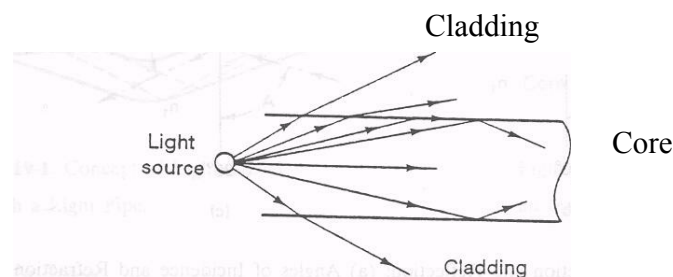


Figure 19-4. Light Directions at Input to Optical Fiber

In optical-fiber operation, the angle of incidence is set by the relationship of the light source to the source end of the fiber. Light from a typical source travels in all directions. The result is that all three of the conditions illustrated in Figure 19-3 are likely to occur in the excitation of an optical fiber, as shown in Figure 19-4. Only those rays that enter the fiber

parallel to its axis or which have incidence angles greater than the critical angle will be propagated in the fiber.

Part I. Comprehension Exercises

A. Put “T” for true and “F” for false statements. Justify your answers.

- 1. Optical systems are used for both analog and digital communications systems.
- 2. Satellite and optical links are identical.
- 3. Optical fibers for communications purposes are as simple as light pumps.
- 4. The index of refraction of a material is directly proportional to the velocity of propagation of light in the material.
- 5. It may be concluded from the text that when the refraction angle is 90° , the ray is absorbed
- 6. The best condition in optical-fiber operation is when the angle of incidence is greater than the critical angle.

B. Choose a, b, c, or d which best completes each item.

- 1. According to the text,
 - a. an optic fiber is quite similar in appearance to a twisted-wire-pair cable
 - b. an optic fiber has the same attenuation characteristics as those of a twisted-wire-pair cable
 - c. an optical transmission line performs the same functions as a microwave link but in a different manner
 - d. an optical transmission line performs function different from those performed by a microwave link
- 2. We may infer from the text that
 - a. optical fibers are used for light transmissions in a manner virtually identical to waveguides at microwave frequencies
 - b. an optic fiber is a piece of very thin, highly pure glass with the same refractive index as the outside cladding
 - c. fiber optic system has fully replaced other communications systems
 - d. fiber optic system is not capable of taking over communication traffic handled by satellite links
- 3. In a fiber-optic communications system, modulators must be used to
 - a. cause the signal to be carried away from the source

- b. cause the light wave to travel down the fiber
 - c. convert the optical signal back into an electrical signal
 - d. impress data or an analog signal on the light beam
4. When the angle of incidence is smaller than the critical angle, the ray
- a. is not able to cross the boundary.
 - b. crosses the boundary and escapes
 - c. is not able to escape
 - d. travels along the boundary
5. As we understand from the text,
- a. the speed of light reduces in materials other than the air and this reduction results in refraction
 - b. the speed of light reduces in materials other than the air but it has nothing to do with refraction
 - c. the angle of refraction increases as the material through which light passes becomes denser
 - d. the angle of incidence of an arriving light ray is directly proportional to its angle of refraction

C. Answer the following questions orally.

1. What does an optical communications system use as a carrier?
2. What are some of the terminal facilities used in an optical system?
3. What principle are optical fibers based on?
4. What does an optical fiber consist of?
5. What are the angles of incidence and refraction?
6. How is the angle of incidence set in an optical-fiber operation?

Part II. Language Practice

A. Choose a, b, c, or d which best completes each item.

1. In two or more messages are simultaneously transmitted over the same transmission path.

a. frequency modulation	b. multiplex operation
c. amplitude modulation	d. fiber-optic operation
2. The electromagnetic waves just below visible light in frequency are the infrared waves. They are increasingly used in Communication schemes.

a. fiber-optic	b. multiplex radio
c. multiplex printing	d. multiple-tuned

3. The ratio of the phase velocity in free space to that in the medium is referred to as refraction
 - a. error
 - b. effect
 - c. index
 - d. loss
4. A is logically equivalent to a multiposition selector switch.
 - a. diplexer
 - b. multiplexer
 - c. pilot carrier
 - d. radio detector
5. The of an arriving light ray is the angle the ray makes with a line perpendicular to the media boundary at the point where the ray makes the boundary.
 - a. angle of incidence
 - b. angle of refraction
 - c. critical angle
 - d. none of the above

B. Fill in the blanks with the appropriate form of the words given.

1. Refract

- a. The speed reduction and subsequent are different for each wavelength.
- b. The index, n , is the ratio of the speed of light in free space to the speed in a given material.
- c. When the refraction angle is 90° , the ray goes along the interface.
- d. Electromagnetic waves traveling from a rarer to a denser medium are toward the perpendicular to the boundary.

2. Bend

- a. The amount of provided by refraction depends on the refractive index of the two materials involved.
- b. Refraction causes the light to be
- c. The ratio of amplitude existing before the introduction of bend-reducing features to that existing afterward is known as bend reduction factor.

3. Carry

- a. Optical systems may be used to digital signals.
- b. The process of extracting the signal information from a modulated wave is called demodulation.
- c. The current associated with a carrier wave is called the current.

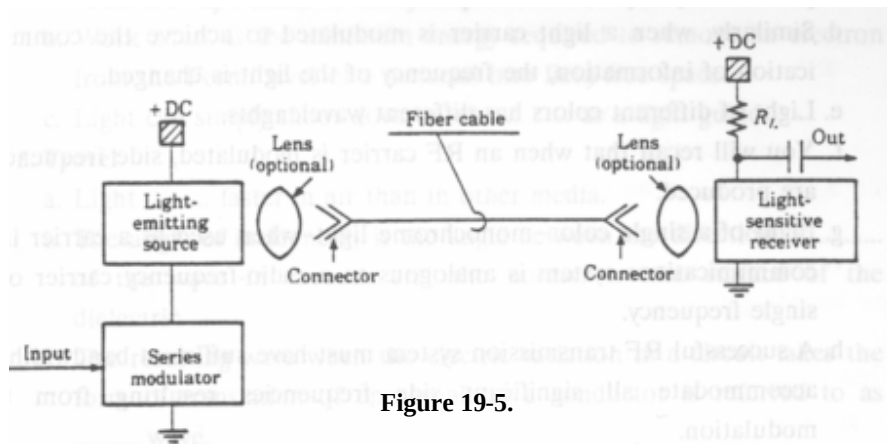


Figure 19-5.

be glass while the clad is plastic, or both may be made from plastic materials of different densities. Fiber optics is founded on the theory of reflection that results at the interface between two materials of different densities. In metallic waveguide, the energy is reflected along the guide when one-half wavelength of energy is shorter than the size of the waveguide. In fiber optics, the energy will reflect down the glass waveguide when the **angle of reflection** remains smaller than a critical angle determined by the ratio of the densities of the core and clad materials.

The cross section of Figure 19-6 illustrates the construction of a cable having a glass core $50\text{ }\mu\text{m}$ in diameter with an index of refraction of 1.45. The

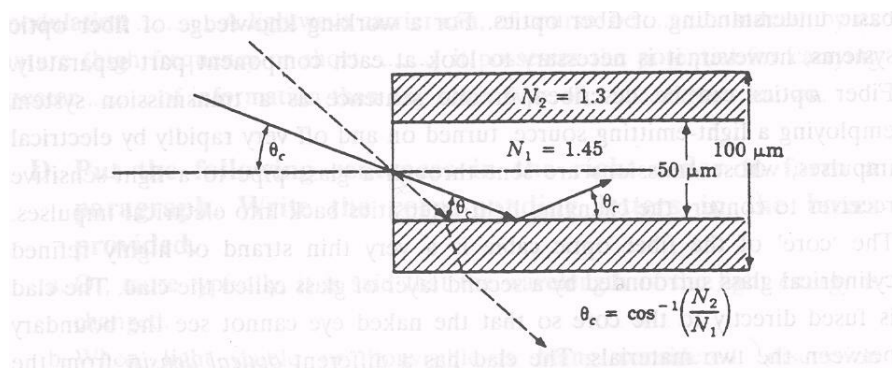


Figure 19-6. Light Reflection Inside the Guide From N_1 and N_2

cladding around the core has an outer diameter of 100 m and an index of 1.3. The clad always has a lower index of refraction than the core material. Reflections with zero loss will take place at the interface surface of the core and clad materials, provided the light energy approaches the interface at an angle that is less than the critical angle θ_c .

Refraction

Light energy traveling through a vacuum will move at a velocity of 3×10^8 m/s. It is considered to have the same velocity in our atmosphere. As light enters any transparent medium, it slows down slightly depending on the optical density of the new material. When comparing the velocity of light in free air to the velocity of light in the given medium, the ratio is a unitless number called the **index of refraction** N :

$$N = \frac{V_c}{V_m}$$

where V_c = velocity of light in air

V_m = velocity of light in the new medium

The index of refraction is a number larger than 1, which means that light in any transparent material moves slower than it does in air.

The term *refraction* identifies a directional change to a ray of light, as well as a velocity change when light crosses between two materials of different refractive indexes. Refraction is also dependent on the angle of penetration. This principle can be demonstrated easily. Figure 19-7 illustrates a person

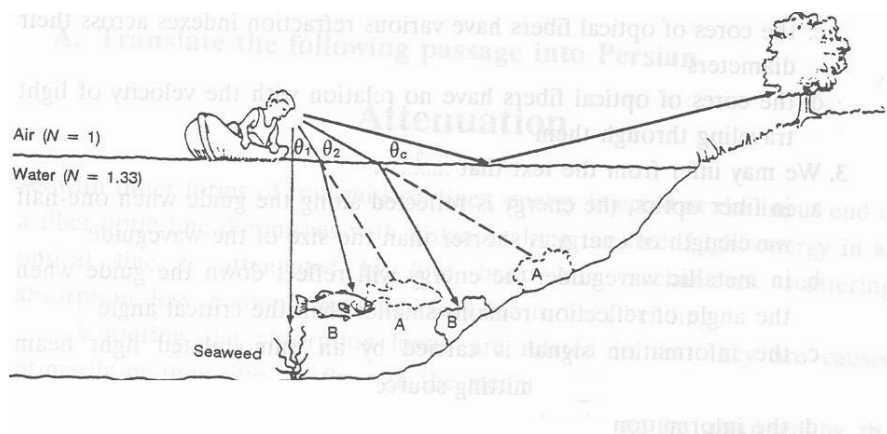


Figure 19-7. Refractions Due to a Change in Index

looking over the side of a boat into a clear pond. The images we see are lightrays that are reflected off of the subject and detected by the eye, so we can use the lines of sight as representing rays of light.

When the person looks straight down into the water, there is a change in velocity but no change in direction; the seaweed appears directly below the boat. As the viewer looks toward the shore, he thinks he sees a fish at A. The angle of penetration is θ_1 , so the light rays refract (bend) as well as reflect, which means that the fish is actually at B. When the viewer looks at the rock, the penetration angle θ_2 is smaller than θ_1 , and the illusion is greater than when the fish was viewed. The rock appears to be at location A but is really at location B. As the viewer gazes closer to the shore line, he finds that he can no longer see into the pond but rather sees the reflection of the tree on the shore. This is because the angle of entry has become smaller than the *critical* angle θ_c .

Comprehension Exercises

A. Choose a, b, c, or d which best completes each item.

1. A simple fiber-optic system
 - a. cannot be considered a transmission system
 - b. is not compatible with other systems
 - c. consists of a light-emitting source, a cable, and a receiver
 - d. consists of two lenses and a sensitive receiver

2. According to the text,
 - a. fiber optics and metallic waveguide have the same mechanism
 - b. fiber optics are based on a theory different from that of metallic waveguide
 - c. the cores of optical fibers have various refraction indexes across their diameters
 - d. the cores of optical fibers have no relation with the velocity of light traveling through them

3. We may infer from the text that
 - a. in fiber optics, the energy is reflected along the guide when one-half wavelength of energy is shorter than the size of the waveguide
 - b. in metallic waveguide, the energy will reflect down the guide when the angle of reflection remains smaller than the critical angle
 - c. the information signal is carried by an unmodulated light beam radiated from a light-emitting source
 - d. the information signal carried by a fiber-optic transmission line is in the form of a modulated light beam

4. Refraction occurs when waves
 - a. pass from one density medium to another
 - b. traveling in straight paths bend around an obstacle
 - c. travel along the earth's surface
 - d. travel through the atmosphere

5. As we understand from the text,
 - a. the angle of entry has no effect on the reflection of the objects in the water
 - b. the angle of entry only affects the reflection of the tree
 - c. the person looking into the water knows exactly where the objects in the water are
 - d. the person looking into the water is unaware of the change in velocity

B. Write the answers to the following questions.

1. What is the function of the receiver in a fiber-optic system?
2. Why should the clad have a different optical density from the core material?
3. What causes the energy to reflect down the glass waveguide?
4. How does the optical density of any material affect the velocity of light?
5. How can an outside cladding of an optic fiber be of the same material as the core?



Section three: translation activities

A. Translate the following passage into Persian.

Attenuation

As with other forms of transmission lines, energy injected at the input end of a fiber-optic line diminishes with distance along the line. Light energy in an optical line is attenuated by four basic loss mechanisms: scattering, absorption, loss in connections, and loss due to fiber bending.

Scattering and absorption losses are related in that they are caused primarily by impurities or flaws in the medium of an optical-fiber core. The ray may be deflected sufficiently to exit the core and be absorbed by the cladding or sheath of the fiber. Or, it may simply be absorbed by a particle of

opaque impurity embedded in the core material. In either case, the energy involved in the particular event is lost to the system.

Connection losses are the result of light rays encountering imperfections in the boundaries of the transmission media where two media are joined to permit the passage of light from one to the other. For example, if there is any surface roughness on the ends of the fibers where they are joined, some light will be refracted by the surface imperfections and lost to the system.

Bending losses are the result of energy lost when light waves are required to make an excessively sharp bend. When analyzing the behavior of light as a wave phenomenon, we must remember that a wave has a 'width' perpendicular to the direction of travel of the wave. When the wave bends around a corner, the outside edge of the wave must travel faster than the inside edge: If it doesn't, it isn't bending. (As an analogy, when a column of marchers goes around a corner, persons in the outside positions must step faster, or persons in the inside positions must mark time, in order for the line to remain straight during the turn.) If the bend is too sharp, part of the wave would have to travel faster than the speed of light, which it obviously cannot do. The result is that some of the light simply exits the fiber and is lost by absorption in the cladding or sheath.

The total losses of an optical communications system includes the sum of all the losses produced by the mechanisms described above. Design for maximum performance requires attention to assure minimization of each type of loss.

B. Find the Persian equivalents of the following terms and expressions and write them in the spaces provided.

- | | |
|-------------------------|-------|
| 1. angle of incidence | |
| 2. angle of reflection | |
| 3. angle of refraction | |
| 4. cladding | |
| 5. conventional | |
| 6. critical angle | |
| 7. demultiplexer | |
| 8. denser | |
| 9. diplexer | |
| 10. index of refraction | |
| 11. monochrome | |
| 12. multiplexer | |

13. optical communication system	
14. optical density	
15. optical fiber	
16. reduction factor	
17. refraction index	
18. Snell's Law	
19. terrestrial	
20. twisted-wire-pair cable	

Unit 20

Section One: Reading Comprehension

The Communications Satellite

The ultimate worldwide communications system will feature a satellite as one of the major components. Although long-range communications took place before the age of satellites, the systems suffered from conditions that required a great deal of effort to overcome. Even today, due to remote location or surrounding terrain, there are isolated communities that are difficult to reach by point-to-point communication systems. The satellite is really nothing more than a radio relay station, but it offers the one advantage that is missing in all other systems; the capability of a direct *line-of-sight* path to the earth's surface.

A satellite travels in space in a direction parallel to the surface of a planet. It has a forward velocity sufficient to create an outward thrust (centrifugal force) equal to the gravitational pull of the planet it orbits. There are three common orbital patterns; the polar orbit, the inclined elliptical orbit, and the equatorial geosynchronous orbit. The following factors apply equally to all orbits.

1. The plane of the orbit must pass through the center of the object to be orbited. For instance, a satellite could *not* orbit the earth around a latitude of 42°N.
2. The time to complete one orbit depends on the mass of the vehicle (as compared to the mass of the earth), the vehicle's velocity (dependent on the initial thrust supplied by the rocket engines and the mass of the payload), and the final orbital altitude.

To place a satellite in a position that appears to be stationary over a selected location on the earth's surface means that the vehicle must move in the same direction as the earth rotates. This final requirement eliminates the polar orbit. An inclined elliptical orbit could be in a direction and at an altitude and velocity that would appear stationary relative to a given longitude, but this orbit shifts its north/south latitudinal position.

The only orbit that meets all of these requirements is the equatorial geosynchronous orbit. It is approximately 22,000 mi or 35,400 km above the earth's surface and in a plane that includes the equator.

Satellite communication allows transoceanic links, and wide bandwidths

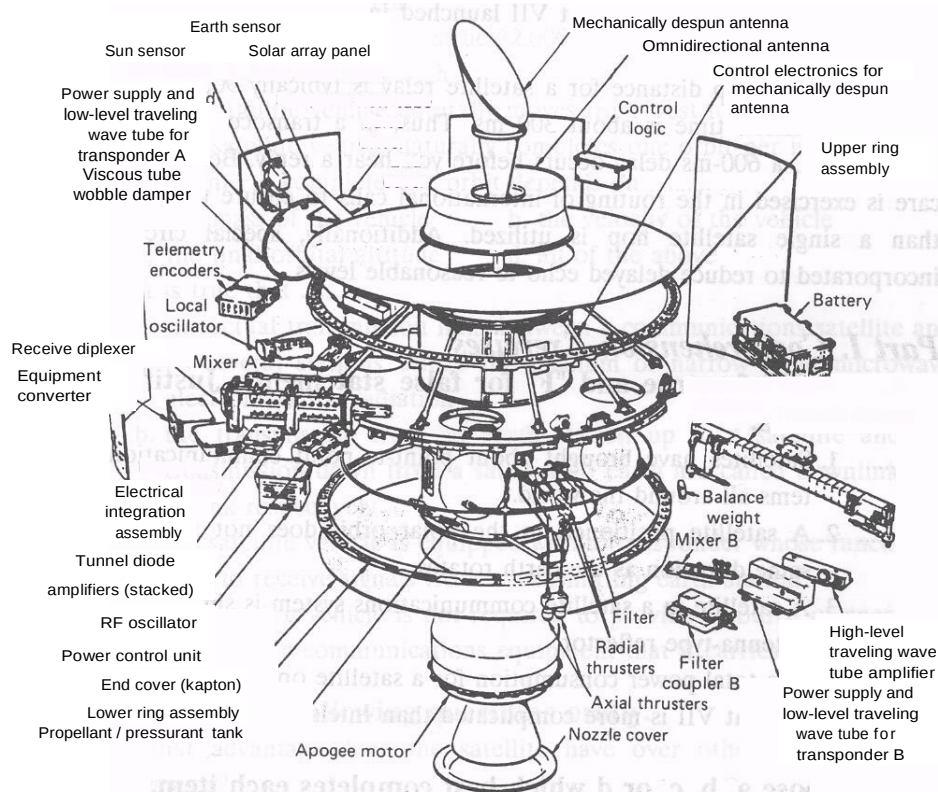


Figure 20-1.

are utilized to allow the multiplexing of a number of different signals. Frequencies used are in excess of 1 GHz. At these high frequencies, the effects of ionospheric refraction and attenuation are negligible. The frequencies used range from about 1 GHz up to 30 GHz. The signals received and subsequently retransmitted by the satellite are at different carrier frequencies. For example, the Intelsat III satellite shown in Figure 20-1 receives signals (the *uplink*) at from 5.93 to 6.42 GHz, amplifies, translates down to 3.705 to 4.195 GHz, and then reamplifies via a TWT output stage to a 7-W level for transmission back to earth (the *downlink*). The frequency translation is to prevent interference between the two signals both at the ground station and satellite. An electronic system performing the reception, frequency translation, and retransmission is called a transponder. The total power consumption for satellite operation is about 150 W. The capacity of

this satellite is for 1200 duplex voice channels or 4 TV broadcasts or any combination thereof. The Intelsat VII launched in 1992 handles 90,000 voice circuits.

The round-trip distance for a satellite relay is typically 90,000 km. The total transmission time is about 300 ms. Thus, in a transoceanic telephone conversation, a 600-ms delay occurs before you hear a reply. Because of this, care is exercised in the routing of international calls to ensure that no more than a single satellite hop is utilized. Additionally, special circuitry is incorporated to reduce delayed echo to reasonable levels.

Part I. Comprehension Exercises

A. Put “T” for true and “F” for false statements. Justify your answers.

- 1. Satellites have brought about point-to-point communication systems all around the world.
- 2. A satellite positioned in the polar orbit does not move in the same direction as the earth rotates.
- 3. A satellite in a satellite communications system is simply a passive antenna-type reflector.
- 4. The total power consumption for a satellite operation is very high.
- 5. Intelsat VII is more complicated than Intelsat III.

B. Choose a, b, c, or d which best completes each item.

- 1. As we understand from the text, a communications satellite
 - a. does not necessarily have to be equipped with highly complicated transmitters and receivers
 - b. is placed into synchronous orbit, that is, its position remains fixed with respect to the earth's rotation
 - c. can be stationed at any altitude above the earth's surface
 - d. must be placed in an orbit compatible with its weight and velocity
- 2. It can be concluded from the text that
 - a. all satellites do not travel around orbits parallel to the surface of the earth
 - b. all satellites are equally energized to have the required velocity
 - c. the plane of an orbit of the latitude of 42°N does not pass through the center of the earth
 - d. an orbit of the latitude of 42°N is a good one for a satellite to be positioned in

3. At a height of approximately 22,000 miles,
 - a. the satellite's speed is just right to keep it in synchrony with the rotation of the earth
 - b. the satellite's speed must be 22,000 mph to keep it in synchrony with the rotation of the earth
 - c. the satellite vehicle naturally moves from west to east
 - d. the satellite vehicle naturally completes one orbit per hour
4. The time to complete one orbit depends on
 - a. the mass of the vehicle
 - b. the velocity of the vehicle
 - c. the final orbital altitude
 - d. all of the above
5. It is true that
 - a. the actual transmission links between a communications satellite and its several stations utilize the medium of narrow-beam microwave electromagnetic radiation
 - b. the transmission from an earth station up to a satellite and the transmission down from a satellite to earth are called downlink and uplink respectively
 - c. the satellite vehicle is equipped with a transponder whose function is only to receive signals from a transmitting earth station
 - d. the satellite vehicle is not required to provide a source of energy to operate the communications equipment that it carries

C. Answer the following questions orally.

1. What advantage does the satellite have over other communication systems?
2. What are the three common orbital patterns?
3. What are the advantages of a geosynchronous orbit over other orbits?
4. What are the characteristics of an inclined elliptical orbit?
5. How does ionosphere affect satellite communication?
6. How is the interference between two signals at the ground station and the satellite prevented?
7. What has been done to reduce delayed echo in transoceanic telephone conversations?

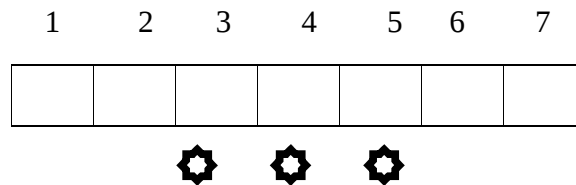
Part II. Language Practice

A. Choose a, b, c, or d which best completes each item.

1. A is a transmitter-receiver facility that transmits signals automatically when the proper interrogation is received.
 - a. relay
 - b. receiver
 - c. transponder
 - d. transmitter

must provide a very large number of equivalent individual communications channels.

- b. The development of digital communications techniques has made the utilization of TDM possible.
- c. Early satcom systems used only frequency-division multiplexing (FDM) to achieve more intensive utilization of these expensive facilities.
- d. Satellite communications systems are extremely expensive as total systems.
- e. New systems are using TDM to increase the information -carrying capacity of equipment.
- f. Ground terminal facilities are also sophisticated and costly.
- g. The launch costs-the cost of the launching rocket, fuel, launching facilities, highly skilled personnel, etc. -are a significant part of the total cost of a satellite system.



Section Two: Further Reading

Time-Division-Multiplexed Earth Terminal

A block diagram showing only the most basic of details of an earth terminal for a TDM digital satellite communications system is shown in Figure 20-2. The diagram is for a system utilizing C-band links: 6-GHz uplink and 20-2GHz downlink. You will observe that the 'sending' side of the terminal includes a multiplexer for selecting, in turn, each of the incoming information signals. For purposes of this diagram it is assumed that all incoming information signals have already been converted to some form of PCM (i.e., to digital form). These signals would typically arrive over twisted-wire pair or coaxial transmission lines from various subscribers located in the vicinity of the earth terminal.

Each incoming signal is allocated a time slot on the uplink carrier. This allocation is the result of the action of the multiplexer. An RF carrier is modulated by the digital signals. After modulation, the carrier frequency is shifted (converted) to that of the uplink-6 GHz-by means of a heterodynetype frequency converter. It is referred to as the *upconverter*. The output

of the upconverter is passed through a bandpass filter (BPF) to ensure that its bandwidth is properly limited. The signal is then amplified to increase its energy level to that sufficient for transmission over the radio link between the earth terminal antenna and the satellite antenna. Amplification is typically by means of a special electronic device for microwave frequencies. The device is called a *traveling-wave tube*, abbreviated TWT.

The antenna for the microwave frequency of the 6-GHz uplink signal is of the parabolic reflector (or 'dish') type. This antenna confines the radiation to a relatively narrow beam. A narrow-beam radiation pattern has at least two significant advantages:

1. It concentrates the radiation so as to increase the ERP (effective radiated power) in the desired direction. This effect is especially important for the downlink since the amount of power available to operate a transmitter in the satellite vehicle is extremely limited.
2. The narrow beam reduces the potential for signals intended for one satellite from interfering with other, nearby satellites.

Refer again to the block diagram of Figure 20-2. Study the portion of

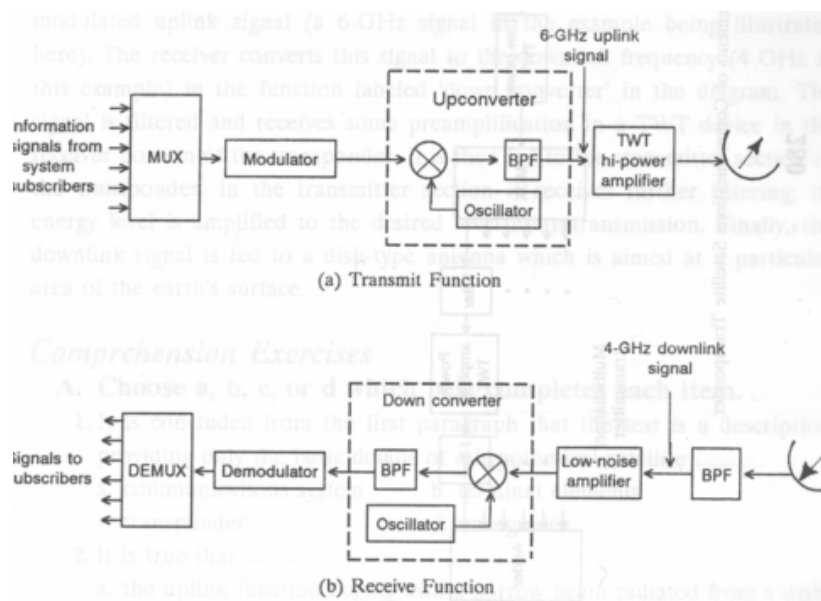


Figure 20-2. Elements of Earth Terminal for Satcom System: (a) Uplink Function; (b) Downlink Function.

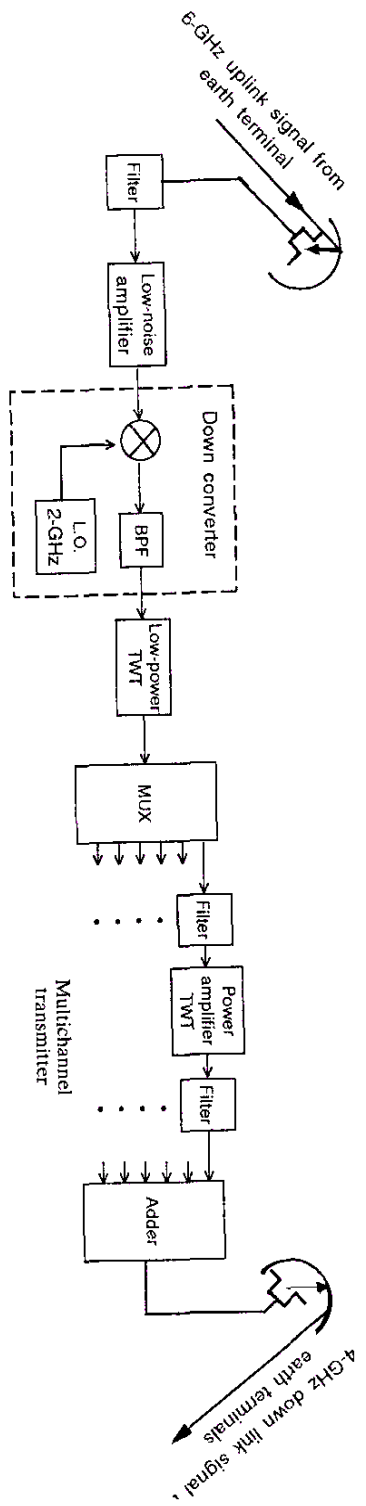


Figure 20-3. Elements of Communications Satellite Transponder.

the diagram that refers to the ‘receive’ function of the terminal. You will observe that from the antenna the signal first passes through a BPF and then a *low-noise amplifier* (LNA). These two items are typically incorporated as part of the antenna assembly. It is important to the success of the system that the downlink signal receive low-noise amplification immediately off the antenna. This measure helps to ensure that the amplitude of what may be a relatively weak signal is boosted before the noise content has become excessive. From the LNA the signal is transmitted, typically through a section of waveguide, to the downlink receiver. The downlink receiver is a form of superheterodyne receiver. The 4-GHz signal is converted in the *down converter* to a lower intermediate frequency (IF), amplified, and finally, demodulated. Remember, the carrier is transporting several information signals by means of time-division multiplexing. These information signals are now separated in a demultiplexer and sent on their way over terrestrial (land-based) facilities to their ultimate destinations.

The satellite transponder portion of a satcom system is represented in Figure 20-3. In brief, the transponder has a receiver section for receiving the modulated uplink signal (a 6-GHz signal in the example being illustrated here). The receiver converts this signal to the downlink frequency (4 GHz in this example) in the function labeled ‘down converter’ in the diagram. The signal is filtered and receives some preamplification in a TWT device in the receiver portion of the transponder. It is then fed to the transmitter section of the transponder. In the transmitter section it receives further filtering; its energy level is amplified to the desired level for retransmission. Finally, the downlink signal is fed to a dish-type antenna which is aimed at a particular area of the earth’s surface.

Comprehension Exercises

A. Choose a, b, c, or d which best completes each item.

1. It is concluded from the first paragraph that the text is a description providing only the basic details of a hypothetical satellite
 - a. communications system
 - b. terminal elements
 - c. transponder
 - d. waveguides
2. It is true that
 - a. the uplink function begins with a narrow beam radiated from a dish
 - b. the multiplexer function is to separate the information signals
 - c. the receiving and the sending sides of the terminal consist of identical elements