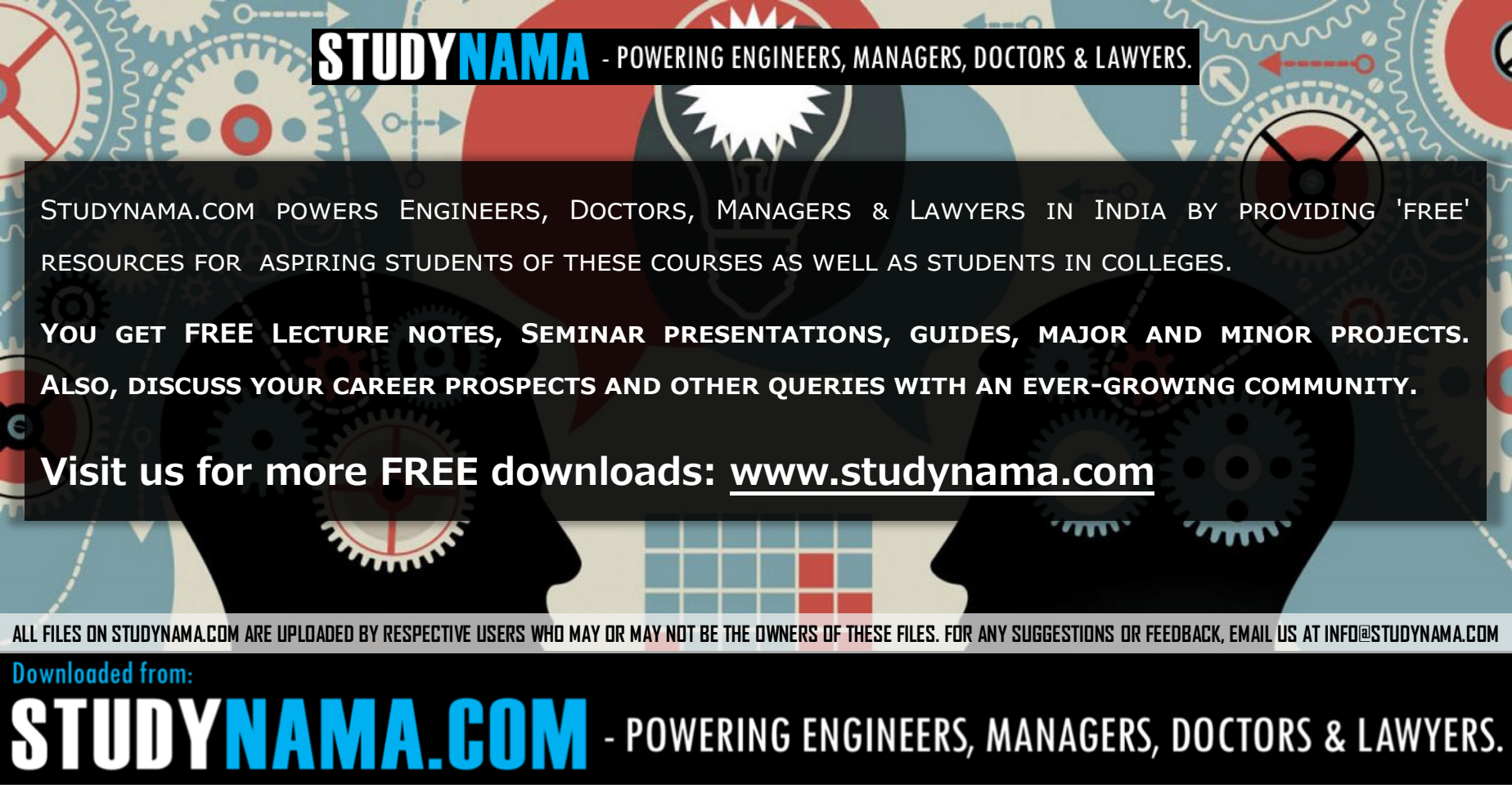




STUDYNAMA - POWERING ENGINEERS, MANAGERS, DOCTORS & LAWYERS.

STUDYNAMA.COM POWERS ENGINEERS, DOCTORS, MANAGERS & LAWYERS IN INDIA BY PROVIDING 'FREE' RESOURCES FOR ASPIRING STUDENTS OF THESE COURSES AS WELL AS STUDENTS IN COLLEGES.

YOU GET FREE LECTURE NOTES, SEMINAR PRESENTATIONS, GUIDES, MAJOR AND MINOR PROJECTS. ALSO, DISCUSS YOUR CAREER PROSPECTS AND OTHER QUERIES WITH AN EVER-GROWING COMMUNITY.

Visit us for more FREE downloads: www.studynama.com

ALL FILES ON STUDYNAMA.COM ARE UPLOADED BY RESPECTIVE USERS WHO MAY OR MAY NOT BE THE OWNERS OF THESE FILES. FOR ANY SUGGESTIONS OR FEEDBACK, EMAIL US AT [INFO@STUDYNAMA.COM](mailto:info@studynama.com)

Downloaded from:

STUDYNAMA.COM - POWERING ENGINEERS, MANAGERS, DOCTORS & LAWYERS.

BE MECH & CIVIL - I/ SEMESTER II

10AEE02

BASIC ELECTRICAL AND ELECTRONICS ENGINEERING

10AEE02 BASIC ELECTRICAL AND ELECTRONICS ENGINEERING

UNIT I ELECTRICAL CIRCUITS AND MEASUREMENTS

Ohm's Law – Kirchhoff's Laws – Steady State Solution of DC Circuits – Introduction to AC Circuits – Waveforms and RMS Value - Power and Power factor – Single Phase and Three Phase Balanced Circuits. Operating Principles of Moving Coil and Moving Iron Instruments (Ammeters and Voltmeters). Dynamometer type Watt meters and Energy meters.

UNIT II ELECTRICAL MACHINES

Construction, Principle of Operation, Basic Equations and Applications of DC Generators, DC Motors, Single Phase Transformer, Single Phase Induction Motor.

UNIT III SEMICONDUCTOR DEVICES AND APPLICATIONS

Characteristics of PN Junction Diode – Zener Effect - Zener Diode and its Characteristics - Half wave and Full wave Rectifiers – Voltage Regulation. Bipolar Junction Transistor – CB, CE, CC Configurations and Characteristics – Elementary Treatment of Signal Amplifier.

UNIT IV DIGITAL ELECTRONICS

Binary Number System – Logic Gates – Boolean Algebra – Half and Full Adders - Flip – Flops - Registers and Counters – A/D and D/A Conversion (simple concepts).

UNIT V FUNDAMENTALS OF COMMUNICATION ENGINEERING

Types of Signals: Analog and Digital Signals – Modulation and Demodulation: Principles of Amplitude and Frequency Modulations. Communication System: Radio, TV, Fax, Microwave, Satellite and Optical Fiber (Block Diagram Approach only).

UNIT I

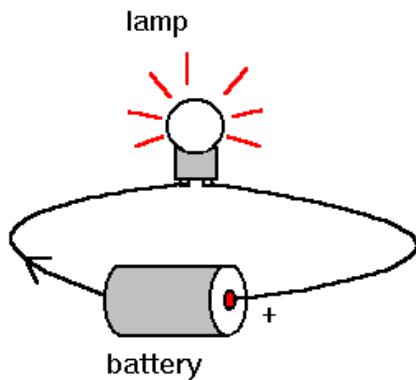
ELECTRICAL CIRCUITS AND MEASUREMENTS

DC Circuits:

- A DC circuit (Direct Current circuit) is an electrical circuit that consists of any combination of constant voltage sources, constant current sources, and resistors. In this case, the circuit voltages and currents are constant, i.e., independent of time. More technically, a DC circuit has no memory. That is, a particular circuit voltage or current does not depend on the past value of any circuit voltage or current. This implies that the system of equations that represent a DC circuit do not involve integrals or derivatives.
- If a capacitor and/or inductor is added to a DC circuit, the resulting circuit is not, strictly speaking, a DC circuit. However, most such circuits have a DC solution. This solution gives the circuit voltages and currents when the circuit is in DC steady state. More technically, such a circuit is represented by a system of differential equations. The solution to these equations usually contain a time varying or transient part as well as constant or steady state part. It is this steady state part that is the DC solution. There are some circuits that do not have a DC solution. Two simple examples are a constant current source connected to a capacitor and a constant voltage source connected to an inductor.
- In electronics, it is common to refer to a circuit that is powered by a DC voltage source such as a battery or the output of a DC power supply as a DC circuit even though what is meant is that the circuit is DC powered.

Electric Current:

Electric current means, depending on the context, a flow of electric charge (a phenomenon) or the rate of flow of electric charge (a quantity). This flowing electric charge is typically carried by moving electrons, in a conductor such as wire; in an electrolyte, it is instead carried by ions, and, in a plasma, by both. The SI unit for measuring the rate of flow of electric charge is the ampere, which is charge flowing through some surface at the rate of one coulomb per second. Electric current is measured using an ammeter.



Current:

The flow of charge is called the current and it is the rate at which electric charges pass through a conductor. The charged particle can be either positive or negative. In order for a charge to flow, it needs a push (a force) and it is supplied by voltage, or potential difference. The charge flows from high potential energy to low potential energy.

$$\text{Current} = I = \frac{Q}{t}$$

Where the symbol **I** to represent the quantity current.

Electro-magnetic force(E.M.F):

Electromotive Force is, the voltage produced by an electric battery or generator in an electrical circuit or, more precisely, the energy supplied by a source of electric power in driving a unit charge around the circuit. The unit is the volt. A difference in charge between two points in a material can be created by an external energy source such as a battery. This causes electrons to move so that there is an excess of electrons at one point and a deficiency of electrons at a second point. This difference in charge is stored as electrical potential energy known as emf. It is the emf that causes a current to flow through a circuit.

Voltage:

Voltage is electric potential energy per unit charge, measured in joules per coulomb (= volts). It is often referred to as "electric potential", which then must be distinguished from electric potential energy by noting that the "potential" is a "per-unit-charge" quantity. Like mechanical potential energy, the zero of potential can be chosen at any point, so the difference in voltage is the quantity which is physically meaningful. The difference in voltage measured when moving from point A to point B is equal to the work which would have to be done, per unit charge, against the electric field to move the charge from A to B.

Electric potential:

A gravitational analogy was relied upon to explain the reasoning behind the relationship between location and potential energy. Moving a positive test charge against the direction of

an electric field is like moving a mass upward within Earth's gravitational field. Both movements would be like *going against nature* and would require work by an external force. This work would in turn increase the potential energy of the object. On the other hand, the movement of a positive test charge in the direction of an electric field would be like a mass falling downward within Earth's gravitational field. Both movements would be like *going with nature* and would occur without the need of work by an external force. This motion would result in the loss of potential energy. Potential energy is the stored energy of position of an object and it is related to the location of the object within a field.

Potential Difference:

A quantity related to the amount of energy needed to move an object from one place to another against various types of forces. The term is most often used as an abbreviation of "electrical potential difference", but it also occurs in many other branches of physics. Only changes in potential or potential energy (not the absolute values) can be measured.

Electrical potential difference is the voltage between two points, or the voltage drop transversely over an impedance (from one extremity to another). It is related to the energy needed to move a unit of electrical charge from one point to the other against the electrostatic field that is present. The unit of electrical potential difference is the volt (joule per coulomb). Gravitational potential difference between two points on Earth is related to the energy needed to move a unit mass from one point to the other against the Earth's gravitational field. The unit of gravitational potential differences is joules per kilogram.

Resistance:

Resistance is the ratio of potential difference across a conductor to the current flowing through it. If energy is used in passing electricity through an object, that object has a resistance.

Electromagnetism:

What is Electromagnetism?

When current passes through a conductor, magnetic field will be generated around the conductor and the conductor becomes a magnet. This phenomenon is called electromagnetism. Since the magnet is produced by electric current, it is called the electromagnet. An electromagnet is a type of magnet in which the magnetic field is produced by a flow of electric current. The magnetic field disappears when the current ceases. In short, when current flows through a conductor, magnetic field will be generated. When the current ceases, the magnetic field disappears.

Applications of Electromagnetism:

- Electromagnetism has numerous applications in today's world of science and physics. The very basic application of electromagnetism is in the use of motors. The motor has a switch that continuously switches the polarity of the outside of motor. An electromagnet does the same thing. We can change the direction by simply reversing the current. The inside of the motor has an electromagnet, but the current is controlled in such a way that the outside magnet repels it.
- Another very useful application of electromagnetism is the "CAT scan machine." This machine is usually used in hospitals to diagnose a disease. As we know that current is present in our body and the stronger the current, the stronger is the magnetic field. This scanning technology is able to pick up the magnetic fields, and it can be easily identified where there is a great amount of electrical activity inside the body.
- The work of the human brain is based on electromagnetism. Electrical impulses cause the operations inside the brain and it has some magnetic field. When two magnetic fields cross each other inside the brain, interference occurs which is not healthy for the brain.

Ohm's Law:

Ohm's law states that the current through a conductor between two points is directly proportional to the potential difference or voltage across the two points, and inversely proportional to the resistance between them. The mathematical equation that describes this relationship is:

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

where I is the current through the resistance in units of amperes, V is the potential difference measured across the resistance in units of volts, and R is the resistance of the conductor in units of ohms. More specifically, Ohm's law states that the R in this relation is constant, independent of the current.

Resistance:

Resistance is the opposition that a substance offers to the flow of electric current. It is represented by the uppercase letter R. The standard unit of resistance is the ohm, sometimes written out as a word, and sometimes symbolized by the uppercase Greek letter omega. When an electric current of one ampere passes through a component across which a potential difference (voltage) of one volt exists, then the resistance of that component is one ohm.

In general, when the applied voltage is held constant, the current in a direct-current (DC) electrical circuit is inversely proportional to the resistance. If the resistance is doubled, the current is cut in half; if the resistance is halved, the current is doubled. This rule also holds true for most low-frequency alternating-current (AC) systems, such as household utility circuits.

In some AC circuits, especially at high frequencies, the situation is more complex, because some components in these systems can store and release energy, as well as dissipating or converting it. The electrical resistance per unit length, area, or volume of a substance is known as resistivity. Resistivity figures are often specified for copper and aluminum wire, in ohms per kilometre. Opposition to AC, but not to DC, is a property known as reactance. In an AC circuit, the resistance and reactance combine vectorially to yield impedance.

Voltage:

Introduction:

The voltage between two points is a short name for the electrical force that would drive an electric current between those points. Specifically, voltage is equal to energy per unit charge. In the case of static electric fields, the voltage between two points is equal to the electrical potential difference between those points. In the more general case with electric and magnetic fields that vary with time, the terms are no longer synonymous.

Electric potential is the energy required to move a unit electric charge to a particular place in a static electric field. Voltage can be measured by a voltmeter. The unit of measurement is the volt.

What is voltage?

Voltage should be more correctly called "potential difference". It is actually the electron moving force in electricity (emf) and the potential difference is responsible for the pushing and pulling of electrons or electric current through a circuit.

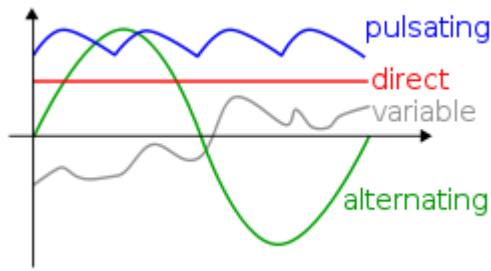
AC Circuits:

Fundamentals of AC:

● An alternating current (AC) is an electrical current, where the magnitude of the current varies in a cyclical form, as opposed to direct current, where the polarity of the current stays constant.

● The usual waveform of an AC circuit is generally that of a sine wave, as this results in the most efficient transmission of energy. However in certain applications different waveforms are used, such as triangular or square waves.

● Used generically, AC refers to the form in which electricity is delivered to businesses and residences. However, audio and radio signals carried on electrical wire are also examples of alternating current. In these applications, an important goal is often the recovery of information encoded (or modulated) onto the AC signal.



Alternating Current (green curve)

AC Instantaneous and RMS:

Instantaneous Value:

The INSTANTANEOUS value of an alternating voltage or current is the value of voltage or current at one particular instant. The value may be zero if the particular instant is the time in the cycle at which the polarity of the voltage is changing. It may also be the same as the peak value, if the selected instant is the time in the cycle at which the voltage or current stops increasing and starts decreasing. There are actually an infinite number of instantaneous values between zero and the peak value.

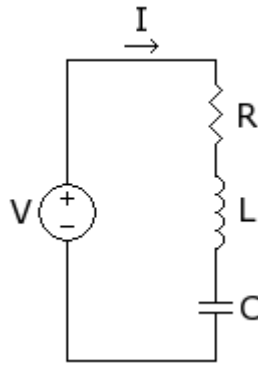
RMS Value:

The average value of an AC waveform is NOT the same value as that for a DC waveforms average value. This is because the AC waveform is constantly changing with time and the heating effect given by the formula ($P = I^2 \cdot R$), will also be changing producing a positive power consumption. The equivalent average value for an alternating current system that provides the same power to the load as a DC equivalent circuit is called the "effective value". This effective power in an alternating current system is therefore equal to: ($I^2 \cdot R \cdot \text{Average}$). As power is proportional to current squared, the effective current, I will be equal to $\sqrt{I^2 \text{ Ave}}$. Therefore, the effective current in an AC system is called the Root Mean Squared or R.M.S.

RLC Series Circuit:

An RLC circuit (or LCR circuit) is an electrical circuit consisting of a resistor, an inductor, and a capacitor, connected in series. I is the current through the circuit.

- $V_R = IR$, voltage drop across R
- $V_L = IX_L$, voltage drop across L
- $V_C = IX_C$, voltage drop across C



RLC Series Circuit

DIFFERENCE BETWEEN AC AND DC:

Current that flows continuously in one direction is called direct current. Alternating current (A.C) is the current that flows in one direction for a brief time then reverses and flows in opposite direction for a similar time. The source for alternating current is called a.c generator or alternator.

Cycle:

One complete set of positive and negative values of an alternating quantity is called cycle.

Frequency:

The number of cycles made by an alternating quantity per second is called frequency. The unit of frequency is Hertz(Hz)

Amplitude or Peak value

The maximum positive or negative value of an alternating quantity is called amplitude or peak value.

Average value:

This is the average of instantaneous values of an alternating quantity over one complete cycle of the wave.

Time period:

The time taken to complete one complete cycle.

Average value derivation:

Let i = the instantaneous value of current

And $i = I_m \sin \theta$

Where, I_m is the maximum value.

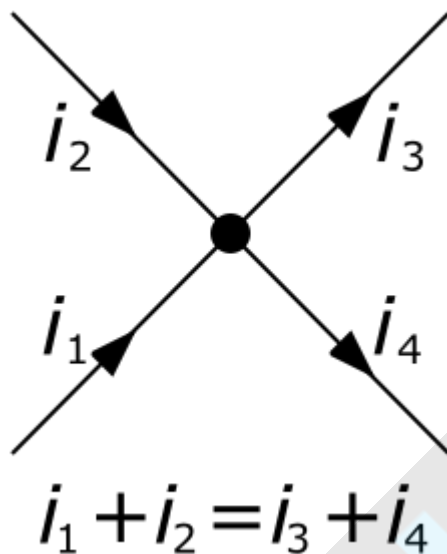
Kirchhoff's law:

Kirchoff's Current Law:

First law (Current law or Point law):

The sum of the currents flowing towards any junction in an electric circuit equal to the sum of currents flowing away from the junction.

Kirchoff's Current law can be stated in words as the sum of all currents flowing into a node is zero. Or conversely, the sum of all currents leaving a node must be zero. As the image below demonstrates, the sum of currents I_b , I_c , and I_d , must equal the total current in I_a . Current flows through wires much like water flows through pipes. If you have a definite amount of water entering a closed pipe system, the amount of water that enters the system must equal the amount of water that exists the system. The number of branching pipes does not change the net volume of water (or current in our case) in the system.



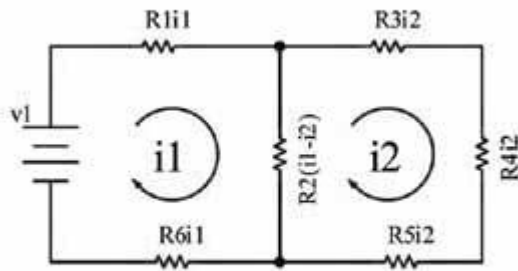
Kirchoff's Voltage Law:

Second law (voltage law or Mesh law):

In any closed circuit or mesh, the algebraic sum of all the electromotive forces and the voltage drops is equal to zero.

Kirchoff's voltage law can be stated in words as the sum of all voltage drops and rises in a closed loop equals zero. As the image below demonstrates, loop 1 and loop 2 are both closed loops within the circuit. The sum of all voltage drops and rises around loop 1 equals zero, and the sum of all voltage drops and rises in loop 2 must also equal zero. A closed loop can be defined as any path in which the originating point in the loop is also the ending point for the loop. No matter how the loop is defined or drawn, the sum of the voltages in the loop must be zero.

Figure 2. Kirchoff's Voltage Law



$$-v_1 + R_1 i_1 + R_2(i_1 - i_2) + R_6 i_2 = 0$$

$$R_2(i_2 - i_1) + R_3 i_2 + R_4 i_2 + R_5 i_2 = 0$$

Problems and Calculations:

Problem 1:

A current of 0.5 A is flowing through the resistance of 10Ω. Find the potential difference between its ends.

Solution:

Current $I = 0.5\text{A}$.

Resistance $R = 10\Omega$

Potential difference $V = ?$

$$V = IR$$

$$= 0.5 \times 10 = 5\text{V}.$$

Problem :2

A supply voltage of 220V is applied to a 100 Ω resistor. Find the current flowing through it.

Solution:

Voltage $V = 220\text{V}$

Resistance $R = 100\Omega$

$$\text{Current } I = \frac{V}{R} = \frac{220}{100} = 2.2 \text{ A}.$$

Problem : 3

Calculate the resistance of the conductor if a current of 2A flows through it when the potential difference across its ends is 6V.

Solution:

Current $I = 2\text{A}$.

Potential difference $= V = 6$.

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

$$= 6/2$$

$$= 3 \text{ ohm.}$$

Problem: 4

Calculate the current and resistance of a 100 W ,200V electric bulb.

Solution:

$$\text{Power, } P = 100\text{W}$$

$$\text{Voltage, } V = 200\text{V}$$

$$\text{Power } p = VI$$

$$\text{Current } I = P/V = 100/200 = 0.5\text{A}$$

$$\text{Resistance } R = V / I = 200/0.5 = 400\text{W.}$$

Problem: 5

Calculate the power rating of the heater coil when used on 220V supply taking 5 Amps.

Solution:

$$\text{Voltage , } V = 220\text{V}$$

$$\text{Current , } I = 5\text{A,}$$

$$\begin{aligned} \text{Power, } P &= VI = 220 \times 5 \\ &= 1100\text{W} = 1.1 \text{ KW.} \end{aligned}$$

Problem: 6

A circuit is made of 0.4Ω wire, a 150Ω bulb and a 120Ω rheostat connected in series. Determine the total resistance of the resistance of the circuit.

Solution:

$$\text{Resistance of the wire} = 0.4\Omega$$

$$\text{Resistance of bulb} = 150\Omega$$

$$\text{Resistance of rheostat} = 120\Omega$$

In series,

$$\begin{aligned} \text{Total resistance , } R &= 0.4 + 150 + 120 \\ &= 270.4\Omega \end{aligned}$$

Problem :7

In the circuit shown in fig .find the current, voltage drop across each resistor and the power dissipated in each resistor.

Solution:

$$\text{Total resistance of the circuit} = 2 + 6 + 7$$

$$R = 15 \Omega$$

$$\begin{aligned}\text{Voltage ,V} &= 45\text{V} \\ \text{Circuit current ,I} &= V/R = 45/15 = 3\text{A}\end{aligned}$$

$$\begin{aligned}\text{Voltage drop across } 2\Omega \text{ resistor } V_1 &= I R_1 \\ &= 3 \times 2 = 6 \text{ Volts.}\end{aligned}$$

$$\begin{aligned}\text{Voltage drop across } 6\Omega \text{ resistor } V_2 &= I R_2 \\ &= 3 \times 6 = 18 \text{ volts.}\end{aligned}$$

$$\begin{aligned}\text{Voltage drop across } 7\Omega \text{ resistor } V_3 &= I R_3 \\ &= 3 \times 7 = 21 \text{ volts.}\end{aligned}$$

$$\begin{aligned}\text{Power dissipated in } R_1 \text{ is } P_1 &= P R_1 \\ &= 3^2 \times 2 = 18 \text{ watts.}\end{aligned}$$

$$\begin{aligned}\text{Power dissipated in } R_2 \text{ is } P_2 &= I^2 R_2. \\ &= 3^2 \times 6 = 54 \text{ watts.}\end{aligned}$$

$$\begin{aligned}\text{Power dissipated in } R_3 \text{ is } P_3 &= I^2 R_3. \\ &= 3^2 \times 7 = 63 \text{ watts.}\end{aligned}$$

Problem : 8

Three resistances of values 2Ω , 3Ω and 5Ω are connected in series across 20 V, D.C supply .Calculate (a) equivalent resistance of the circuit (b) the total current of the circuit (c) the voltage drop across each resistor and (d) the power dissipated in each resistor.

Solution:

$$\begin{aligned}\text{Total resistance } R &= R_1 + R_2 + R_3. \\ &= 2 + 3 + 5 = 10\Omega\end{aligned}$$

$$\text{Voltage} = 20\text{V}$$

$$\text{Total current } I = V/R = 20/10 = 2\text{A.}$$

$$\begin{aligned}\text{Voltage drop across } 2\Omega \text{ resistor } V_1 &= I R_1 \\ &= 2 \times 2 = 4 \text{ volts.}\end{aligned}$$

$$\begin{aligned}\text{Voltage drop across } 3\Omega \text{ resistor } V_2 &= I R_2 \\ &= 2 \times 3 = 6 \text{ volts.}\end{aligned}$$

$$\begin{aligned}\text{Voltage drop across } 5\Omega \text{ resistor } V_3 &= I R_3 \\ &= 2 \times 5 = 10 \text{ volts.}\end{aligned}$$

$$\begin{aligned}\text{Power dissipated in } 2\Omega \text{ resistor is } P_1 &= I^2 R_1 \\ &= 2^2 \times 2 = 8 \text{ watts.}\end{aligned}$$

$$\begin{aligned}\text{Power dissipated in } 3 \text{ resistor is } P_2 &= I^2 R_2. \\ &= 2^2 \times 3 = 12 \text{ watts.}\end{aligned}$$

$$\begin{aligned}\text{Power dissipated in } 5 \text{ resistor is } P_3 &= I^2 R_3 \\ &= 2^2 \times 5 = 20 \text{ watts.}\end{aligned}$$

Problem: 9

A lamp can work on a 50 volt mains taking 2 amps. What value of the resistance must be connected in series with it so that it can be operated from 200 volt mains giving the same power.

Solution:

Lamp voltage, $V = 50V$

Current, $I = 2$ amps.

Resistance of the lamp $= V/I = 50/2 = 25 \Omega$

Resistance connected in series with lamp $= r$.

Supply voltage $= 200$ volt.

Circuit current $I = 2A$

Total resistance $R_t = V/I = 200/2 = 100\Omega$

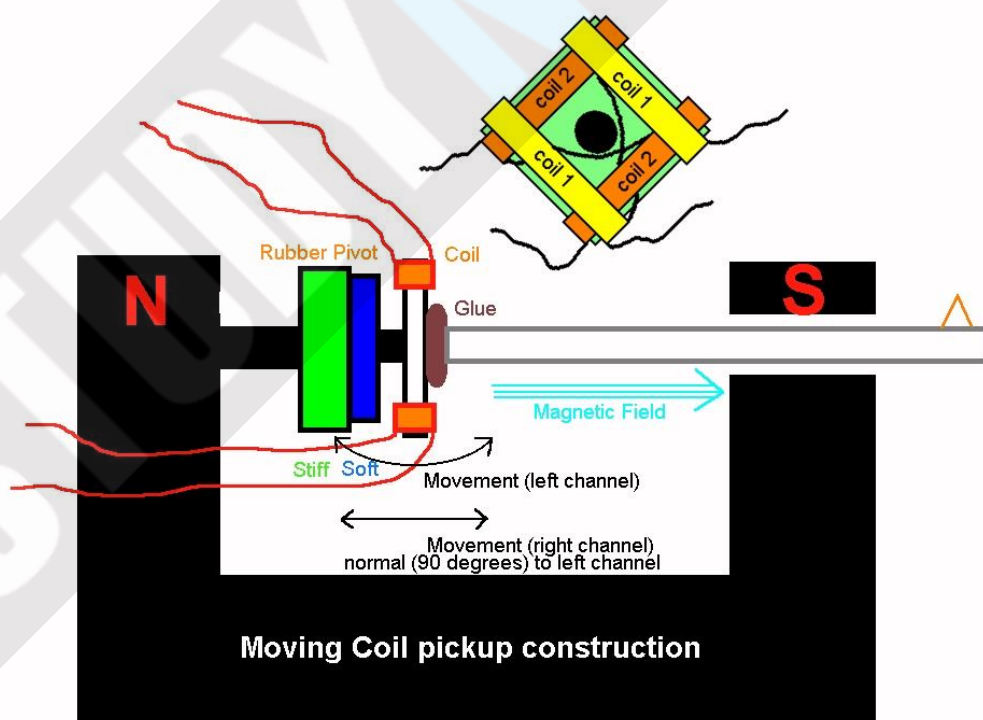
$$R_t = R + r$$

$$100 = 25 + r$$

$$r = 75\Omega$$

Moving Coil

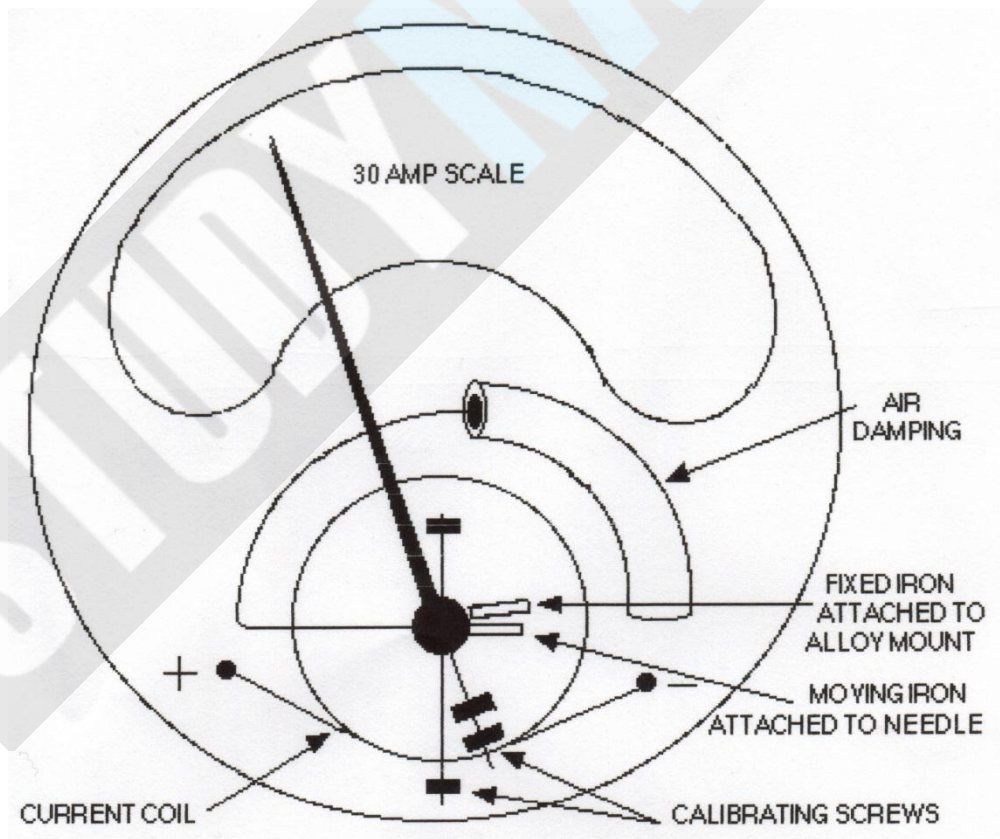
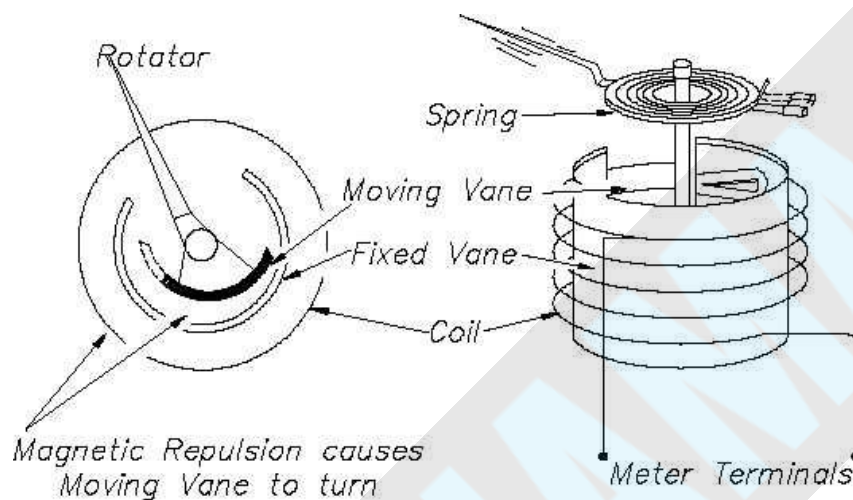
Moving Coil Instruments are used for measuring DC quantities. They can be used on AC systems when fed through bridge rectifiers. Center magnet system is incorporated in our moving coil instruments which completely shields the movement from the effect of external magnetic fields. The movement is pivoted between synthetic sapphire jewel bearings for frictionless operation.



Fig(Moving Coil)

Moving Iron

Moving Iron Instruments are generally used for measuring AC Voltage and Currents. A feature of the moving element is that it is fitted with synthetic sapphire jewels. The movement is light, quick acting, but extremely robust. An efficient system of fluid damping is employed. The movement is efficiently shielded against the effect of external magnetic fields.





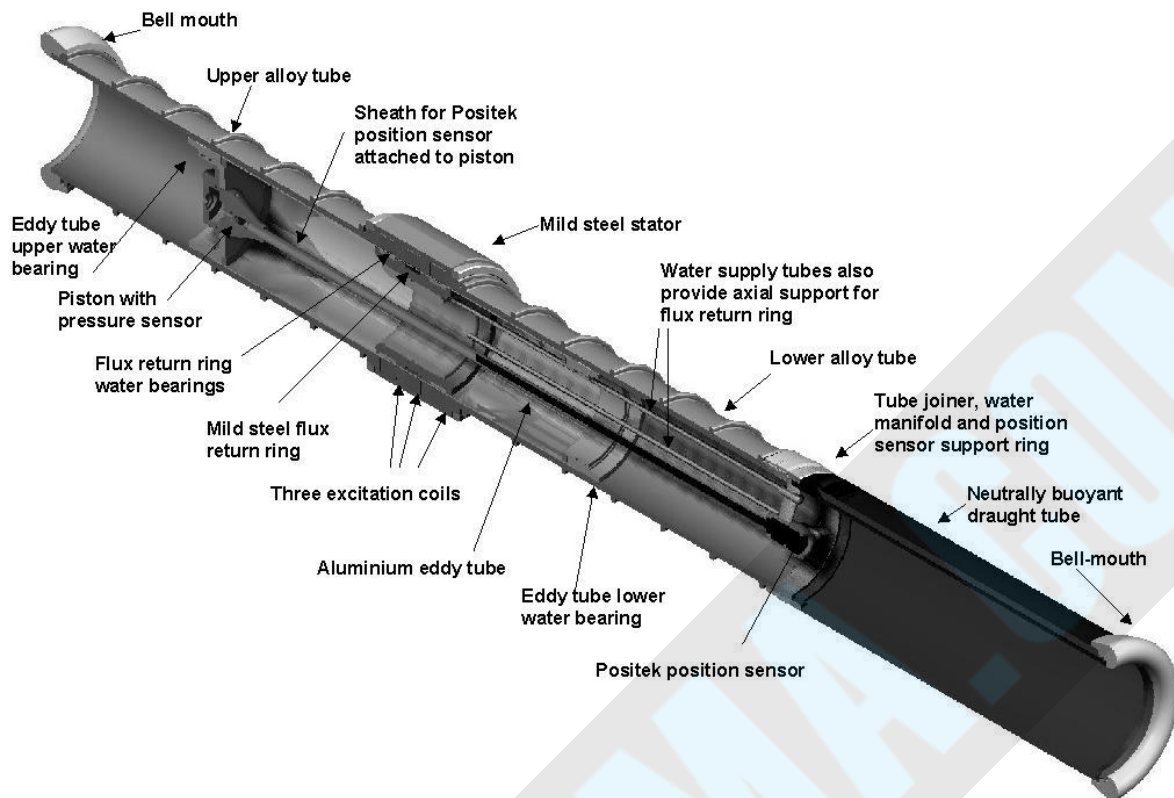
Fig(Moving Iron)

Dynamometer:

A **dynamometer** or "**dyno**" for short, is a device for measuring force, moment of force (torque), or power. For example, the power produced by an engine, motor or other rotating prime mover can be calculated by simultaneously measuring torque and rotational speed (rpm).

A dynamometer can also be used to determine the torque and power required to operate a driven machine such as a pump. In that case, a motoring or driving dynamometer is used. A dynamometer that is designed to be driven is called an absorption or passive dynamometer. A dynamometer that can either drive or absorb is called a universal or active dynamometer.

Dynamometers are specialized instruments used to measure an engine's revolutions per minute (RPM) and torque. RPM is a measurement of the number of times the crankshaft revolves inside an engine. The more revolutions a crankshaft makes each minute, the faster and more powerful the engine is.



Fig(dynamometer Outline)

Energy Meter:

Introduction

The energy meter is an electrical measuring device, which is used to record Electrical Energy Consumed over a specified period of time in terms of units. Electric meters are typically calibrated in billing units, the most common one being the kilowatt hour. A periodic reading of electric meters establishes billing cycles and energy used during a cycle.

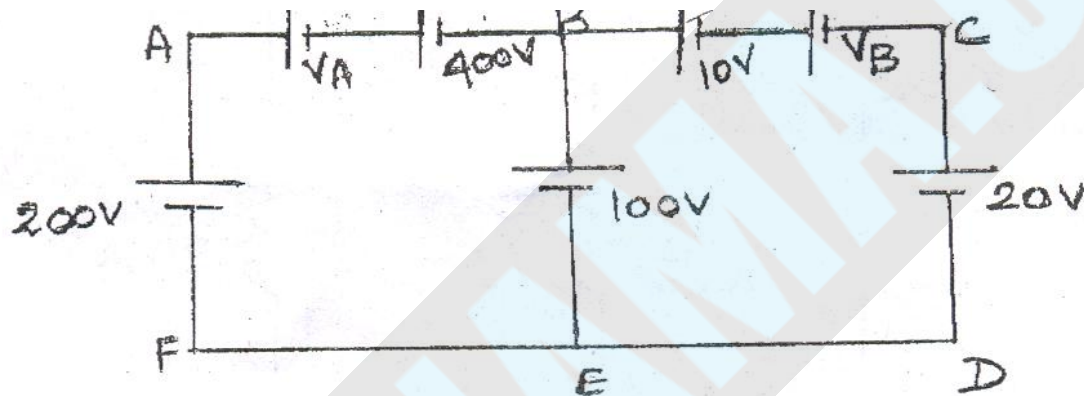
Features:

- * Display of current time (24 hours type), week, load power and cost tariff.
- * Display of total on time, total used energy and accrued energy cost.
- * Display of total record time, total on time and percentage.
- * Dual programmable power tariffs.
- * Connection, operation settings.

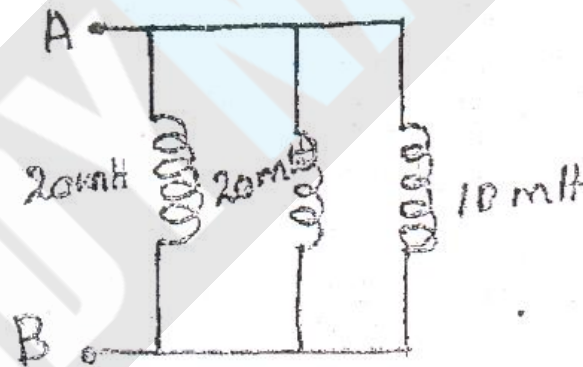
Question Bank

Part-A(2 Marks)

1. Explain star and delta connection of impedances.
2. State Superposition theorem.
3. State Thevenin's theorem.
4. State Norton's theorem.
5. State Maximum Power transfer theorem
6. State and explain KCL.
7. State and explain KVL.
8. Explain duality.
9. Determine the Missing Voltage across the elements in the circuit

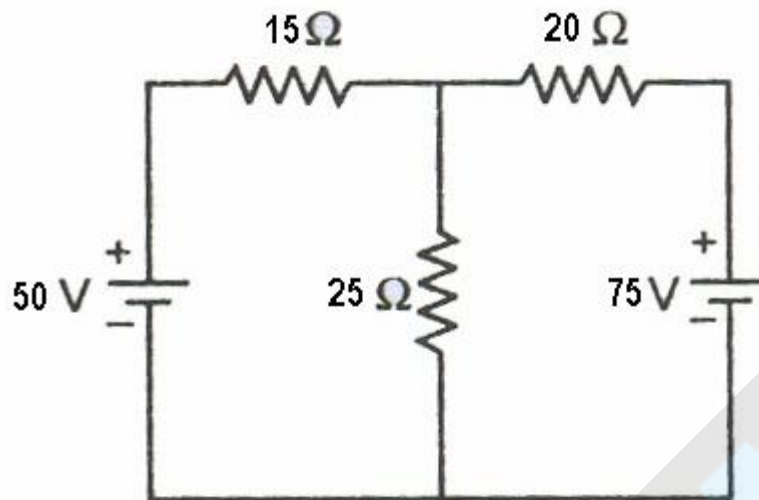


10. Find total Inductance

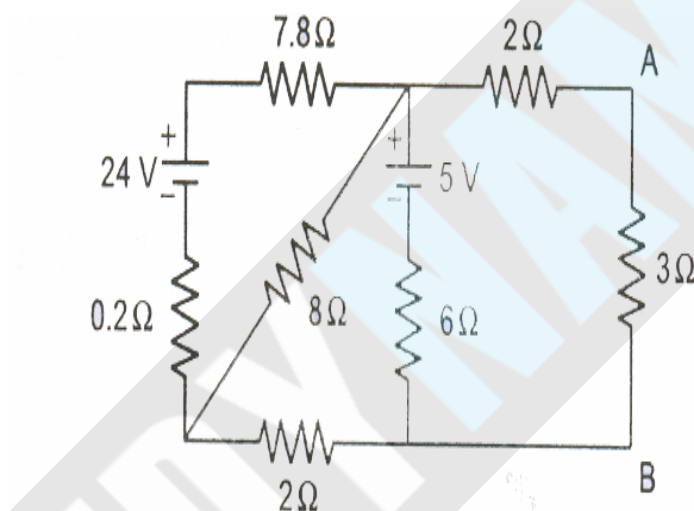


PART – B (16 Marks)

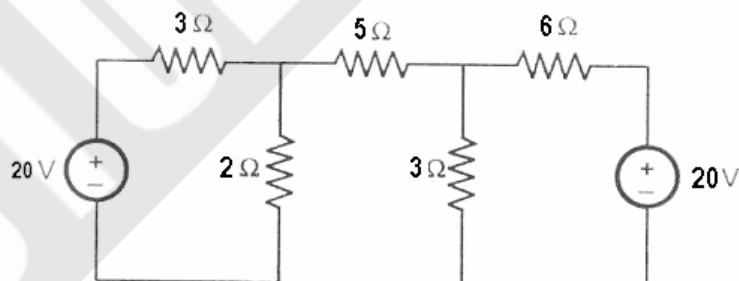
1. Apply KCL and KVL to the circuit shown in fig.



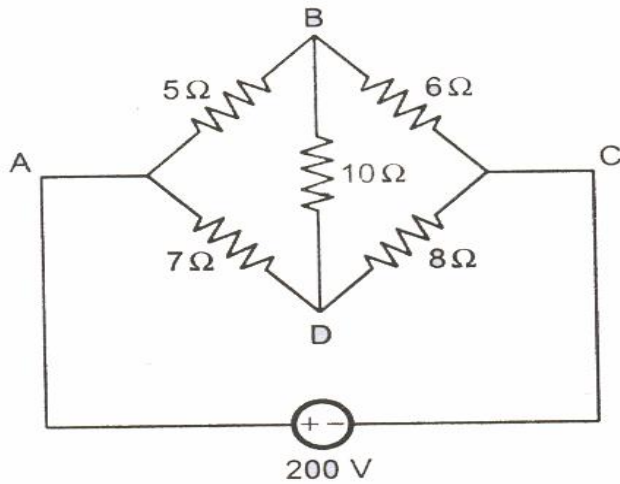
2. Find the current through branch AB by using superposition theorem.



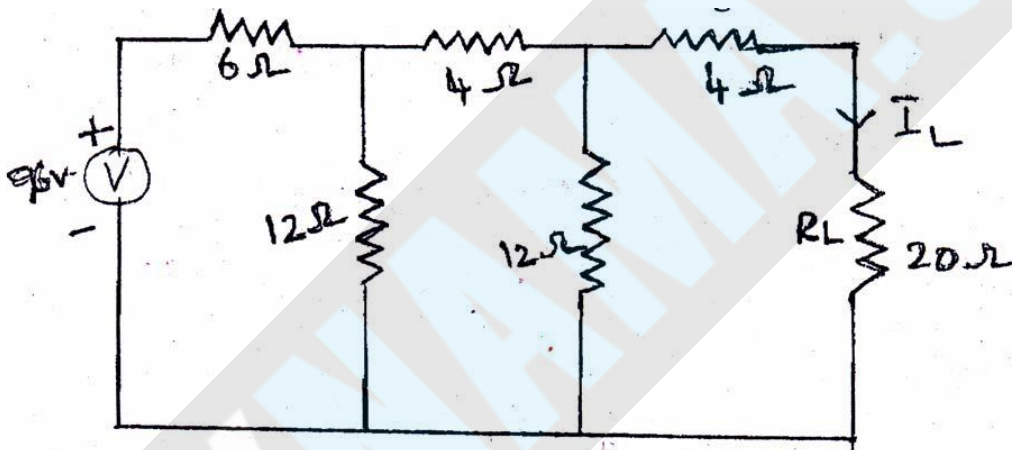
3. Find the current through 5 ohm resistance using Superposition theorem.



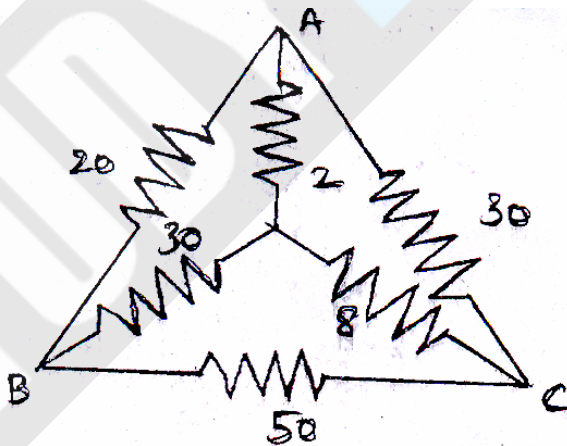
4. Find the current through 10 ohm resistance using Norton's theorem



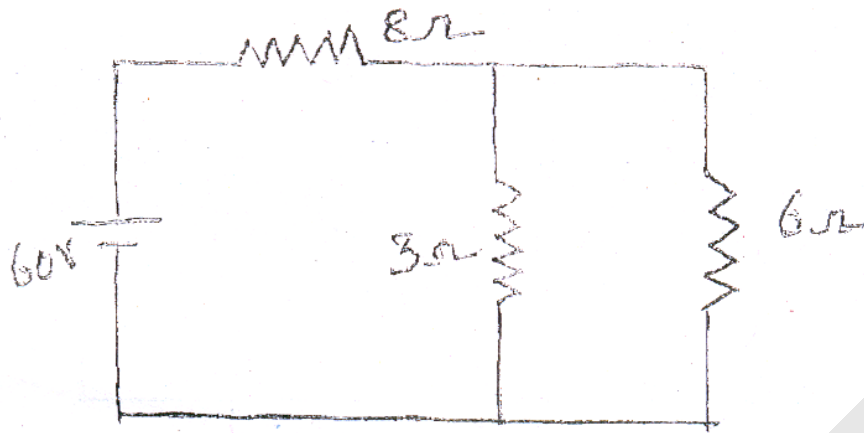
5. Find the Current (I) in 20Ω Resistance using Thevenin's theorem



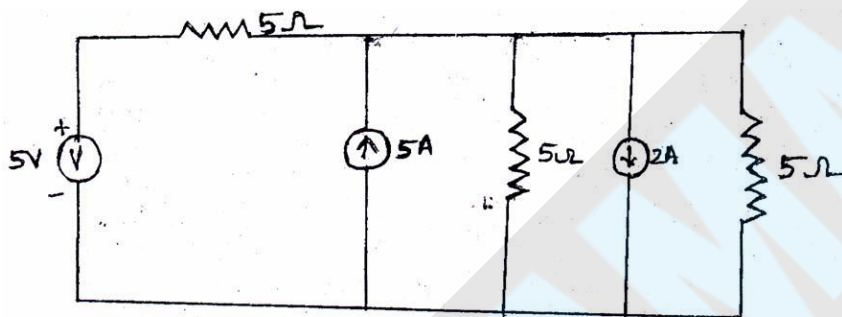
6. Find the resistance between A & B, A & C



7. Consider the following network as shown in figure. Determine the power observed by the 6Ω .



8. Find the total Current and total Resistance in the circuit given



UNIT II

ELECTRICAL MACHINES

DC Generator:

To change the Simple Generator into a direct-current generator, two things must be done: (1) The current must be conducted from the rotating loop of wire (2) The current must be made to move in only one direction. A device called a commutator performs both tasks.

DC generator construction:

What is Generator?

An electrical generator is a device that converts mechanical energy to electrical energy, generally using electromagnetic induction. The source of mechanical energy may be a reciprocating or turbine steam engine, water falling through a turbine or waterwheel, an internal combustion engine, a wind turbine, a hand crank, or any other source of mechanical energy.

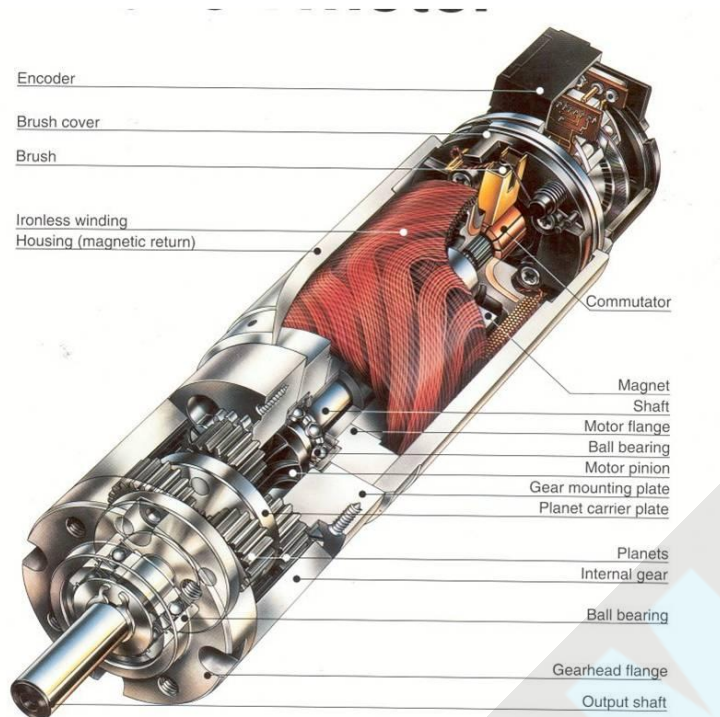
The Dynamo was the first electrical generator capable of delivering power for industry. The dynamo uses electromagnetic principles to convert mechanical rotation into an alternating electric current. A dynamo machine consists of a stationary structure which generates a strong magnetic field, and a set of rotating windings which turn within that field. On small machines the magnetic field may be provided by a permanent magnet; larger machines have the magnetic field created by electromagnets.

The energy conversion in generator is based on the principle of the production of dynamically induced e.m.f. whenever a conductor cuts magnetic flux, dynamically induced e.m.f. is produced in it according to Faraday's Laws of Electromagnetic induction. This e.m.f. causes a current to flow if the conductor circuit is closed. Hence, two basic essential parts of an electrical generator are (i) a magnetic field and (ii) a conductor or conductors which can so move as to cut the flux.

Here is the construction diagram of dc generator:

Generator Construction:

Simple loop generator is having a single-turn rectangular copper coil rotating about its own axis in a magnetic field provided by either permanent magnet or electro magnets. In case of without commutator the two ends of the coil are joined to slip rings which are insulated from each other and from the central shaft. Two collecting brushes (of carbon or copper) press against the slip rings. Their function is to collect the current induced in the coil. In this case the current waveform we obtain is alternating current (you can see in fig). In case of with commutator the slip rings are replaced by split rings. In this case the current is unidirectional.



Components of a generator:

Yoke: Yoke is a outer frame. It serves two purposes.

- (i) It provides mechanical support for the poles and acts as a protecting cover for the whole machine and
- (ii) It carries the magnetic flux produced by the poles.

In small generators where cheapness rather than weight is the main consideration, yokes are made of cast iron. But for large machines usually cast steel or rolled steel is employed. The modern process of forming the yoke consists of rolling a steel slab round a cylindrical mandrel and then welding it at the bottom. The feet and the terminal box etc., are welded to the frame afterwards. Such yokes possess sufficient mechanical strength and have high permeability.

Rotor: In its simplest form, the rotor consists of a single loop of wire made to rotate within a magnetic field. In practice, the rotor usually consists of several coils of wire wound on an armature.

Armature: The armature is a cylinder of laminated iron mounted on an axle. The axle is carried in bearings mounted in the external structure of the generator. Torque is applied to the axle to make the rotor spin.

Coil: Each coil usually consists of many turns of copper wire wound on the armature. The two ends of each coil are connected either to two slip rings (AC) or two opposite bars of a split-ring commutator (DC).

Stator: The stator is the fixed part of the generator that supplies the magnetic field in which the coils rotate. It may consist of two permanent magnets with opposite poles facing and shaped to fit around the rotor. Alternatively, the magnetic field may be provided by two electromagnets.

Field electromagnets: Each electromagnet consists of a coil of many turns of copper wire wound on a soft iron core. The electromagnets are wound, mounted and shaped in such a way that opposite poles face each other and wrap around the rotor.

Brushes: The brushes are carbon blocks that maintain contact with the ends of the coils via the slip rings (AC) or the split-ring commutator (DC), and conduct electric current from the coils to the external circuit.

Principle of operation:

DC generator converts mechanical energy into electrical energy. when a conductor move in a magnetic field in such a way conductors cuts across a magnetic flux of lines and emf produces in a generator and it is defined by faradays law of electromagnetic induction :emf causes current to flow if the conductor circuit is closed.

Applications of DC generator:

1. Shunt generators are extensively used for general light and power supply, and for charging of batteries, since, in conjunction with a field regulator, a constant terminal voltage can be maintained at all loads.
2. Series generators are mainly used as animation boosters in dc transmission system, in order to compensate for the drop of voltage due to the resistance of transmission conductors.
3. Over-compounded generators find use in dc transmission, since it is possible to keep on a constant voltage at the load end, by generating a larger voltage so as to overcome the line drop.

DC Motor:

A DC motor is a device which converts electrical energy into mechanical energy. D.C. motors are motors that run on Direct Current from a battery or D.C. power supply. Direct Current is the term used to describe electricity at a constant voltage. A.C. motors run on Alternating Current, which oscillates with a fixed cycle between a positive and negative value. Electrical outlets provide A.C. power.

In a brushed DC motor, the brushes make mechanical contact with a set of electrical contacts provided on a commutator secured to an armature, forming an electrical circuit between the DC electrical source and coil windings on the armature. As the armature rotates on an axis, the stationary brushes come into contact with different sections of the rotating commutator.

Permanent magnet DC motors utilize two or more brushes contacting a commutator which provides the direct current flow to the windings of the rotor, which in turn provide the desired magnetic repulsion/attraction with the permanent magnets located around the periphery of the motor.

The brushes are conventionally located in brush boxes and utilize a U-shaped spring which biases the brush into contact with the commutator. Permanent magnet brushless dc motors are widely used in a variety of applications due to their simplicity of design, high efficiency, and low noise. These motors operate by electronic commutation of stator windings rather than the conventional mechanical commutation accomplished by the pressing engagement of brushes against a rotating commutator.

A brushless DC motor basically consists of a shaft, a rotor assembly equipped with one or more permanent magnets arranged on the shaft, and a stator assembly which incorporates a stator component and phase windings. Rotating magnetic fields are formed by the currents applied to the coils.

The rotator is formed of at least one permanent magnet surrounded by the stator, wherein the rotator rotates within the stator. Two bearings are mounted at an axial distance to each other on the shaft to support the rotor assembly and stator assembly relative to each other. To achieve electronic commutation, brushless dc motor designs usually include an electronic controller for controlling the excitation of the stator windings.

How DC motors work?

There are different kinds of D.C. motors, but they all work on the same principles. When a permanent magnet is positioned around a loop of wire that is hooked up to a D.C. power source, we have the basics of a D.C. motor. In order to make the loop of wire spin, we have to connect a battery or DC power supply between its ends, and support it so it can spin about its axis. To allow the rotor to turn without twisting the wires, the ends of the wire loop are connected to a set of contacts called the commutator, which rubs against a set of conductors called the brushes. The brushes make electrical contact with the commutator as it spins, and are connected to the positive and negative leads of the power source, allowing electricity to flow through the loop. The electricity flowing through the loop creates a magnetic field that interacts with the magnetic field of the permanent magnet to make the loop spin.

Principles of Operation:

It is based on the principle that when a current-carrying conductor is placed in a magnetic field, it experiences a mechanical force whose direction is given by Fleming's Left-hand rule and whose magnitude is given by

Force, $F = B I l$ newton

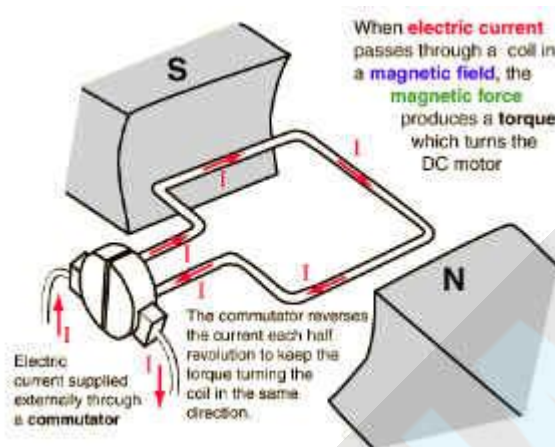
Where B is the magnetic field in weber/m².

I is the current in amperes and

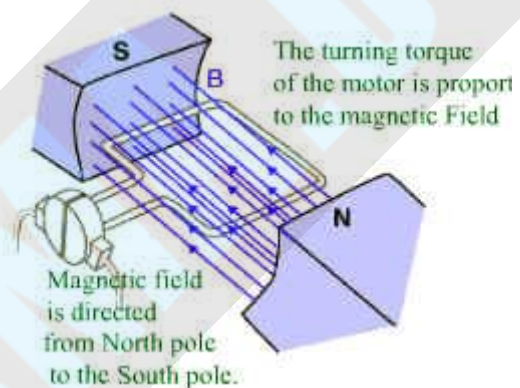
l is the length of the coil in meter.

The force, current and the magnetic field are all in different directions.

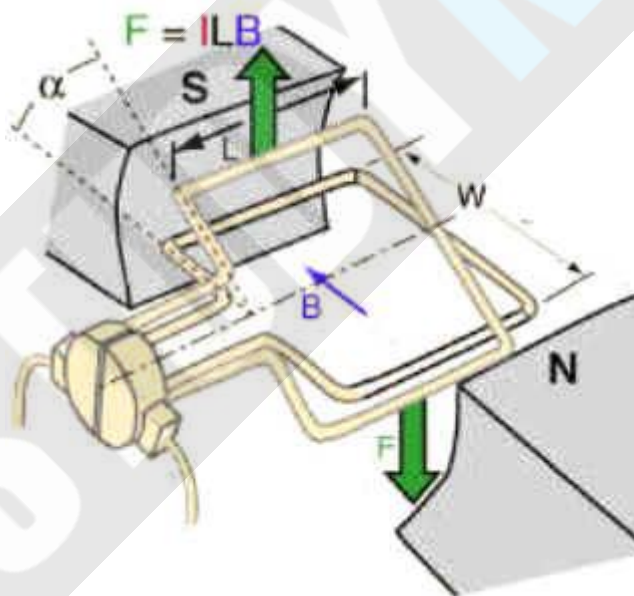
If an Electric current flows through two copper wires that are between the poles of a magnet, an upward force will move one wire up and a downward force will move the other wire down.



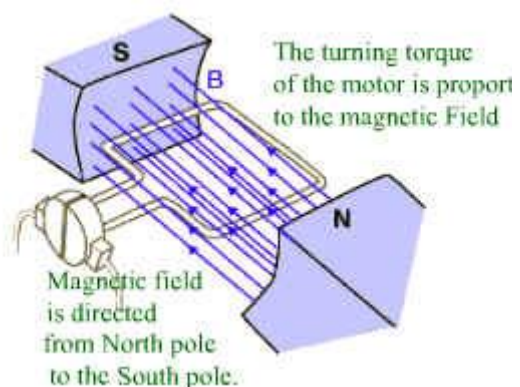
Force in DC Motor



Magnetic Field in DC Motor

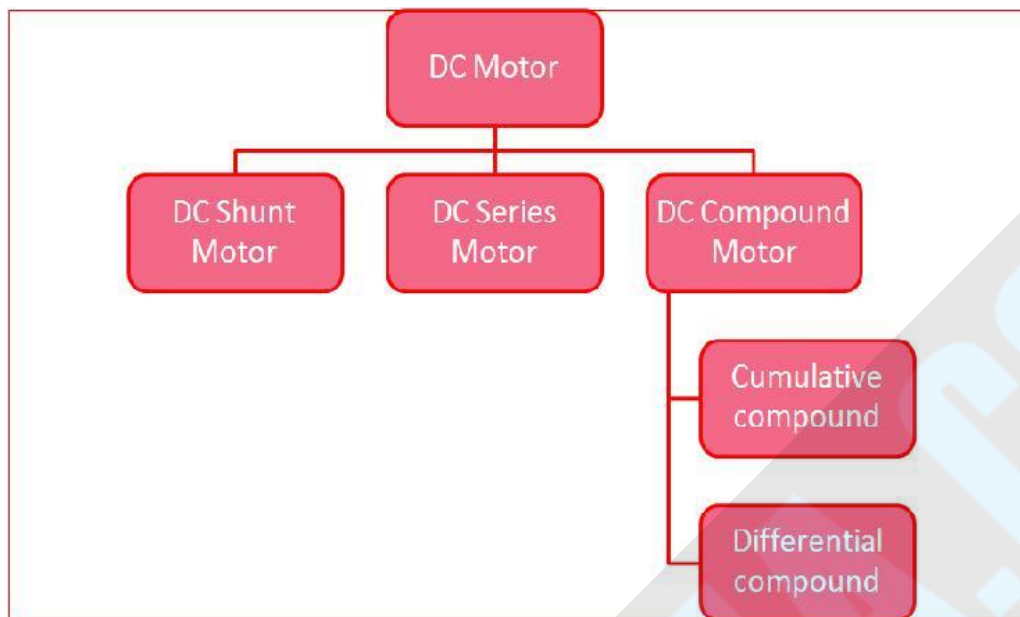


Torque in DC Motor



Current Flow in DC Motor

Classification of motor:



DC motors are more common than we may think. A car may have as many as 20 DC motors to drive fans, seats, and windows. They come in three different types, classified according to the electrical circuit used. In the shunt motor, the armature and field windings are connected in parallel, and so the currents through each are relatively independent. The current through the field winding can be controlled with a field rheostat (variable resistor), thus allowing a wide variation in the motor speed over a large range of load conditions. This type of motor is used for driving machine tools or fans, which require a wide range of speeds.

In the series motor, the field winding is connected in series with the armature winding, resulting in a very high starting torque since both the armature current and field strength run at their maximum. However, once the armature starts to rotate, the counter EMF reduces the current in the circuit, thus reducing the field strength. The series motor is used where a large starting torque is required, such as in automobile starter motors, cranes, and hoists.

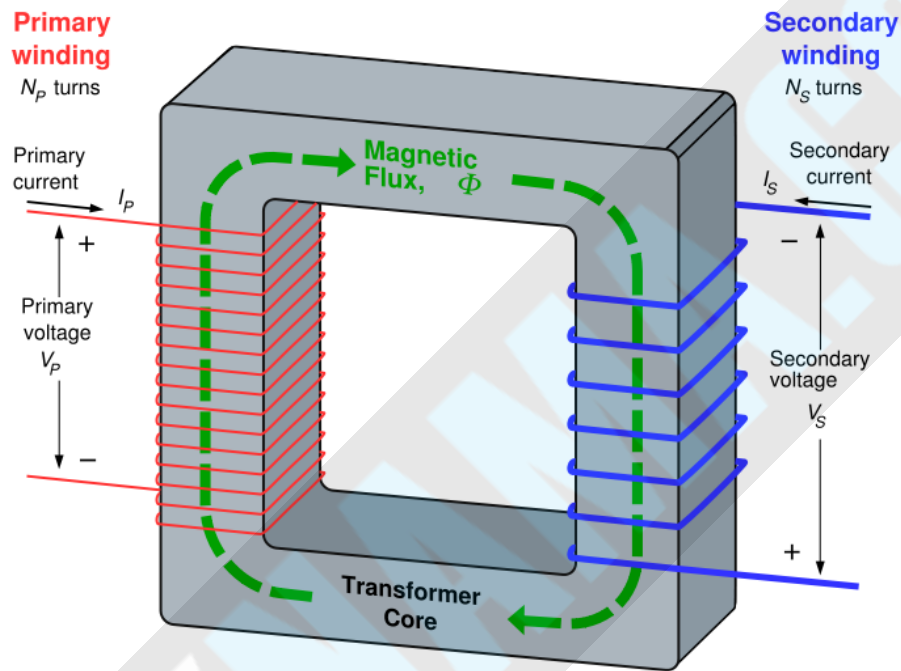
The compound motor is a combination of the series and shunt motors, having parallel and series field windings. This type of motor has a high starting torque and the ability to vary the speed and is used in situations requiring both these properties such as punch presses, conveyors and elevators.

DC motor advantages:

- Easy to understand design
- Easy to control speed
- Easy to control torque
- Simple, cheap drive design

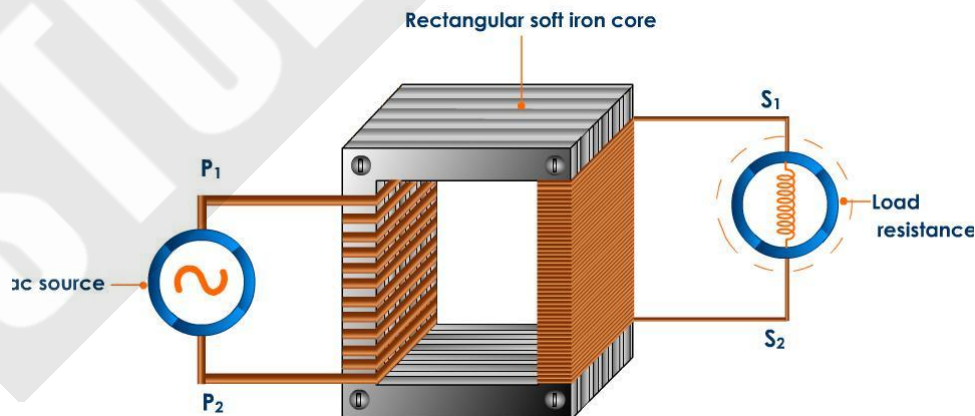
Transformer:

A transformer transfers electrical energy between two circuits. It usually consists of two wire coils wrapped around a core. These coils are called primary and secondary windings. Energy is transferred by mutual induction caused by a changing electromagnetic field. If the coils have different number of turns around the core, the voltage induced in the secondary coil will be different to the first.



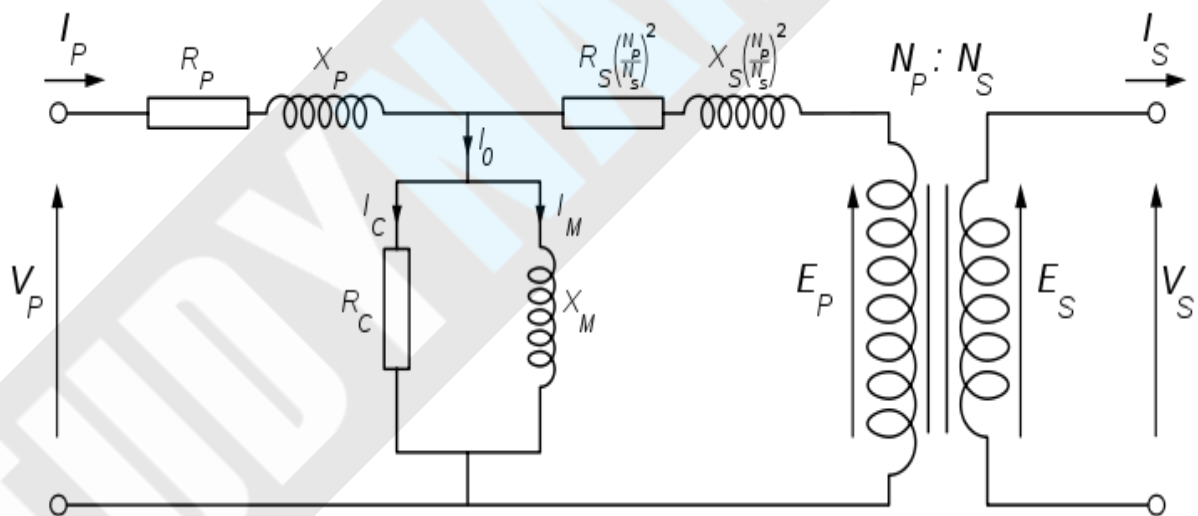
The device which is used to stepping up or stepping down of voltages is known as **transformer**.

For transformer working principle and its uses:



Equivalent circuit:

- The physical limitations of the practical transformer may be brought together as an equivalent circuit model built around an ideal lossless transformer. Power loss in the windings is current-dependent and is represented as in-series resistances R_P and R_S . Flux leakage results in a fraction of the applied voltage dropped without contributing to the mutual coupling, and thus can be modelled as reactance of each leakage inductance X_P and X_S in series with the perfectly coupled region.
- Iron losses are caused mostly by hysteresis and eddy current effects in the core, and are proportional to the square of the core flux for operation at a given frequency. Since the core flux is proportional to the applied voltage, the iron loss can be represented by a resistance R_C in parallel with the ideal transformer.
- A core with finite permeability requires a magnetizing current I_M to maintain the mutual flux in the core. The magnetizing current is in phase with the flux; saturation effects cause the relationship between the two to be non-linear, but for simplicity this effect tends to be ignored in most circuit equivalents. With a sinusoidal supply, the core flux lags the induced EMF by 90° and this effect can be modeled as a magnetizing reactance (reactance of an effective inductance) X_M in parallel with the core loss component. R_C and X_M are sometimes together termed the magnetizing branch of the model. If the secondary winding is made open-circuit, the current I_0 taken by the magnetizing branch represents the transformer's no-load current.
- The secondary impedance R_S and X_S is frequently moved to the primary side after multiplying the components by the impedance scaling factor $(N_P/N_S)^2$.



Transformer equivalent circuit, with secondary impedances referred to the primary side

- The resulting model is sometimes termed the "exact equivalent circuit", though it retains a number of approximations, such as an assumption of linearity. Analysis may be simplified by moving the magnetizing branch to the left of the primary impedance, an implicit assumption that the magnetizing current is low, and then summing primary and referred secondary impedances, resulting in so-called equivalent impedance.

Application of transformer:

Transformers are frequently used in power applications to interconnect systems operating at different voltage classes, for example 138 kV to 66 kV for transmission. Another application is in industry to adapt machinery built (for example) for 480 V supplies to operate on a 600 V supply. They are also often used for providing conversions between the two common domestic mains voltage bands in the world (100-130 and 200-250). The links between the UK 400kV and 275kV 'Super Grid' networks are normally three phase autotransformers with taps at the common neutral end.

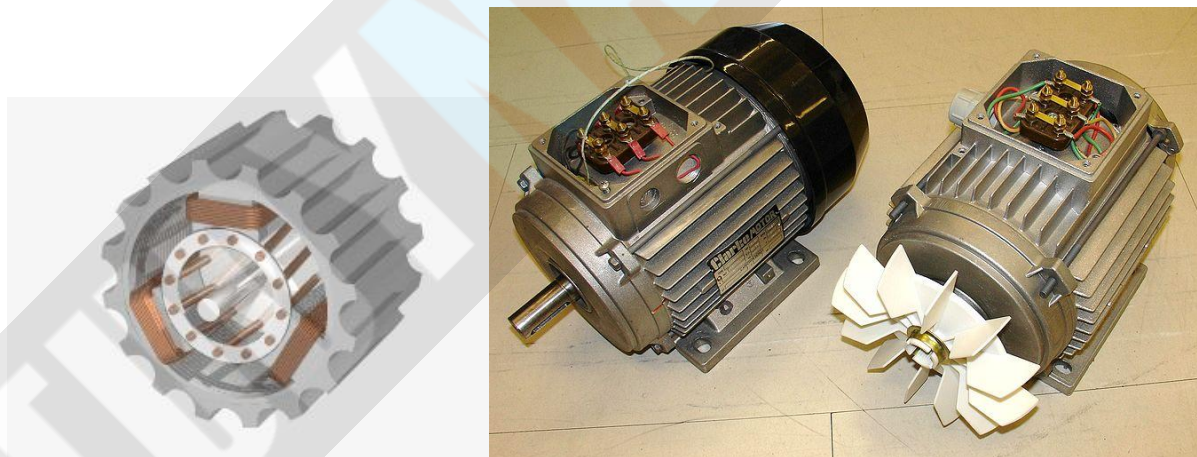
On long rural power distribution lines, special autotransformers with automatic tap-changing equipment are inserted as voltage regulators, so that customers at the far end of the line receive the same average voltage as those closer to the source. The variable ratio of the autotransformer compensates for the voltage drop Voltage drop along the line.

In audio applications, tapped autotransformers are used to adapt speakers to constant-voltage audio distribution systems, and for impedance matching such as between a low-impedance microphone and a high-impedance amplifier input.

Induction motors:

Definition:

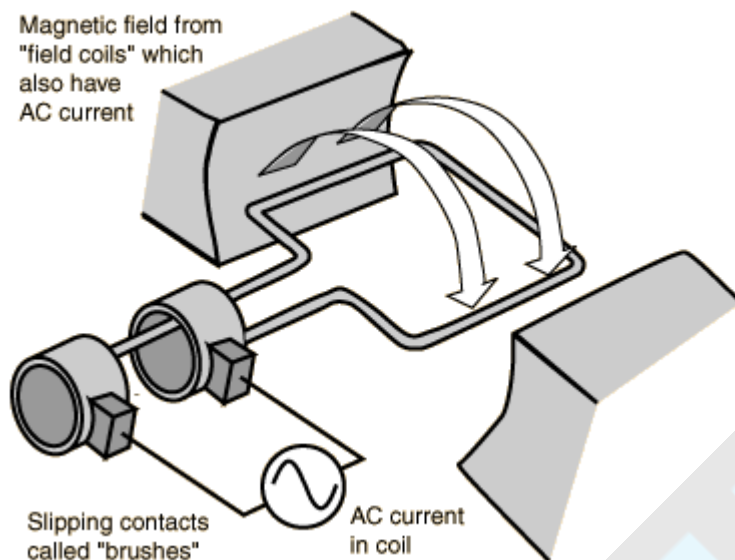
An induction motor (or asynchronous motor) is a type of alternating current motor where power is supplied to the rotor by means of electromagnetic induction.



An AC motor is an electric motor that is driven by an alternating current. It consists of two basic parts, an outside stationary stator having coils supplied with alternating current to produce a rotating magnetic field, and an inside rotor attached to the output shaft that is given a torque by the rotating field.

There are two types of AC motors, depending on the type of rotor used. The first is the synchronous motor, which rotates exactly at the supply frequency or a submultiple of the supply frequency. The magnetic field on the rotor is either generated by current delivered through slip rings or by a permanent magnet.

The second type is the induction motor, which turns slightly slower than the supply frequency. The magnetic field on the rotor of this motor is created by an induced current.



Induction ac motors are the simplest and most rugged electric motor and consists of two basic electrical assemblies: the wound stator and the rotor assembly. The induction ac motor derives its name from currents flowing in the secondary member (rotor) that are induced by alternating currents flowing in the primary member (stator). The combined electromagnetic effects of the stator and rotor currents produce the force to create rotation.

AC motors typically feature rotors, which consist of a laminated, cylindrical iron core with slots for receiving the conductors. The most common type of rotor has cast-aluminum conductors and short-circuiting end rings. This ac motor "squirrel cage" rotates when the moving magnetic field induces a current in the shorted conductors. The speed at which the ac motor magnetic field rotates is the synchronous speed of the ac motor and is determined by the number of poles in the stator and the frequency of the power supply: $n_s = 120f/p$, where n_s = synchronous speed, f = frequency, and p = the number of poles.

Synchronous speed is the absolute upper limit of ac motor speed. If the ac motor's rotor turns exactly as fast as the rotating magnetic field, then no lines of force are cut by the rotor conductors, and torque is zero. When ac motors are running, the rotor always rotates slower than the magnetic field. The ac motor's rotor speed is just slow enough to cause the proper amount of rotor current to flow, so that the resulting torque is sufficient to overcome windage and friction losses, and drive the load. The speed difference between the ac motor's rotor and magnetic field, called slip, is normally referred to as a percentage of synchronous speed: $s = 100 (n_s - n_a)/n_s$, where s = slip, n_s = synchronous speed, and n_a = actual speed.

Construction of 1 phase induction motor:

The induction motor essentially consists of two parts:

- Stator
- Rotor.

The supply is connected to the stator and the rotor received power by induction caused by the stator rotating flux, hence the motor obtains its name -induction motor. The stator consists of a cylindrical laminated & slotted core placed in a frame of rolled or cast steel. The frame provides mechanical protection and carries the terminal box and the end covers with bearings. In the slots of a 3-phase winding of insulated copper wire is distributed which can be wound for 2,4,6 etc. poles.

The rotor consists of a laminated and slotted core tightly pressed on the shaft. There are two general types of rotors:

- The squirrel-cage rotor,
- The wound (or slip ring) rotor.

In the squirrel-cage rotor, the rotor winding consists of single copper or aluminium bars placed in the slots and short-circuited by end-rings on both sides of the rotor.

In the wound rotor, an insulated 3-phase winding similar to the stator winding and for the same number of poles is placed in the rotor slots. The ends of the star-connected rotor winding are brought to three slip rings on the shaft so that connection can be made to it for starting or speed control.

Application of induction motor:

- Speed variation.
- Heavy load inertia starting.
- High starting torque requirements.
- Low starting current requirements.
- High efficiency at low speed.
- High power factor.

PART -A

1. What is op-amp?
2. Define slew rate.
3. Convert decimal 9 to binary
4. State Demorgan theorem
5. What is a Demultiplexer?

6. What is decoder?
7. How many FFS are required to construct a Decade counter?
8. What is race around condition?
9. What is a volatile memory?
10. What is meant by quantization?

Part B

11. State the characteristics of an ideal op. amp.
12. Draw the Logic diagram for the Boolean function $AB + C$.
13. State the Truth Table of a HALF Adder and FULL Adder.
14. What are the differences between Ring counter and Johnson counter?
15. Draw the circuit diagram of a 4 bit weighted – Resistor D/A converter.

UNIT III

SEMICONDUCTOR DEVICES AND APPLICATIONS

Passive Component:

A passive component is one that contributes no power gain to a circuit or system. It has no control action and does not require any input other than a signal to perform its function. The most commonly used passive circuit components in electronic and electrical applications are resistors e.g. Resistor, Capacitor, Inductor etc...

Resistors:

A resistor is a two-terminal electronic component designed to oppose an electric current by producing a voltage drop between its terminals in proportion to the current, that is, in accordance with Ohm's law:

$$V = IR$$

Physical material resist the flow of electrical current to some extent. Certain materials such as copper offer very low resistance to current flow, and hence they are called conductors. Other materials such as ceramic which offer extremely high resistance to current flow are called *insulators*. In electric and electronic circuits, there is a need for materials with specific values of resistance in the range between that of conductor and an insulator,. These materials are called resistors and their values of resistance expressed in ohms.

Resistors are used as part of electrical networks and electronic circuits. They are extremely commonplace in most electronic equipment. Practical resistors can be made of various compounds and films, as well as resistance wire (wire made of a high-resistivity alloy, such as nickel/chrome).

Resistor characteristics:

The primary characteristics of resistors are their resistance and the power they can dissipate. Other characteristics include temperature coefficient, noise, and inductance. Less well-known is critical resistance, the value below which power dissipation limits the maximum permitted current flow, and above which the limit is applied voltage. Critical resistance depends upon the materials constituting the resistor as well as its physical dimensions; it's determined by design.

Resistors can be integrated into hybrid and printed circuits, as well as integrated circuits. Size, and position of leads (or terminals) are relevant to equipment

designers; resistors must be physically large enough not to overheat when dissipating their power.

Inductors:

An inductor is an electronic device which consists of a coil of wire which may have a metallic or ferrite core. If it appears to have no core, it is considered to have an air core. The core material will greatly affect the value of the inductor. The unit of measure is the 'henry', but since that is such a large value of inductance, the value is usually stated in millihenries. One henry is equal to one thousand millihenries. If everything else stays constant, increasing the number of turns of wire around the core, will increase the value of the inductor.



This is the schematic symbol for an inductor.

Inductor symbol:



An "ideal inductor" has inductance, but no resistance or capacitance, and does not dissipate energy. A real inductor is equivalent to a combination of inductance, some resistance due to the resistivity of the wire, and some capacitance. At some frequency, usually much higher than the working frequency, a real inductor behaves as a resonant circuit (due to its self capacitance). In addition to dissipating energy in the resistance of the wire, magnetic core inductors may dissipate energy in the core due to hysteresis, and at high currents may show other departures from ideal behavior due to nonlinearity.

Application of Inductor:

Inductors are used as surge protectors because they block strong current changes.

They are used as telephone line filters, to remove high frequency broadband signals and are placed on the ends of the cables to reduce signal noise.

Inductors and capacitors are used together in audio circuits to filter or amplify specific frequencies.

Chokes are small inductors that block alternating current and are used to reduce electrical and radio interference. A basic transformer is just two inductors wound around a large steel core. Their magnetic fields are coupled, because the core forces them to flow through both coils.

When an alternating current flows in one coil, it induces an alternating current in the other coil.

How Inductors Works:

The inductor has as a coil of copper conductors wound around a central core. When current is passed through the coil a magnetic flux is created around the coil due to the properties of electromotive force. The resistance increases when a core is placed in the coil and this increases the inductance by hundreds of times. The core can be made of different materials but cores made of ferrite produce the maximum inductance. The current to voltage lag is 90° but with the use of resistive substance a resistive and inductive circuit is formed, the phase angle lag becomes smaller and is based on the frequency that is constant.

Inductance is the circuit's resistance to change in current. Inductance tolerance is the amount of variation that is permitted within the nominal value. The frequency for which the distributed capacitance starts resonating with the inductance and canceling the capacitance is called the self resonant frequency or SRF. At SRF, the inductor works as a high impedance, resistive element. Quality factor (Q value) is the measure of relative losses of the inductor and is expressed as capacitive resistance divided by the equivalent serial resistance.

Inductors are used

1. In tuning circuits
2. In filter
3. In timing circuits
4. In oscillator tank circuits

Capacitors:

A capacitor or condenser is a passive electronic component consisting of a pair of conductors separated by a dielectric. When a voltage potential difference exists between the conductors, an electric field is present in the dielectric. This field stores energy and produces a mechanical force between the plates. The effect is greatest between wide, flat, parallel, narrowly separated conductors.

An ideal capacitor is characterized by a single constant value, capacitance, which is measured in farads. This is the ratio of the electric charge on each conductor to the potential difference between them. In practice, the dielectric between the plates passes a small amount of leakage current. The conductors and leads introduce an equivalent series resistance and the dielectric has an electric field strength limit resulting in a breakdown voltage.

The properties of capacitors in a circuit may determine the resonant frequency and quality factor of a resonant circuit, power dissipation and operating frequency in a digital logic circuit, energy capacity in a high-power system, and many other important aspects

Active components:

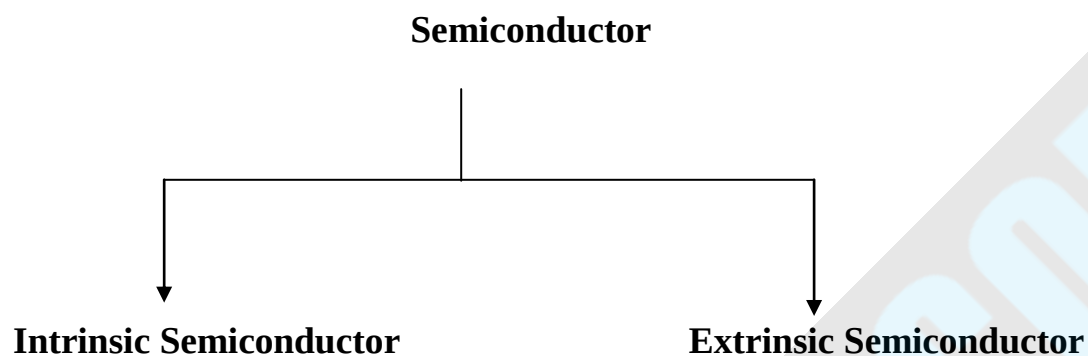
An electronic component is any physical entity in an electronic system whose intention is to affect the electrons or their associated fields in a desired manner consistent with the intended function of the electronic system. Components are generally intended to be in mutual electromechanical contact, usually by being soldered to a printed circuit board (PCB), to create an electronic circuit with a particular function (for example an amplifier, radio receiver, or oscillator). Components may be packaged singly or in more complex groups as integrated circuits. Some common electronic components are capacitors, resistors, diodes, transistors, etc.

Semiconductor

A semiconductor is a solid material that has electrical conductivity between those of a conductor and an insulator.

- A material with electrical conductivity due to electron flow intermediate in magnitude between that of a conductor and an insulator.
- This means a conductivity roughly in the range of 10^3 to 10^{-8} siemens per centimeter.
- Silicon is the most widely used semiconductor material.
- The number of electrons in the valence orbit is the key to conductivity.
- Conductors have one valence electron, semiconductors have four valence electrons, and insulators have eight valence electrons.

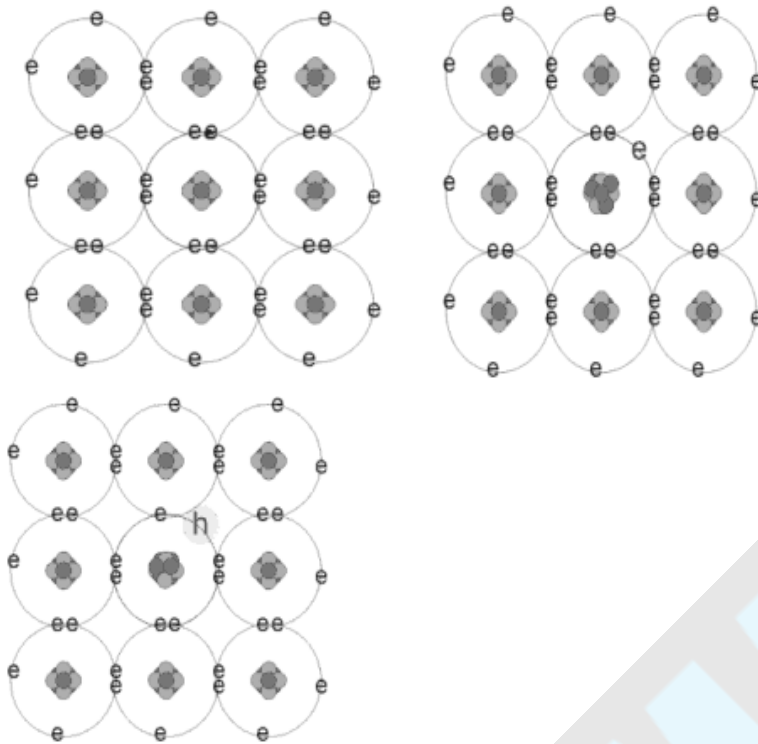
Classification:



Intrinsic Semiconductor:

- An intrinsic semiconductor also called an undoped semiconductor or i-type semiconductor.
- It is a pure semiconductor without any significant dopant species present.
- The number of charge carriers determined by the properties of the material itself instead of the amount of impurities.
- In intrinsic semiconductors the number of excited electrons and the number of holes are equal: $n = p$.

$$\begin{array}{llll} \text{Free electron concentration} & = & \text{hole concentration} & = \text{intrinsic electron} \\ & & & \text{Concentration} \\ n \text{ (electron / cm}^3\text{)} & = & p \text{ (hole / cm}^3\text{)} & = n_i \end{array}$$



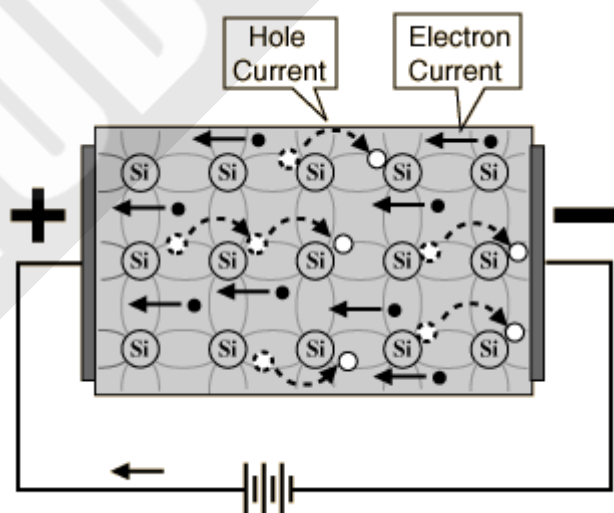
Intrinsic Semiconductor

n-type

p-type

Conductivity of Intrinsic semiconductor:

- The electrical conductivity of intrinsic semiconductors can be due to crystal defects or to thermal excitation.
- Both electrons and holes contribute to current flow in an intrinsic semiconductor.



- The current which will flow in an intrinsic semiconductor consists of both electron and hole current.
- That is, the electrons which have been freed from their lattice positions into the conduction band can move through the material.
- In addition, other electrons can hop between lattice positions to fill the vacancies left by the freed electrons.
- This additional mechanism is called hole conduction because it is as if the holes are migrating across the material in the direction opposite to the free electron movement.
- The current flow in an intrinsic semiconductor is influenced by the density of energy states which in turn influences the electron density in the conduction band.
- This current is highly temperature dependent.

Thermal excitation:

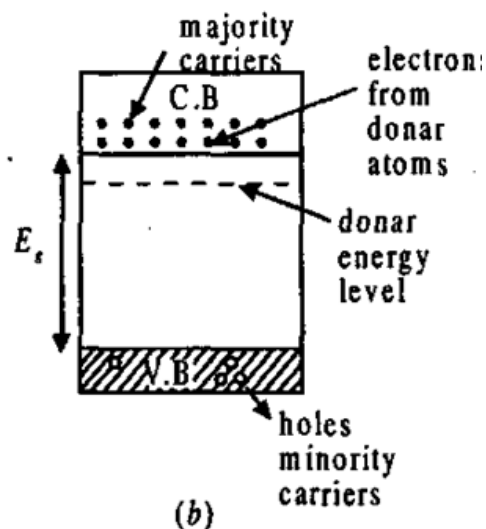
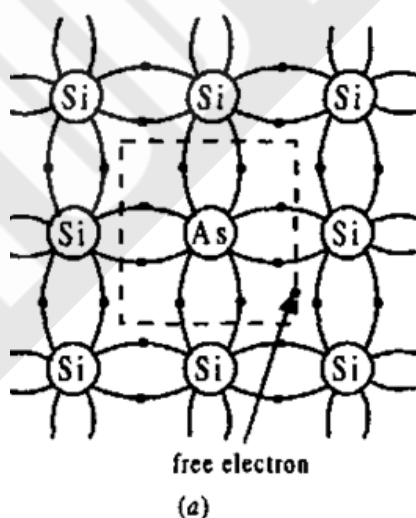
- In an intrinsic semiconductor like silicon at temperatures above absolute zero, there will be some electrons which are excited across the band gap into the conduction band and which can produce current.
- When the electron in pure silicon crosses the gap, it leaves behind an electron vacancy or "hole" in the regular silicon lattice.
- Under the influence of an external voltage, both the electron and the hole can move across the material.
- In n-type semiconductor:

- The dopant contributes extra electrons, dramatically increasing the conductivity.
- In p-type semiconductor,
- The dopant produces extra vacancies or holes, which likewise increase the conductivity.

Extrinsic Semiconductor

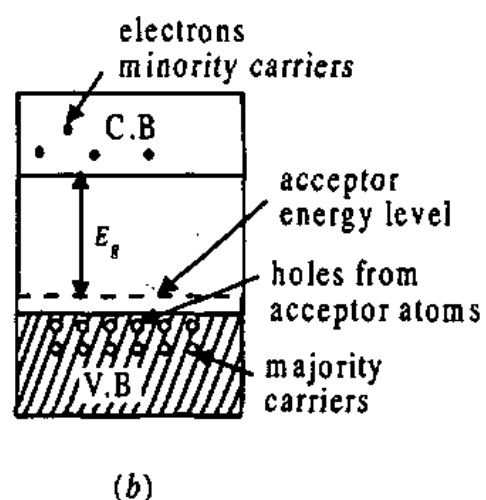
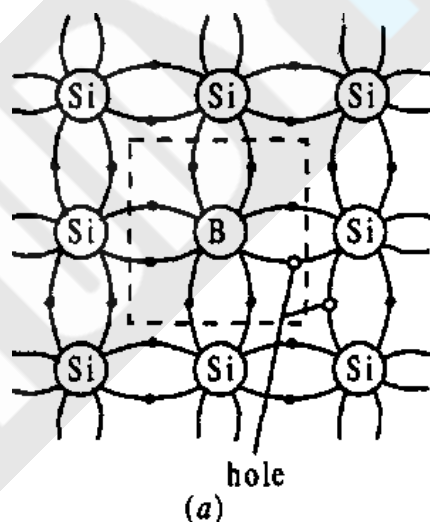
- The electrical conductivity of a pure semiconductor is very small.
- To increase the conductivity, impurities are added.
- The impurity added semiconductor is called extrinsic semiconductor.
- The process of adding impurity is called doping.
- The added impurity is called dopant.
- Usually one or two atoms of impurity is added per 10^6 atoms of a semiconductor.
- There are two types (i) p-type and (ii) n-type semiconductors.

(i) n-type semiconductor:



- When an impurity, from V group elements like arsenic (As), antimony having 5 valence electrons is added to Ge (or Si), the impurity atom donates one electron to Ge (or Si).
- The 4 electrons of the impurity atom is engaged in covalent bonding with Si atom.
- The fifth electron is free. This increases the conductivity.
- The impurities are called donors.
- The impurity added semiconductor is called n-type semiconductor, because their increased conductivity is due to the presence of the negatively charged electrons, which are called the majority carriers.
- The energy band of the electrons donated by the impurity atoms is just below the conduction band.
- The electrons absorb thermal energy and occupy the conduction band.
- Due to the breaking of covalent bond, there will be a few holes in the valence band at this temperature.
- These holes in n-type are called minority carriers.

(ii) p-type semiconductor:

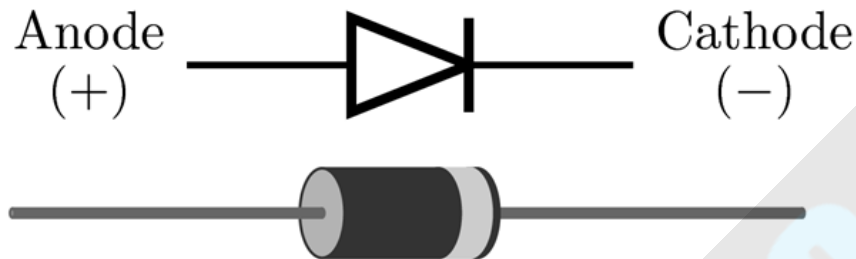


- If a III group element, like indium (In), boron (B), aluminium (Al) etc., having three valence electrons, is added to a semiconductor say Si, the three electrons form covalent bond.
- There is a deficiency of one electron to complete the 4th covalent bond and is called a hole.
- The presence of the hole increases the conductivity because these holes move to the nearby atom, at the same time the electrons move in the opposite direction.
- The impurities added semiconductor is called p-type semiconductor.
- The impurities are called acceptors as they accept electrons from the semiconductor
- Holes are the majority carriers and the electrons produced by the breaking of bonds are the minority carriers.

PN Junction Diode

- ✦ A p-n junction is formed by joining P-type and N-type semiconductors together in very close contact.
- ✦ The term junction refers to the boundary interface where the two regions of the semiconductor meet.
- ✦ Diode is a two-terminal electronic component that conducts electric current in only one direction.

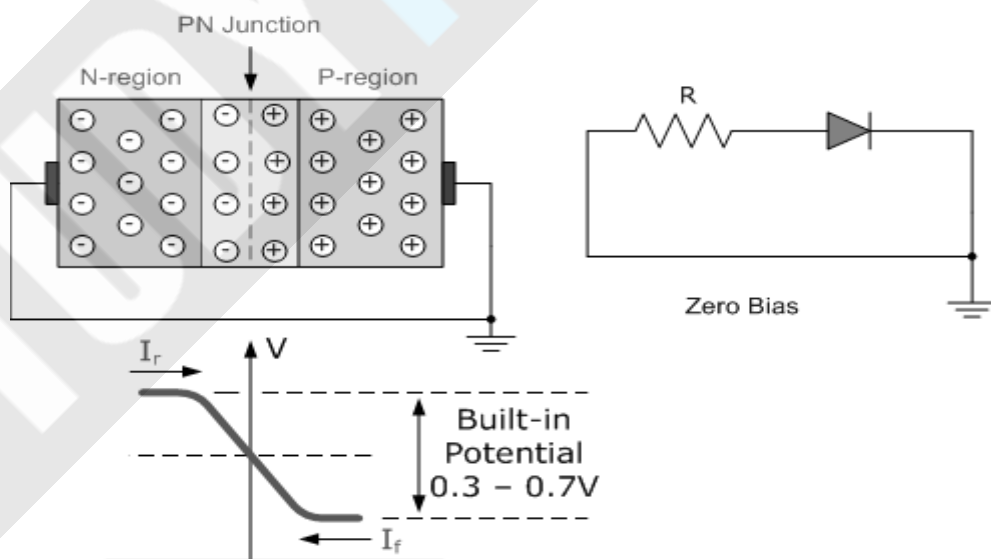
- ✦ The crystal conducts conventional current in a direction from the p-type side (called the anode) to the n-type side (called the cathode), but not in the opposite direction.



Biasing:

There are two operating regions and three possible "biasing" conditions for the standard Junction Diode and these are:

- ✦ Zero Bias - No external voltage potential is applied to the PN-junction.
 - When a diode is **Zero Biased** no external energy source is applied and a natural **Potential Barrier** is developed across a depletion layer.

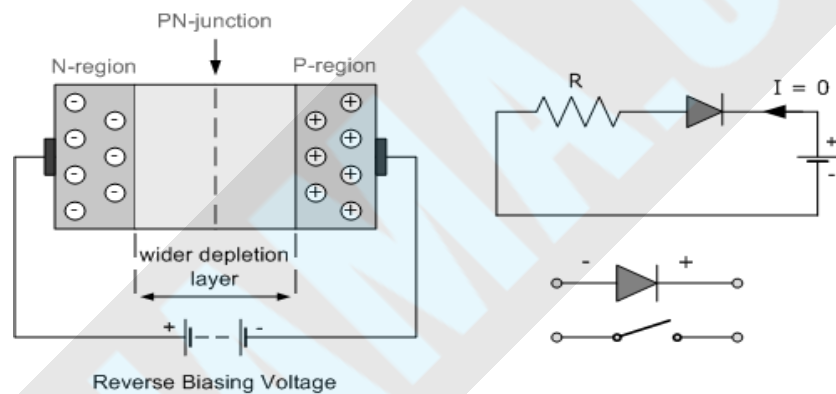


- ✦ Reverse Bias - The voltage potential is connected negative, (-ve) to the P-

type material and positive, (+ve) to the N-type material across the diode which has the effect of Increasing the PN-

junction width.

-When a junction diode is **Forward Biased** the thickness of the depletion region reduces and the diode acts like a short circuit allowing full current to flow.

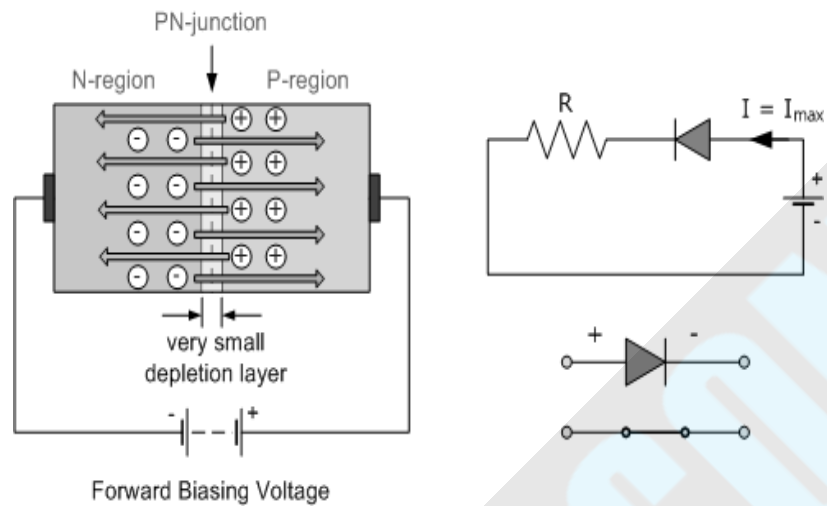


⊕ Forward Bias - The voltage potential is connected positive, (+ve) to the P-t

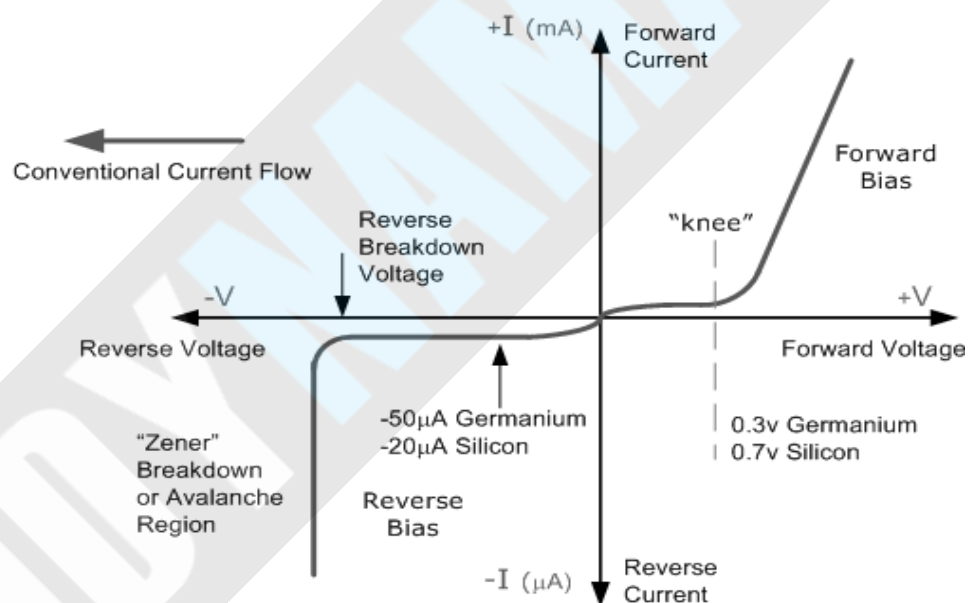
ype material and negative, (-ve) to the N-type material across the diode which has the effect of Decreasing the PN-

junction width.

- When a junction diode is **Reverse Biased** the thickness of the depletion region increases and the diode acts like an open circuit blocking any current flow, (only a very small leakage current).



V-I Characteristics:



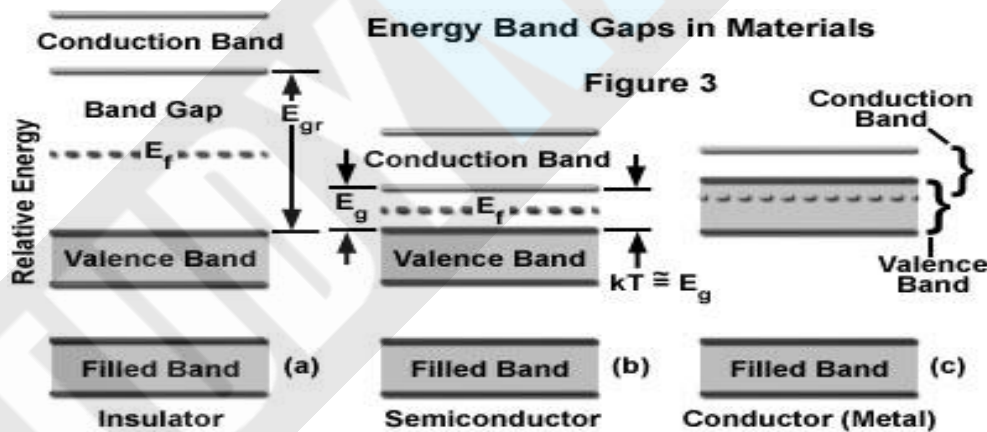
Forward Bias:

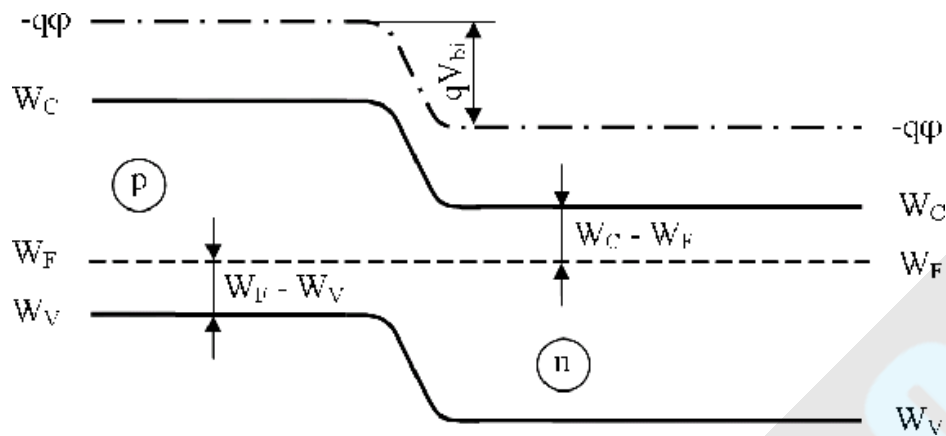
- ⊕ The application of a **forward biasing** voltage on the junction diode results in the depletion layer becoming very thin and narrow which represents a low impedance path through the junction thereby allowing high currents to flow.
- ⊕ The point at which this sudden increase in current takes place is represented on the static I-V characteristics curve above as the "knee" point.

Reverse Bias:

- ✦ In Reverse biasing voltage a high resistance value to the PN junction and practically zero current flows through the junction diode with an increase in bias voltage.
- ✦ However, a very small **leakage current** does flow through the junction which can be measured in microamperes, (μA).
- ✦ One final point, if the reverse bias voltage V_r applied to the diode is increased to a sufficiently high enough value, it will cause the PN junction to overheat and fail due to the avalanche effect around the junction.
- ✦ This may cause the diode to become shorted and will result in the flow of maximum circuit current, and this shown as a step downward slope in the reverse static characteristics curve below.

Energy Band Structure:





Energy Band Diagram of a PN-Junction

Diffusion capacitance and Space Charge capacitance

Diffusion capacitance

- ❖ As a p-n diode is forward biased, the minority carrier distribution in the quasi-neutral region increases dramatically.
- ❖ In addition, to preserve quasi-neutrality, the majority carrier density increases by the same amount.
- ❖ This effect leads to an additional capacitance called the diffusion capacitance.
- ❖ The diffusion capacitance is calculated from the change in charge with voltage:

$$C = \frac{d\Delta Q}{dV_a}$$

- ❖ Where the charge, ΔQ , is due to the excess carriers.
- ❖ Unlike a parallel plate capacitor, the positive and negative charge is not spatially separated. Instead, the electrons and holes are

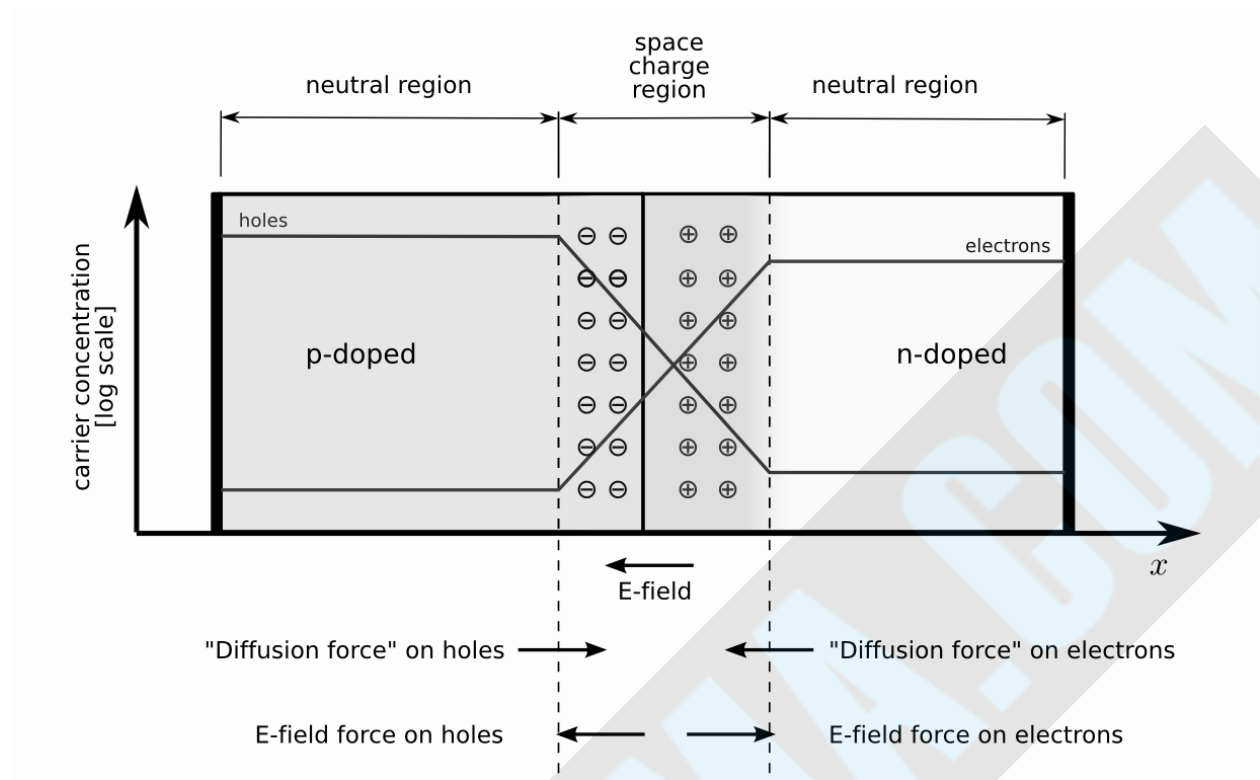
separated by the energy bandgap.

- ❖ Nevertheless, these voltage dependent charges yield a capacitance just as the one associated with a parallel plate capacitor.
- ❖ The excess minority-carrier charge is obtained by integrating the charge density over the quasi-neutral region:

$$\Delta Q_p = \int_{x_n}^{w_n} qA(p_n - p_{n0}) dx$$

Space Charge capacitance:

- ❖ After joining p-type and n-type semiconductors, electrons near the p–n interface tend to diffuse into the p region.
- ❖ As electrons diffuse, they leave positively charged ions (donors) in the n region.
- ❖ Similarly, holes near the p–n interface begin to diffuse into the n-type region leaving fixed ions (acceptors) with negative charge.
- ❖ The regions nearby the p–n interfaces lose their neutrality and become charged, forming the space charge capacitance.

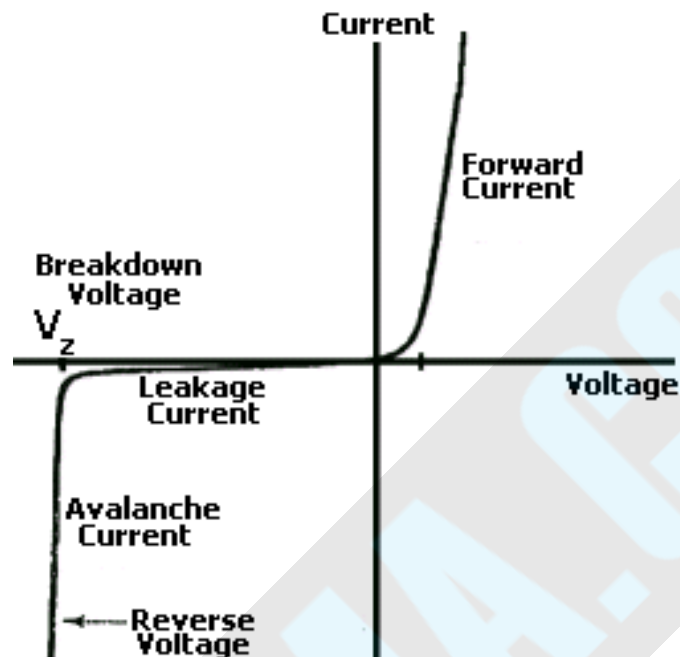


Zener Diode

A Zener diode is a special kind of diode which permits current to flow in the forward direction as normal, but will also allow it to flow in the reverse direction when the voltage is above a certain value-the breakdown voltage known as the Zener voltage.

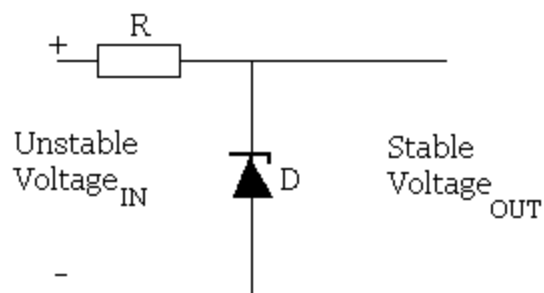


V-I Characteristics of the Zener Diode:



- The illustration above shows this phenomenon in a current vs voltage graph with a zener diode connected in the forward direction .It behaves exactly as a standard diode.
- In the reverse direction however there is a very small leakage current between 0v and the zener voltage –i.e. just a tiny amount of current is able to flow.
- Then, when the voltage reaches the breakdown voltage(v_z),suddenly current can flow freely through it.

Zenor Diode voltage regulator circuit



- Since the voltage dropped across a Zener diode is a known and fixed value, Zener diodes are typically used to regulate the voltage in electric circuits.
- Using a resistor to ensure that the current passing through the Zener diode is at least 5mA.
- The voltage drop across the diode is exactly equal to the Zener voltage of the diode.

BJT

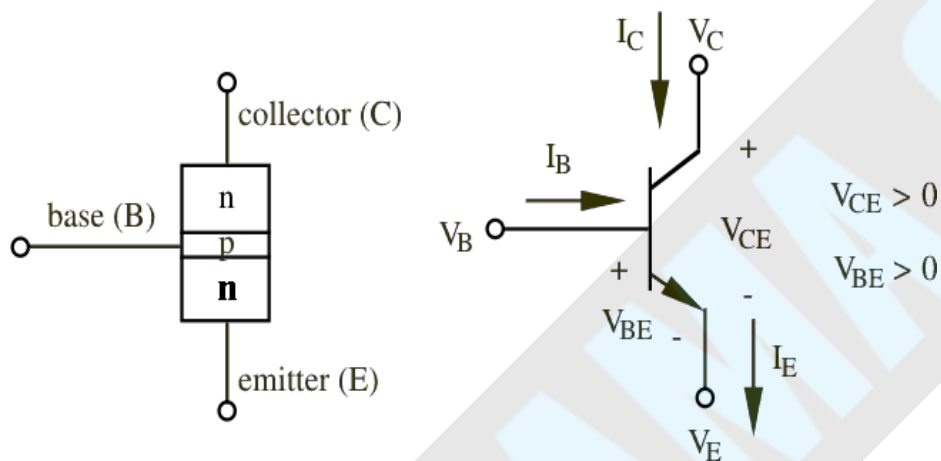
- ◆ A bipolar junction transistor (BJT) is a three-terminal electronic device constructed of doped semiconductor material.
- ◆ It may be used in amplifying or switching applications.
- ◆ Bipolar transistors are so named because their operation involves both electrons and holes.
- ◆ Charge flow in a BJT is due to bidirectional diffusion of charge carriers across a junction between two regions of different charge concentrations.
- ◆ 3 adjacent regions of doped Si (each connected to a lead):
 - Base. (thin layer, less doped).
 - Collector.

- Emitter.

◆ 2 types of BJT:

- npn.
- pnp.

npn bipolar junction transistor:



npn bipolar junction transistor

- 1 thin layer of p-type, sandwiched between 2 layers of n-type.
- N-type of emitter: more heavily doped than collector.
- With $V_C > V_B > V_E$:
 - Base-Emitter junction forward biased, Base-Collector reverse biased.
 - Electrons diffuse from Emitter to Base (from n to p).
 - There's a depletion layer on the Base-Collector junction → no flow of e^- allowed.
 - BUT the Base is thin and Emitter region is n^+ (heavily doped) → electrons have enough momentum to cross the Base into the Collector.
 - The small base current I_B controls a large current I_C

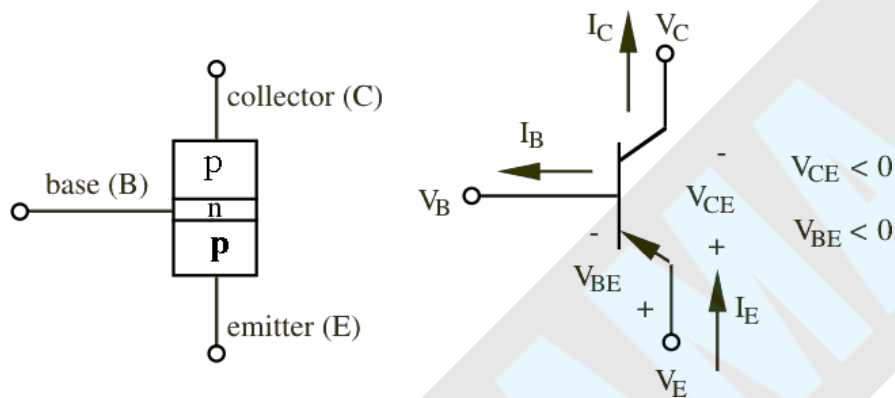
$$I_E = I_C + I_B$$

$$V_{BE} = V_B - V_E$$

$$V_{CE} = V_C - V_E$$

$$I_C = \beta I_B$$

pn bipolar junction transistor:



pn bipolar junction transistor

- 1 thin layer of n-type, sandwiched between 2 layers of p-type.
- P-type of emitter: more heavily doped than collector.
- The voltages v_{EB} and v_{CB} are positive when they forward bias their respective pn junctions.
- Collector current and base current exit the transistor terminals and emitter current enters the device.
- Base current is given by:

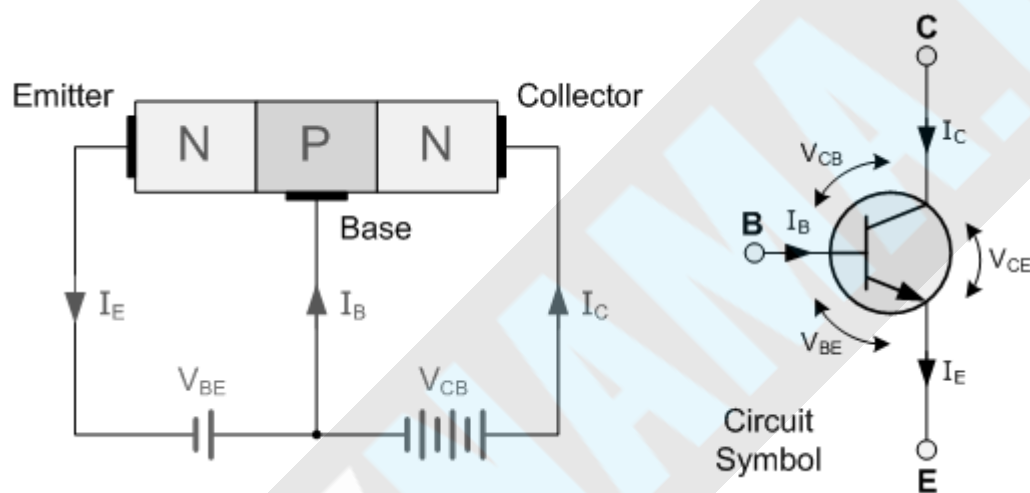
$$i_B = \frac{i_F}{\beta_F} = \frac{I_S}{\beta_F} \left[\exp\left(\frac{v_{EB}}{V_T}\right) - 1 \right]$$

- Emitter current is given by:

$$i_E = i_C + i_B = I_S \left[1 + \frac{1}{\beta_F} \right] \left[\exp\left(\frac{v_{EB}}{V_T}\right) - 1 \right]$$

Operation of NPN and PNP Transistor

Operation of NPN Transistor:



Forward-biased junction:

- ✦ With the emitter-to-base junction in the figure biased in the forward direction, electrons leave the negative terminal of the battery and enter the n - material (emitter).
- ✦ Since electrons are majority current carriers in the n - material, they pass easily through the emitter, cross over the junction, and combine with holes in the p material (base).

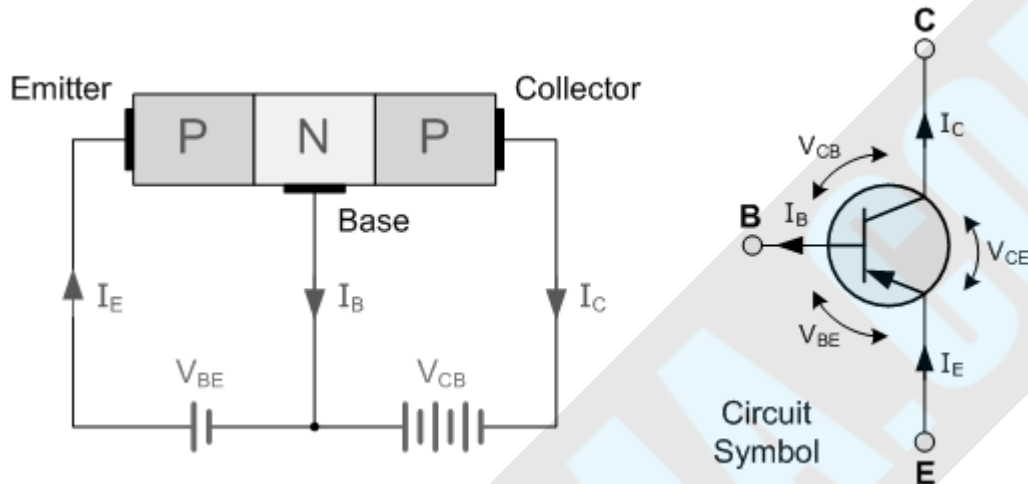
- ◆ For each electron that fills a hole in the p material, another electron will leave the p material (creating a new hole) and enter the positive terminal of the battery.

Reverse-biased junction:

- ◆ The second pn - junction (base-to-collector), or reverse-biased junction as it is called blocks the majority current carriers from crossing the junction.
- ◆ However, there is a very small current, that does pass through this junction.
- ◆ This current is called minority current, or reverse current. This current was produced by the electron-hole pairs.
- ◆ The minority carriers for the reverse-biased pn - junction are the electrons in the p material and the holes in the n - material.
- ◆ These minority carriers actually conduct the current for the reverse-biased junction when electrons from the p material enter the n material, and the holes from the n - material enter the p - material.
- ◆ However, the minority current electrons play the most important part in the operation of the nnp transistor.
- ◆ Total current flow in the nnp transistor is through the emitter lead. Therefore, in terms of percentage, I_E is 100 percent.
- ◆ On the other hand, since the base is very thin and lightly doped, a smaller percentage of the total current (emitter current) will flow in the base circuit than in the collector circuit.
- ◆ Usually no more than 2 to 5 percent of the total current is base current (I_B) while the remaining 95 to 98 percent is collector current (I_C).
- ◆ A very basic relationship exists between these two currents:

$$I_E = I_B + I_C$$

Operation of PNP Transistor:



Forward Biased Junction:

- ◆ With the bias setup shown, the positive terminal of the battery repels the emitter holes toward the base, while the negative terminal drives the base electrons toward the emitter.
- ◆ When an emitter hole and a base electron meet, they combine.
- ◆ For each electron that combines with a hole, another electron leaves the negative terminal of the battery, and enters the base.
- ◆ At the same time, an electron leaves the emitter, creating a new hole, and enters the positive terminal of the battery.
- ◆ This movement of electrons into the base and out of the emitter constitutes base current flow (I_B), and the path these electrons take is referred to as the emitter-base circuit.

Reverse Biased Junction:

- ◆ In the reverse-biased junction the negative voltage on the collector and the positive voltage on the base block the majority current carriers from crossing the junction.
- ◆ However, this same negative collector voltage acts as forward bias for the minority current holes in the base, which cross the junction and enter the collector.
- ◆ The minority current electrons in the collector also sense forward bias-the positive base voltage-and move into the base.
- ◆ The holes in the collector are filled by electrons that flow from the negative terminal of the battery.
- ◆ At the same time the electrons leave the negative terminal of the battery, other electrons in the base break their covalent bonds and enter the positive terminal of the battery.
- ◆ Although there is only minority current flow in the reverse-biased junction, it is still very small because of the limited number of minority current carriers.

Transistor Configuration

CE, CB, CC Configurations:

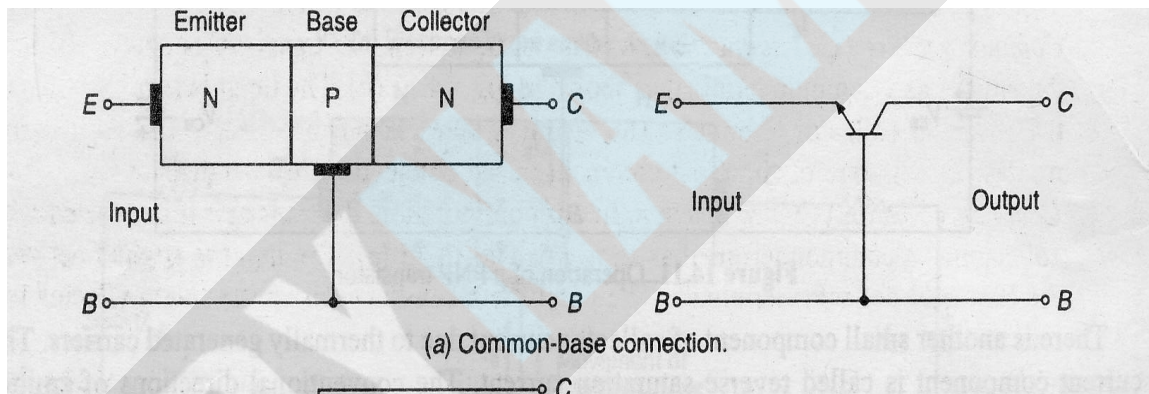
- ◆ We know that transistor has three terminals namely emitter(E), base(B), collector(C).
- ◆ However, when a transistor is connected in a circuit, we require four terminals (ie) two terminals for input and two terminals for output.
- ◆ This difficulty is overcome by using one of the terminals as common terminal.

- ◆ Depending upon the terminals which are used as a common terminal to the input and output terminals, the transistors can be connected in the following three different configuration.

1. Common base configuration
2. Common emitter configuration
3. Common collector configuration

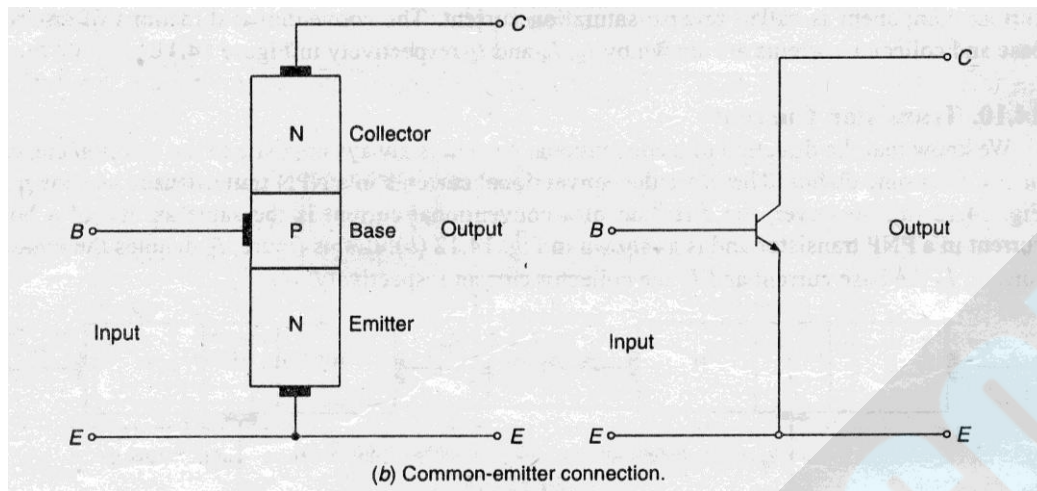
Common base configuration:

- ◆ In this configuration base terminal is connected as a common terminal.
- ◆ The input is applied between the emitter and base terminals. The output is taken between the collector and base terminals.



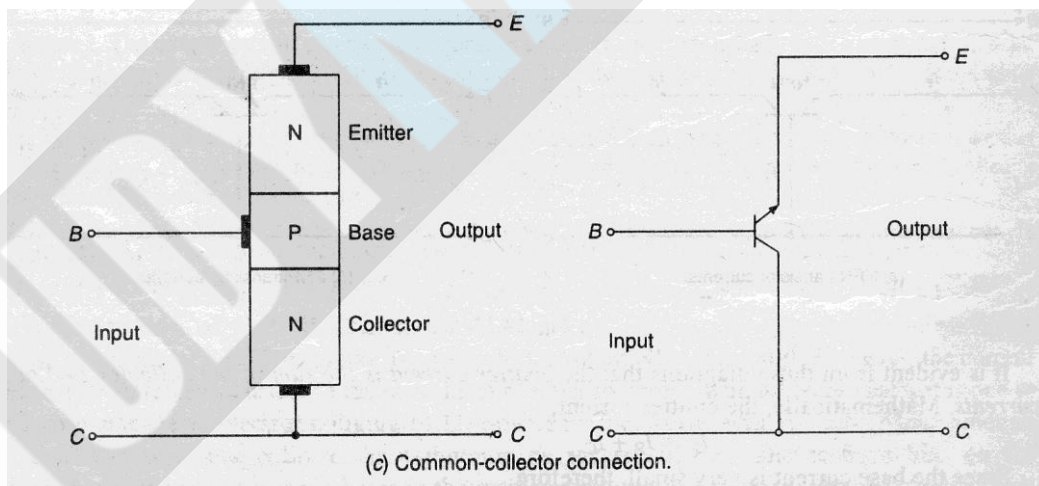
Common emitter configuration:

- ◆ In this configuration emitter terminal is connected as a common terminal.
- ◆ The input is applied between the base and emitter terminals. The output is taken between the collector and base terminals.



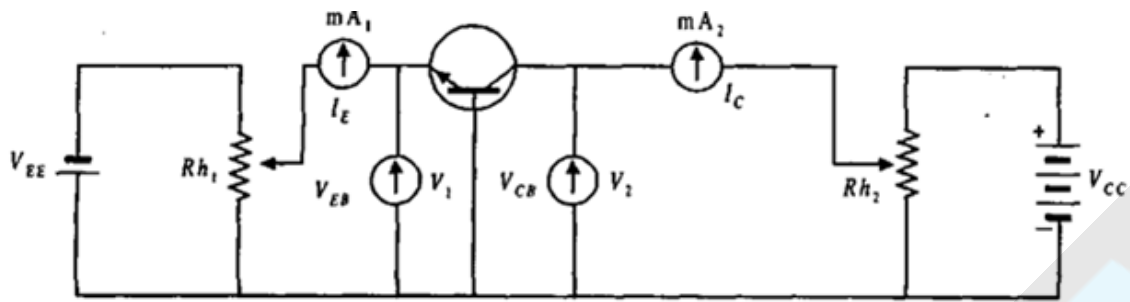
Common collector configuration:

- ◆ In this configuration collector terminal is connected as a common terminal.
- ◆ The input is applied between the base and collector terminals. The output is taken between the emitter and collector terminals

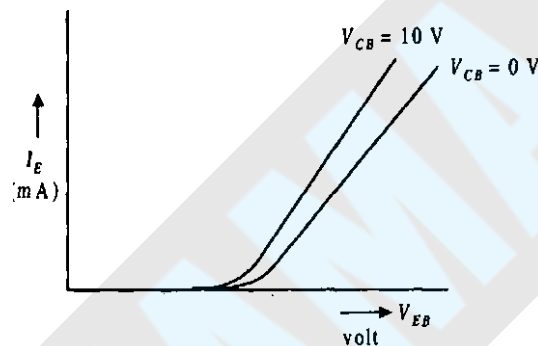


CB, CE, CC Characteristics :

1. Common base Characteristics :



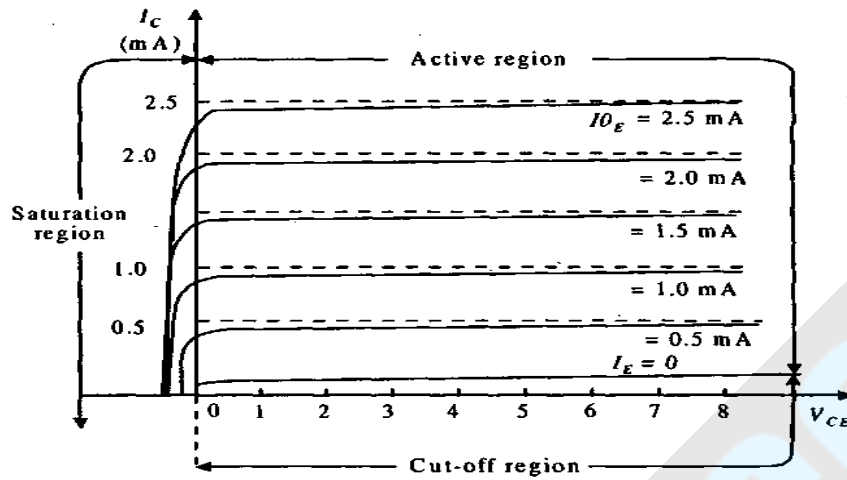
Input characteristics:



- ❖ The output(CB) voltage is maintained constant and the input voltage (EB) is set at several convenient levels. For each level of input voltage, the input current I_E is recorded.
- ❖ I_E is then plotted versus V_{EB} to give the common-base input characteristics.
- ❖ The EB junction is essentially the same as a forward biased diode, therefore the current-voltage characteristics is essentially the same as that of a diode:

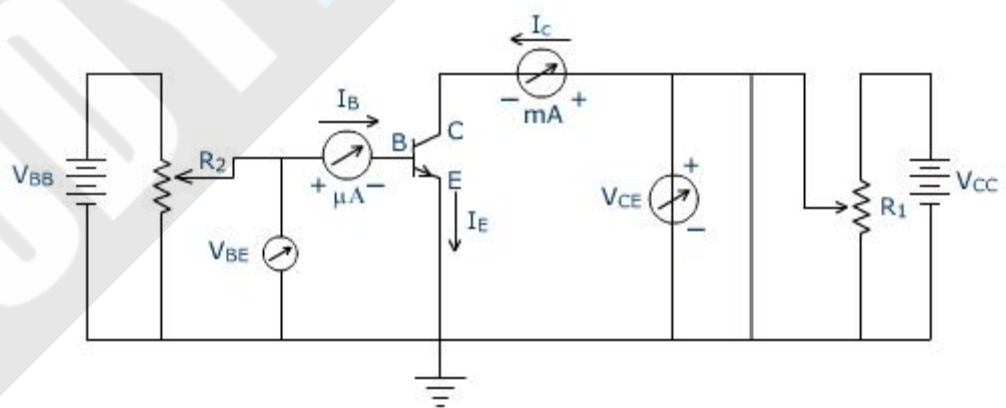
$$I_E = I_0(e^{V_{BE}/V_T} - 1)$$

Output characteristics:

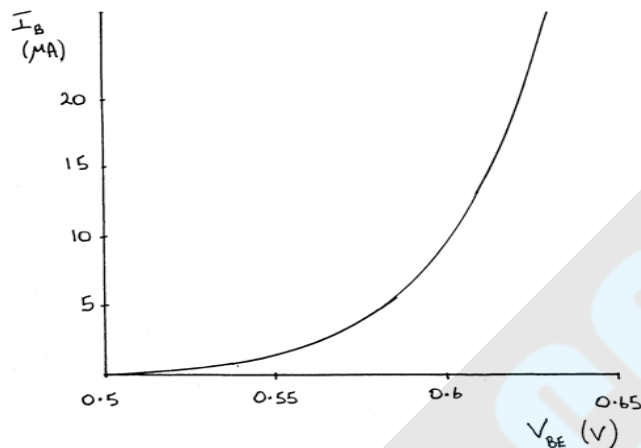


- ❖ The emitter current I_E is held constant at each of several fixed levels. For each fixed value of I_E , the output voltage V_{CB} is adjusted in convenient steps and the corresponding levels of collector current I_C are recorded.
- ❖ For each fixed value of I_E , I_C is almost equal to I_E and appears to remain constant when V_{CB} is increased.

2.Common-Emitter Characteristics :

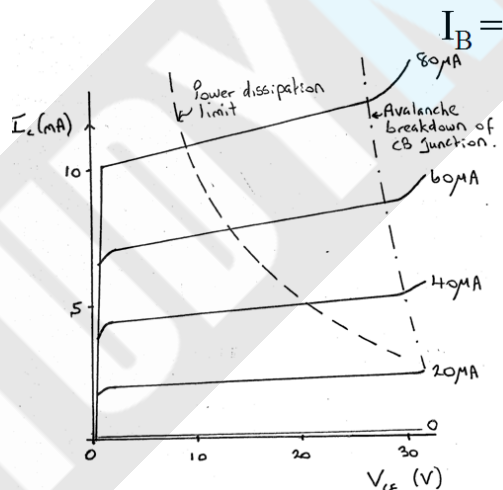


Input characteristics:



- ❖ The output voltage V_{CE} is maintained constant and the input voltage V_{BE} is set at several convenient levels. For each level of input voltage, the input current I_B is recorded.
- ❖ I_B is then plotted versus V_{BE} to give the common-base input characteristics.

Output characteristics:



Ideal Output characteristic

Actual Output Characteristic

- ❖ The Base current I_B is held constant at each of several fixed levels. For each fixed value of I_B , the output voltage V_{CE} is

adjusted in convenient steps and the corresponding levels of collector current I_C are recorded

- ❖ .For each fixed value of I_B , I_C level is Recorded at each V_{CE} step. For each I_B level, I_C is plotted versus V_{CE} to give a family of characteristics.

3.Common-Collector Characteristics :

Input characteristics:

- ❖ The common-collector input characteristics are quite different from either common base or common-emitter input characteristics.
- ❖ The difference is due to the fact that the input voltage (V_{BC}) is largely determined by (V_{EC}) level .

$$V_{EC} = V_{EB} + V_{BC}$$

$$V_{EB} = V_{EC} - V_{BC}$$

Output characteristics:

- ❖ The operation is much similar to that of C-E configuration. When the base current is I_{CO} , the emitter current will be zero and consequently no current will flow in the load.
- ❖ When the base current is increased, the transistor passes through active region and eventually reaches saturation.

Under the saturation conditions all the supply voltage, except for a very small drop across the transistor will appear across the load resistor.

Early Effect and Early Voltage

Base-Width Modulation:

- ❖ As reverse-bias across the collector-base junction increases, the width of the collector-base depletion layer increases and the effective width of base decreases.
- ❖ This is called “base-width modulation”.

Early effect:

When the output characteristics are extrapolated back to where the i_C curves intersect at common point, $v_{CE} = -V_A$ (Early voltage), which lies between 15 V and 150 V.

Comparison of the CB, CE, and CC Configurations

	CB	CE	CC
Voltage Gain	High	Medium	Low
Current Gain	Low	Medium	High
Power Gain	Low	High	Medium
Input Resistance	Low	Medium	High
Output Resistance	High	Medium	Low

PART A

1. Why do we choose q point at the center of the loadline?
2. Name the two techniques used in the stability of the q point ,explain.
3. Why is the operating point selected at the Centre of the active region?
4. List out the different types of biasing.
5. What do you meant by thermal runaway?
6. Why is the transistor called a current controlled device?
7. Define current amplification factor?
8. What are the requirements for biasing circuits?
9. When does a transistor act as a switch?
10. What is biasing?
11. What is operating point?
12. What is stability factor?
13. What is d.c load line?
14. What are the advantages of fixed bias circuit?
15. Explain about the various regions in a transistor?
16. Explain about the characteristics of a transistor?
17. What is the necessary of the coupling capacitor?
18. What is reverse saturation current?

PART B

1. Explain the need for biasing , Stability factor and Fixed bias circuit
2. Explain in detail different types of biasing circuits
3. Explain the advantage of self bias (voltage divider bias) over other types of biasing
4. Explain the various types of bias compensation techniques.
5. i) Explain biasing of FET
ii) Explain biasing of MOSFET

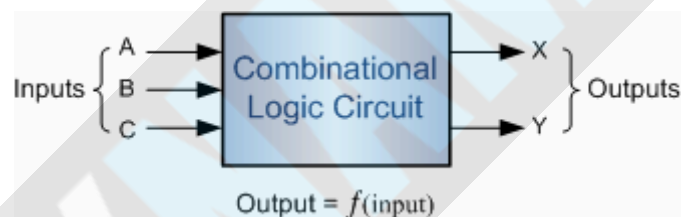
UNIT IV

DIGITAL ELECTRONICS

Combinational Logic Circuits

Unlike **Sequential Logic Circuits** whose outputs are dependant on both their present inputs and their previous output state giving them some form of **Memory**, the outputs of **Combinational Logic Circuits** are only determined by the logical function of their current input state, logic "0" or logic "1", at any given instant in time as they have no feedback, and any changes to the signals being applied to their inputs will immediately have an effect at the output. In other words, in a **Combinational Logic Circuit**, the output is dependant at all times on the combination of its inputs and if one of its inputs condition changes state so does the output as combinational circuits have "no memory", "timing" or "feedback loops".

Combinational Logic



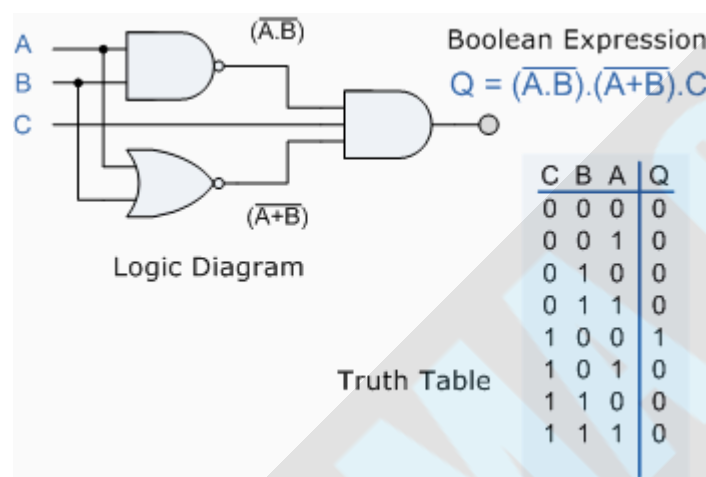
Combination Logic Circuits are made up from basic logic **NAND**, **NOR** or **NOT** gates that are "combined" or connected together to produce more complicated switching circuits. These logic gates are the building blocks of combinational logic circuits. An example of a combinational circuit is a decoder, which converts the binary code data present at its input into a number of different output lines, one at a time producing an equivalent decimal code at its output.

Combinational logic circuits can be very simple or very complicated and any combinational circuit can be implemented with only **NAND** and **NOR** gates as these are classed as "universal" gates. The three main ways of specifying the function of a combinational logic circuit are:

- **Truth Table** Truth tables provide a concise list that shows the output values in tabular form for each possible combination of input variables.
-

- **Boolean Algebra** Forms an output expression for each input variable that represents a logic "1"
-
- **Logic Diagram** Shows the wiring and connections of each individual logic gate that implements the circuit.

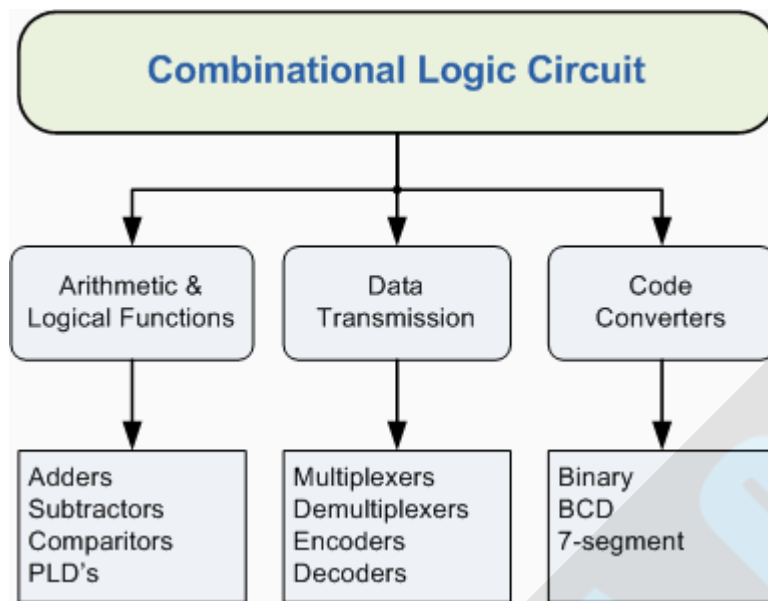
and all three are shown below.



As combination logic circuits are made up from individual logic gates only, they can also be considered as "decision making circuits" and combinational logic is about combining logic gates together to process two or more signals in order to produce at least one output signal according to the logical function of each logic gate. Common combinational circuits made up from individual logic gates that carry out a desired application include **Multiplexers, De-multiplexers, Encoders, Decoders, Full and Half Adders** etc.

Classification of Combinational Logic

One of the most common uses of combination logic is in **Multiplexer** and **De-multiplexer** type circuits. Here, multiple inputs or outputs are connected to a common signal line and logic gates are used to decode an address to select a single data input or output switch. A multiplexer consist of two separate components, a logic decoder and some solid state switches, but before we can discuss multiplexers, decoders and de-multiplexers in more detail we first need to understand how these devices use these "solid state switches" in their design.

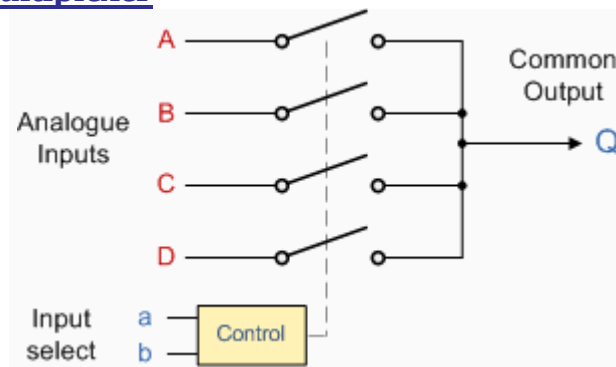


The Multiplexer

A data selector, more commonly called a **Multiplexer**, shortened to "Mux" or "MPX", are combinational logic switching devices that operate like a very fast acting multiple position rotary switch. They connect or control, multiple input lines called "channels" consisting of either 2, 4, 8 or 16 individual inputs, one at a time to an output. Then the job of a multiplexer is to allow multiple signals to *share* a single common output. For example, a single 8-channel multiplexer would connect one of its eight inputs to the single data output. Multiplexers are used as one method of reducing the number of logic gates required in a circuit or when a single data line is required to carry two or more different digital signals.

Digital **Multiplexers** are constructed from individual **analogue switches** encased in a single IC package as opposed to the "mechanical" type selectors such as normal conventional switches and relays. Generally, multiplexers have an even number of data inputs, usually an even power of two, n^2 , a number of "control" inputs that correspond with the number of data inputs and according to the binary condition of these control inputs, the appropriate data input is connected directly to the output. An example of a **Multiplexer** configuration is shown below.

4-to-1 Channel Multiplexer



Addressing		Input Selected
b	a	
0	0	A
0	1	B
1	0	C
1	1	D

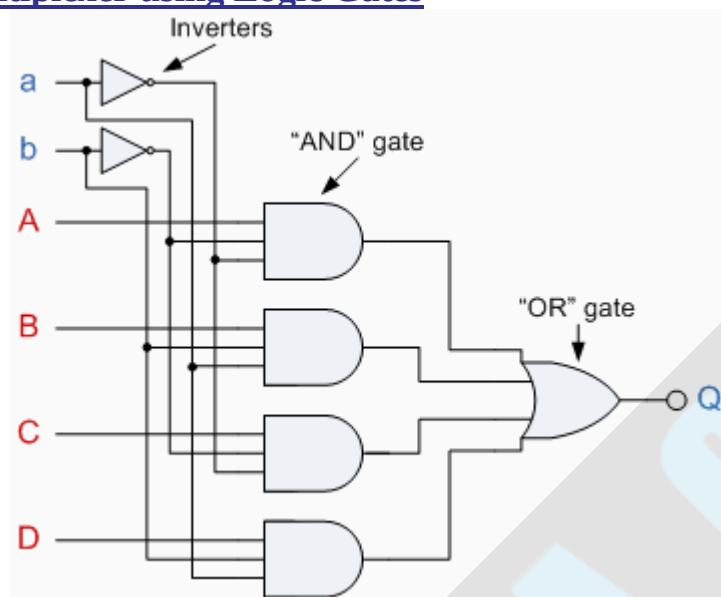
The Boolean expression for this 4-to-1 **Multiplexer** above with inputs *A* to *D* and data select lines *a*, *b* is given as:

$$Q = abA + abB + abC + abD$$

In this example at any one instant in time only ONE of the four analogue switches is closed, connecting only one of the input lines *A* to *D* to the single output at *Q*. As to which switch is closed depends upon the addressing input code on lines "*a*" and "*b*", so for this example to select input *B* to the output at *Q*, the binary input address would need to be "*a*" = logic "1" and "*b*" = logic "0". Adding more control address lines will allow the multiplexer to control more inputs but each control line configuration will connect only ONE input to the output.

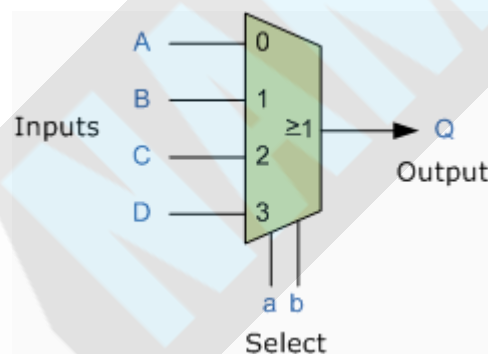
Then the implementation of this Boolean expression above using individual logic gates would require the use of seven individual gates consisting of **AND**, **OR** and **NOT** gates as shown.

4 Channel Multiplexer using Logic Gates



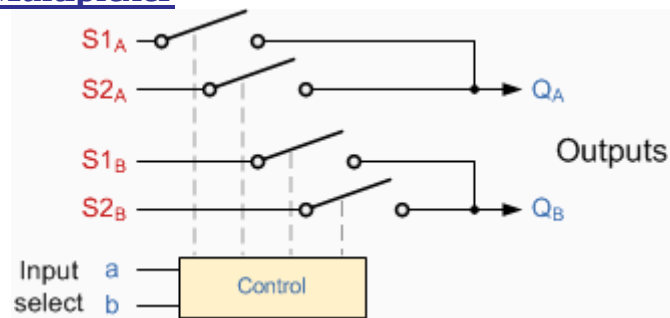
The symbol used in logic diagrams to identify a multiplexer is as follows.

Multiplexer Symbol



Multiplexers are not limited to just switching a number of different input lines or channels to one common single output. There are also types that can switch their inputs to multiple outputs and have arrangements or 4 to 2, 8 to 3 or even 16 to 4 etc configurations and an example of a simple Dual channel 4 input multiplexer (4 to 2) is given below:

4-to-2 Channel Multiplexer

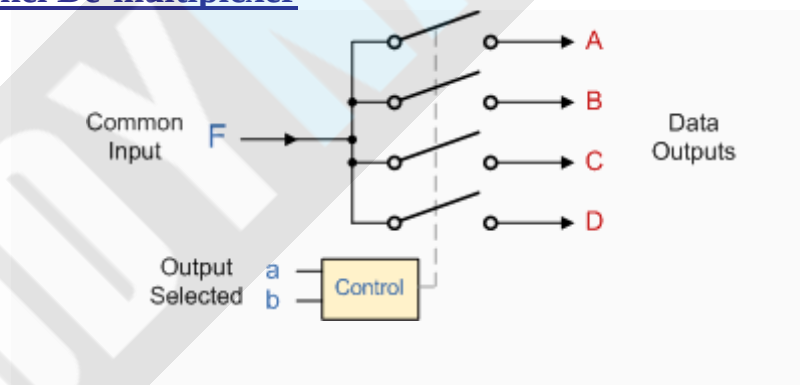


Here in this example the 4 input channels are switched to 2 individual output lines but larger arrangements are also possible. This simple 4 to 2 configuration could be used for example, to switch audio signals for stereo pre-amplifiers or mixers.

The Demultiplexer

The data distributor, known more commonly as a **Demultiplexer** or "Demux", is the exact opposite of the **Multiplexer** we saw in the previous tutorial. The demultiplexer takes one single input data line and then switches it to any one of a number of individual output lines one at a time. The **demultiplexer** converts a serial data signal at the input to a parallel data at its output lines as shown below.

1-to-4 Channel De-multiplexer



Addressing		Output Selected
b	a	
0	0	A
0	1	B
1	0	C
1	1	D

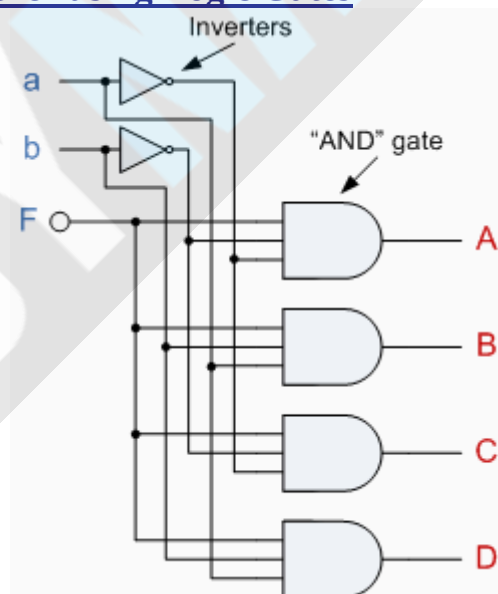
The Boolean expression for this 1-to-4 **Demultiplexer** above with outputs **A** to **D** and data select lines **a**, **b** is given as:

$$F = ab A + abB + abC + abD$$

The function of the **Demultiplexer** is to switch one common data input line to any one of the 4 output data lines **A** to **D** in our example above. As with the multiplexer the individual solid state switches are selected by the binary input address code on the output select pins "**a**" and "**b**" and by adding more address line inputs it is possible to switch more outputs giving a 1-to-2ⁿ data line outputs. Some standard demultiplexer IC's also have an "enable output" input pin which disables or prevents the input from being passed to the selected output. Also some have latches built into their outputs to maintain the output logic level after the address inputs have been changed. However, in standard decoder type circuits the address input will determine which single data output will have the same value as the data input with all other data outputs having the value of logic "0".

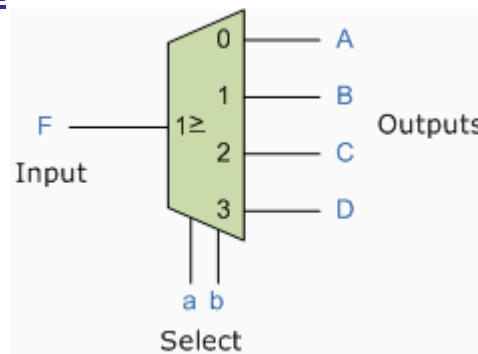
The implementation of the Boolean expression above using individual logic gates would require the use of six individual gates consisting of **AND** and **NOT** gates as shown.

4 Channel Demultiplexer using Logic Gates



The symbol used in logic diagrams to identify a demultiplexer is as follows.

Demultiplexer Symbol



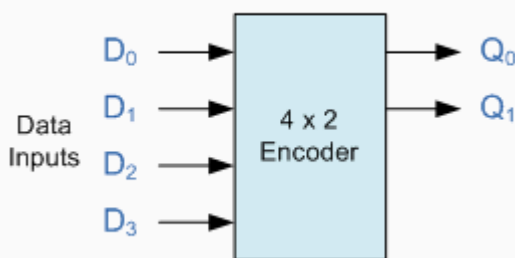
Standard **Demultiplexer** IC packages available are the TTL 74LS138 1 to 8-output demultiplexer, the TTL 74LS139 Dual 1-to-4 output demultiplexer or the CMOS CD4514 1-to-16 output demultiplexer. Another type of demultiplexer is the 24-pin, 74LS154 which is a 4-bit to 16-line demultiplexer/decoder. Here the individual output positions are selected using a 4-bit binary coded input. Like multiplexers, demultiplexers can also be cascaded together to form higher order demultiplexers.

Unlike multiplexers which convert data from a single data line to multiple lines and demultiplexers which convert multiple lines to a single data line, there are devices available which convert data to and from multiple lines and in the next tutorial about combinational logic devices, we will look at **Encoders** which convert multiple input lines into multiple output lines, converting the data from one form to another.

The Digital Encoder

Unlike a multiplexer that selects one individual data input line and then sends that data to a single output line or switch, a **Digital Encoder** more commonly called a **Binary Encoder** takes *ALL* its data inputs one at a time and then converts them into a single encoded output. So we can say that a binary encoder, is a multi-input combinational logic circuit that converts the logic level "1" data at its inputs into an equivalent binary code at its output. Generally, digital encoders produce outputs of 2-bit, 3-bit or 4-bit codes depending upon the number of data input lines. An "n-bit" binary encoder has 2^n input lines and n-bit output lines with common types that include 4-to-2, 8-to-3 and 16-to-4 line configurations. The output lines of a digital encoder generate the binary equivalent of the input line whose value is equal to "1" and are available to encode either a decimal or hexadecimal input pattern to typically a binary or B.C.D. output code.

4-to-2 Bit Binary Encoder



Inputs				Outputs	
D ₃	D ₂	D ₁	D ₀	Q ₁	Q ₀
0	0	0	1	0	0
0	0	1	0	0	1
0	1	0	0	1	0
1	0	0	0	1	1
0	0	0	0	x	x

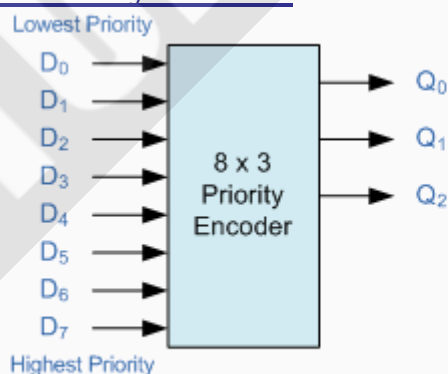
One of the main disadvantages of standard digital encoders is that they can generate the wrong output code when there is more than one input present at logic level "1". For example, if we make inputs D_1 and D_2 HIGH at logic "1" at the same time, the resulting output is neither at "01" or at "10" but will be at "11" which is an output binary number that is different to the actual input present. Also, an output code of all logic "0"s can be generated when all of its inputs are at "0" OR when input D_0 is equal to one.

One simple way to overcome this problem is to "Prioritise" the level of each input pin and if there was more than one input at logic level "1" the actual output code would only correspond to the input with the highest designated priority. Then this type of digital encoder is known commonly as a **Priority Encoder** or **P-encoder** for short.

Priority Encoders

Priority Encoders solve the problem mentioned above by allocating a priority level to each input. The encoder output corresponds to the currently active input with the highest priority. So when an input with a higher priority is present, all other inputs with a lower priority will be ignored. Priority encoders come in many forms with an example of an 8-input priority encoder along with its truth table shown below.

8-to-3 Bit Priority Encoder



Inputs								Outputs		
D ₇	D ₆	D ₅	D ₄	D ₃	D ₂	D ₁	D ₀	Q ₂	Q ₁	Q ₀
0	0	0	0	0	0	0	1	0	0	0
0	0	0	0	0	0	1	x	0	0	1
0	0	0	0	0	1	x	x	0	1	0
0	0	0	0	1	x	x	x	0	1	1
0	0	0	1	x	x	x	x	1	0	0
0	0	1	x	x	x	x	x	1	0	1
0	1	x	x	x	x	x	x	1	1	0
1	x	x	x	x	x	x	x	1	1	1

X = dont care

Priority encoders are available in standard IC form and the TTL 74LS148 is an 8 to 3 bit priority encoder which has eight active LOW (logic "0") inputs and provides a 3-bit code of the highest ranked input at its output. Priority encoders output the highest order input first for example, if input lines "D2", "D3" and "D5" are applied simultaneously the output code would be for input "D5" ("101") as this has the highest order out of the 3 inputs. Once input "D5" had been removed the next highest output code would be for input "D3" ("011"), and so on.

The Boolean expression for this 8-to-3 encoder above with inputs D_0 to D_7 and outputs Q_0 , Q_1 , Q_2 is given as:

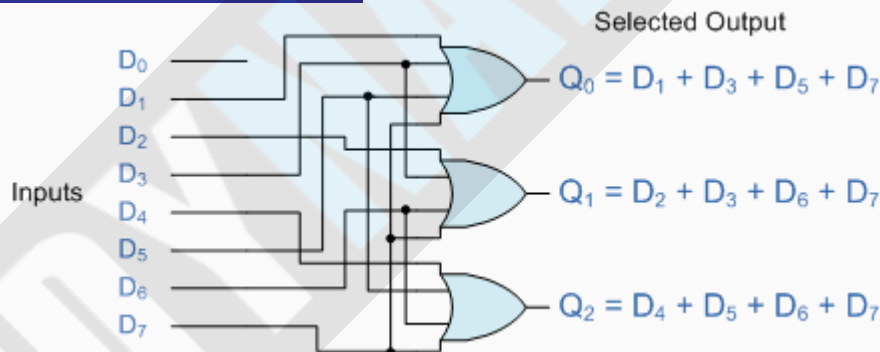
$$Q_0 = D_1 + D_3 + D_5 + D_7$$

$$Q_1 = D_2 + D_3 + D_6 + D_7$$

$$Q_2 = D_4 + D_5 + D_6 + D_7$$

Then the implementation of these Boolean expression outputs above using individual OR gates is as follows.

Digital Encoder using Logic Gates

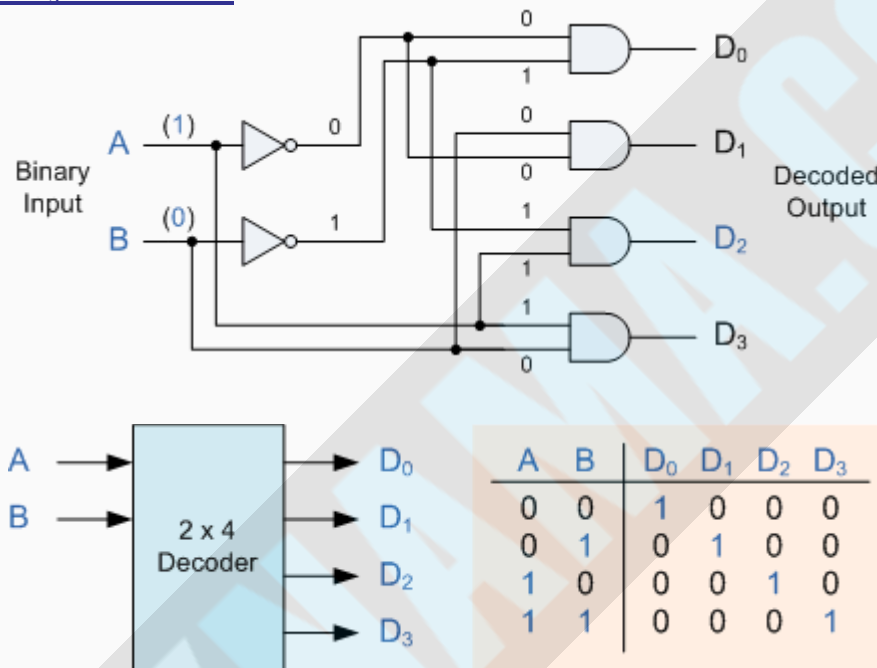


Binary Decoder

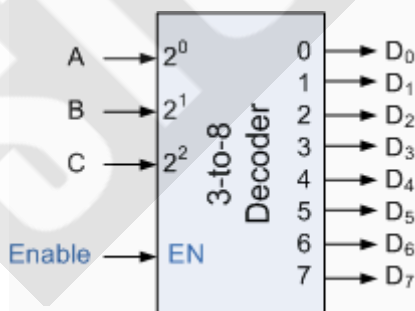
A **Decoder** is the exact opposite to that of an "Encoder" we looked at in the last tutorial. It is basically, a combinational type logic circuit that converts the binary code data at its input into one of a number of different output lines, one at a time producing an equivalent decimal code at its output. **Binary Decoders** have inputs of 2-bit, 3-bit or 4-bit codes depending upon the number of data input lines, and a n -bit decoder has 2^n output lines. Therefore, if it receives n inputs (usually grouped as a binary or Boolean number) it activates one and only one of its 2^n outputs based on that input with all other outputs deactivated. A decoder's output code normally has more bits than its input code and practical binary decoder circuits include, 2-to-4, 3-to-8 and 4-to-16 line configurations.

A binary decoder converts coded inputs into coded outputs, where the input and output codes are different and decoders are available to "decode" either a Binary or BCD (8421 code) input pattern to typically a Decimal output code. Commonly available BCD-to-Decimal decoders include the TTL 7442 or the CMOS 4028. An example of a 2-to-4 line decoder along with its truth table is given below. It consists of an array of four **NAND** gates, one of which is selected for each combination of the input signals **A** and **B**.

A 2-to-4 Binary Decoders.



In this simple example of a 2-to-4 line binary decoder, the binary inputs **A** and **B** determine which output line from **D₀** to **D₃** is "HIGH" at logic level "1" while the remaining outputs are held "LOW" at logic "0" so only one output can be active (HIGH) at any one time. Therefore, whichever output line is "HIGH" identifies the binary code present at the input, in other words it "de-codes" the binary input and these types of binary decoders are commonly used as **Address Decoders** in microprocessor memory applications.

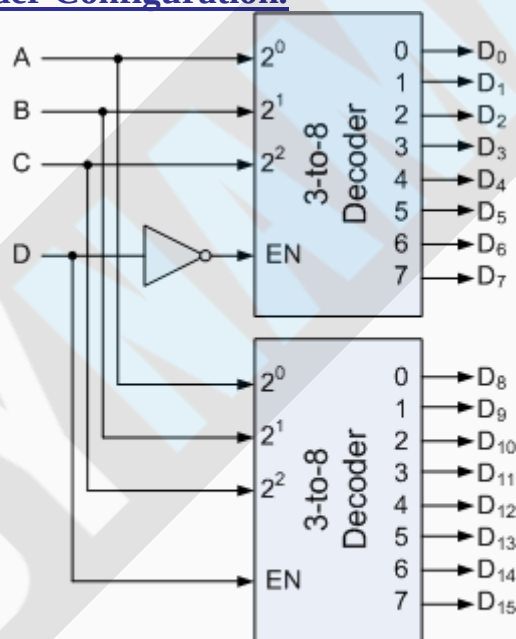


74LS138 Binary Decoder

Some binary decoders have an additional input labelled "Enable" that controls the outputs from the device. This allows the decoders outputs to be turned "ON" or "OFF" and we can see that the logic diagram of the basic decoder is identical to that of the basic demultiplexer. Therefore, we say that a demultiplexer is a decoder with an additional data line that is used to enable the decoder. An alternative way of looking at the decoder circuit is to regard inputs **A**, **B** and **C** as address signals. Each combination of **A**, **B** or **C** defines a unique address which can access a location having that address.

Sometimes it is required to have a **Binary Decoder** with a number of outputs greater than is available, or if we only have small devices available, we can combine multiple decoders together to form larger decoder networks as shown. Here a much larger 4-to-16 line binary decoder has been implemented using two smaller 3-to-8 decoders.

A 4-to-16 Binary Decoder Configuration.



4-to-16 Line Decoder Implemented
with two 3-to-8 Decoders

Inputs **A**, **B**, **C** are used to select which output on either decoder will be at logic "1" (HIGH) and input **D** is used with the enable input to select which encoder either the first or second will output the "1"

BCD to 7-Segment Display Decoder

As we saw in the previous tutorial, a **Decoder IC**, is a device which converts one digital format into another and the most commonly used device for doing this is the

Binary Coded Decimal (BCD) to 7-Segment Display Decoder. 7-segment **LED** (Light Emitting Diode) or **LCD** (Liquid Crystal) displays, provide a very convenient way of displaying information or digital data in the form of numbers, letters or even alpha-numerical characters and they consist of 7 individual LED's (the segments), within one single display package.

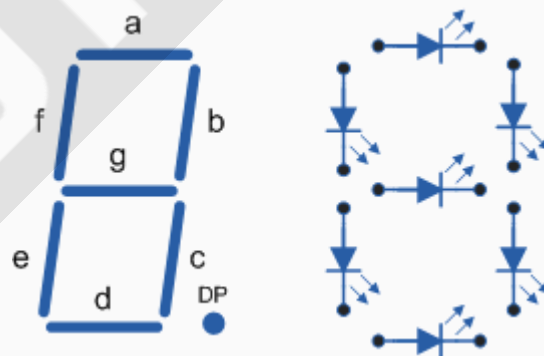
In order to produce the required numbers or HEX characters from 0 to 9 and A to F respectively, on the display the correct combination of LED segments need to be illuminated and **BCD to 7-segment Display Decoders** such as the 74LS47 do just that. A standard 7-segment LED display generally has 8 input connections, one for each LED segment and one that acts as a common terminal or connection for all the internal segments. Some single displays have an additional input pin for the decimal point in their lower right or left hand corner.

There are two important types of 7-segment LED digital display.

The Common Cathode Display (CCD) - In the common cathode display, all the cathode connections of the LED's are joined together to logic "0" and the individual segments are illuminated by application of a "HIGH", logic "1" signal to the individual Anode terminals.

The Common Anode Display (CAD) - In the common anode display, all the anode connections of the LED's are joined together to logic "1" and the individual segments are illuminated by connecting the individual Cathode terminals to a "LOW", logic "0" signal.

7-Segment Display Format



Truth Table for a 7-segment display

Individual Segments							Display
a	b	c	d	e	f	g	
x	x	x	x	x	x		0

Individual Segments							Display
a	b	c	d	e	f	g	
x	x	x	x	x	x	x	8

	x	x					1
x	x		x	x		x	2
x	x	x	x			x	3
	x	x			x	x	4
x		x	x		x	x	5
x		x	x	x	x	x	6
x	x	x					7

x	x	x			x	x	9
x	x	x		x	x	x	A
		x	x	x	x	x	b
x			x	x	x		C
	x	x	x	x		x	d
x			x	x	x	x	E
x				x	x	x	F



7-Segment Display Elements for all Numbers.

It can be seen that to display any single digit number from 0 to 9 or letter from A to F, we would need 7 separate segment connections plus one additional connection for the LED's "common" connection. Also as the segments are basically a standard light emitting diode, the driving circuit would need to produce up to 20mA of current to illuminate each individual segment and to display the number 8, all 7 segments would need to be lit resulting a total current of nearly 140mA, (8 x 20mA). Obviously, the use of so many connections and power consumption is impractical for some electronic or microprocessor based circuits and so in order to reduce the number of signal lines required to drive just one single display, display decoders such as the BCD to 7-Segment Display Decoder and Driver IC's are used instead.

Binary Coded Decimal

Binary Coded Decimal (BCD or "8421" BCD) numbers are made up using just 4 data bits (a nibble or half a byte) similar to the **Hexadecimal** numbers we saw in the binary tutorial, but unlike hexadecimal numbers that range in full from 0 through to F, BCD numbers only range from 0 to 9, with the binary number patterns of 1010 through to 1111 (A to F) being invalid inputs for this type of display and so are not used as shown below.

Decimal	Binary Pattern				BCD
	8	4	2	1	
0	0	0	0	0	0

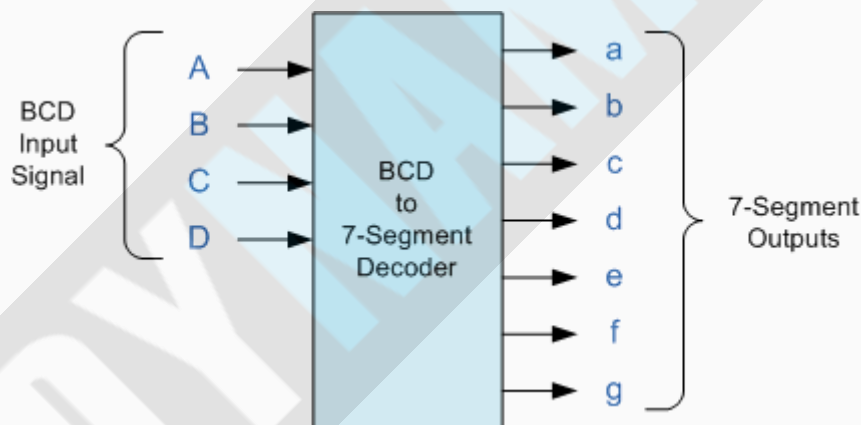
Decimal	Binary Pattern				BCD
	8	4	2	1	
8	1	0	0	0	8

1	0	0	0	1	1	9	1	0	0	1	9
2	0	0	1	0	2	10	1	0	1	0	Invalid
3	0	0	1	1	3	11	1	0	1	1	Invalid
4	0	1	0	0	4	12	1	1	0	0	Invalid
5	0	1	0	1	5	13	1	1	0	1	Invalid
6	0	1	1	0	6	14	1	1	1	0	Invalid
7	0	1	1	1	7	15	1	1	1	1	Invalid

BCD to 7-Segment Display Decoders

A binary coded decimal (BCD) to 7-segment display decoder such as the TTL 74LS47 or 74LS48, have 4 BCD inputs and 7 output lines, one for each LED segment. This allows a smaller 4-bit binary number (half a byte) to be used to display all the denary numbers from 0 to 9 and by adding two displays together, a full range of numbers from 00 to 99 can be displayed with just a single byte of 8 data bits.

BCD to 7-Segment Decoder



The use of **packed** BCD allows two BCD digits to be stored within a single byte (8-bits) of data, allowing a single data byte to hold a BCD number in the range of 00 to 99.

The Binary Adder

Another common and very useful combinational logic circuit which can be constructed using just a few basic logic gates and adds together binary numbers is the **Binary Adder** circuit. The Binary Adder is made up from standard **AND** and **Ex-OR** gates and allow us to "add" together single bit binary numbers, **a** and **b** to produce two outputs, the **SUM** of the addition and a **CARRY** called the **Carry-out**, (**Cout**) bit. One of the main uses for the **Binary Adder** is in arithmetic and counting

circuits.

Consider the addition of two denary (base 10) numbers below.

$$\begin{array}{r} 123 \quad A \quad (\text{Augend}) \\ + 789 \quad B \quad (\text{Addend}) \\ \hline 912 \quad \text{SUM} \end{array}$$

Each column is added together starting from the right hand side and each digit has a weighted value depending upon its position in the columns. As each column is added together a carry is generated if the result is greater or equal to ten, the base number. This carry is then added to the result of the addition of the next column to the left and so on, simple school math's addition. The adding of binary numbers is basically the same as that of adding decimal numbers but this time a carry is only generated when the result in any column is greater or equal to "2", the base number of binary.

Binary Addition

Binary Addition follows the same basic rules as for the denary addition above except in binary there are only two digits and the largest digit is "1", so any "SUM" greater than 1 will result in a "CARRY". This carry 1 is passed over to the next column for addition and so on. Consider the single bit addition below.

$$\begin{array}{r} 0 \quad 0 \quad 1 \quad 1 \\ + 0 \quad + 1 \quad + 0 \quad + 1 \\ \hline 0 \quad 1 \quad 1 \quad 10 \end{array}$$

The single bits are added together and "0 + 0", "0 + 1", or "1 + 0" results in a sum of "0" or "1" until you get to "1 + 1" then the sum is equal to "2". For a simple 1-bit addition problem like this, the resulting carry bit could be ignored which would result in an output truth table resembling that of an **Ex-OR Gate** as seen in the Logic Gates section and whose result is the sum of the two bits but without the carry. An **Ex-OR** gate only produces an output "1" when either input is at logic "1", but not both. However, all microprocessors and electronic calculators require the carry bit to correctly calculate the equations so we need to rewrite them to include 2 bits of output data as shown below.

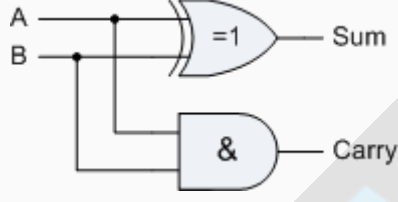
$$\begin{array}{r} 00 \quad 00 \quad 01 \quad 01 \\ + 00 \quad + 01 \quad + 00 \quad + 01 \end{array}$$

00 01 01 10

From the above equations we know that an **Ex-OR** gate will only produce an output "1" when "EITHER" input is at logic "1", so we need an additional output to produce a carry output, "1" when "BOTH" inputs "A" and "B" are at logic "1" and a standard **AND Gate** fits the bill nicely. By combining the **Ex-OR** gate with the **AND** gate results in a simple digital binary adder circuit known commonly as the "**Half Adder**" circuit.

The Half Adder Circuit

1-bit Adder with Carry-Out

Symbol	Truth Table			
	A	B	SUM	CARRY
	0	0	0	0
	0	1	1	0
	1	0	1	0
	1	1	0	1
Boolean Expression: Sum = $A \oplus B$ Carry = $A \cdot B$				

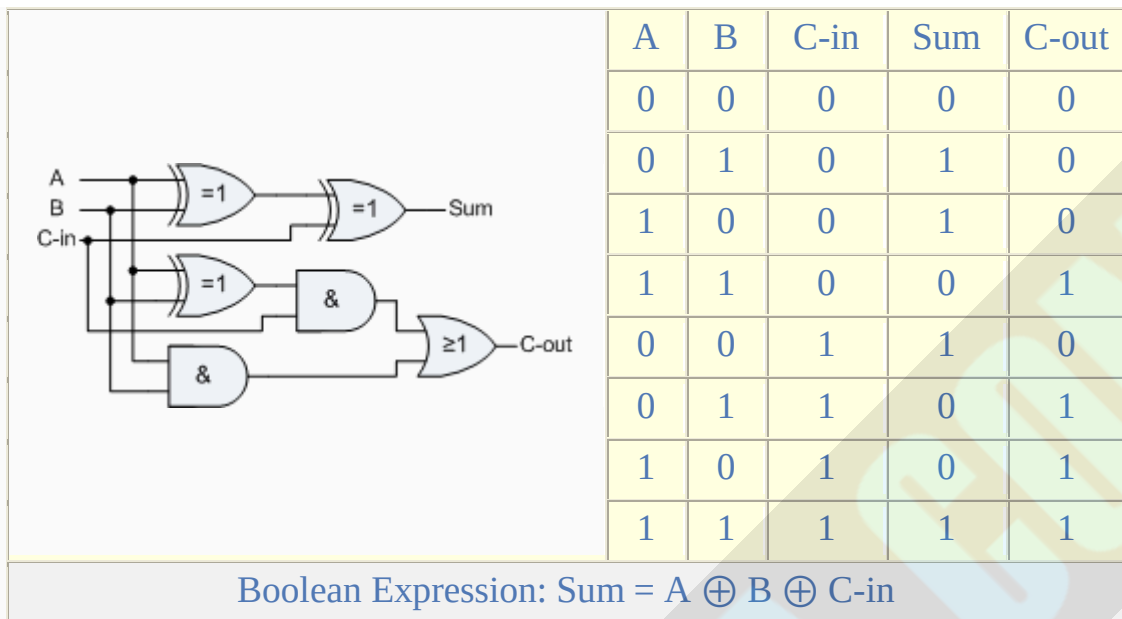
From the truth table we can see that the SUM (S) output is the result of the **Ex-OR** gate and the Carry-out (**Cout**) is the result of the **AND** gate. One major disadvantage of the Half Adder circuit when used as a binary adder, is that there is no provision for a "Carry-in" from the previous circuit when adding together multiple data bits. For example, suppose we want to add together two 8-bit bytes of data, any resulting carry bit would need to be able to "ripple" or move across the bit patterns starting from the least significant bit (LSB). The most complicated operation the half adder can do is "1 + 1" but as the half adder has no carry input the resultant added value would be incorrect. One simple way to overcome this problem is to use a **Full Adder** type binary adder circuit.

The Full Adder Circuit

The main difference between the **Full Adder** and the previous seen **Half Adder** is that a full adder has three inputs, the same two single bit binary inputs **A** and **B** as before plus an additional **Carry-In (C-in)** input as shown below.

Full Adder with Carry-In

Symbol	Truth Table
--------	-------------

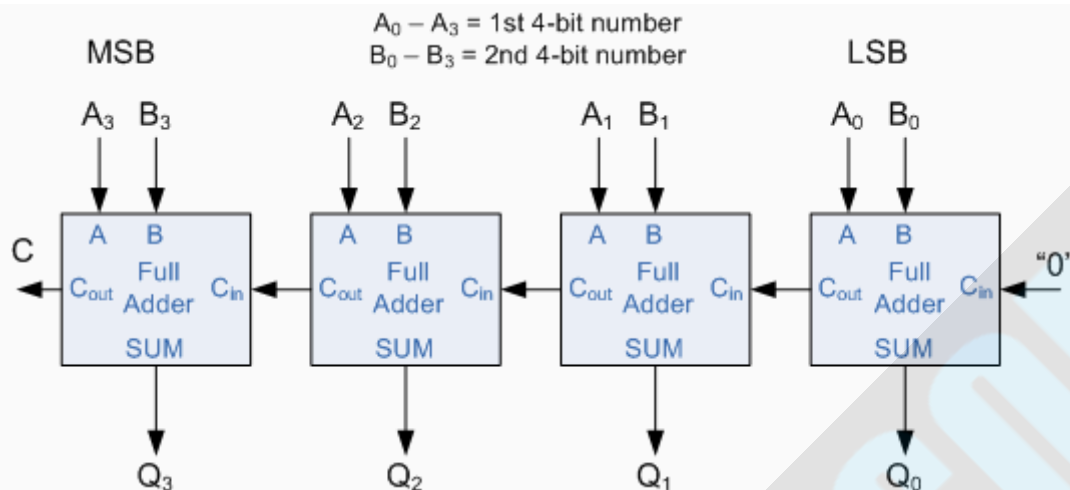


The 1-bit **Full Adder** circuit above is basically two half adders connected together and consists of three **Ex-OR** gates, two **AND** gates and an **OR** gate, six logic gates in total. The truth table for the full adder includes an additional column to take into account the Carry-in input as well as the summed output and carry-output. 4-bit full adder circuits are available as standard IC packages in the form of the TTL 74LS83 or the 74LS283 which can add together two 4-bit binary numbers and generate a **SUM** and a **CARRY** output. But what if we wanted to add together two **n-bit** numbers, then **n** 1-bit full adders need to be connected together to produce what is known as the **Ripple Carry Adder**.

The 4-bit Binary Adder

The **Ripple Carry Binary Adder** is simply **n**, full adders cascaded together with each full adder represents a single weighted column in the long addition with the carry signals producing a "ripple" effect through the binary adder from right to left. For example, suppose we want to "add" together two 4-bit numbers, the two outputs of the first full adder will provide the first place digit sum of the addition plus a carry-out bit that acts as the carry-in digit of the next binary adder. The second binary adder in the chain also produces a summed output (the 2nd bit) plus another carry-out bit and we can keep adding more full adders to the combination to add larger numbers, linking the carry bit output from the first full binary adder to the next full adder, and so forth. An example of a 4-bit adder is given below.

A 4-bit Binary Adder

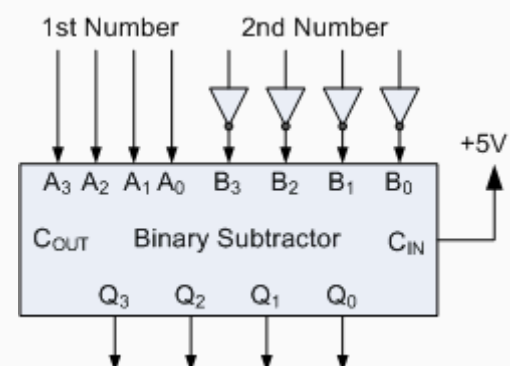


One main disadvantage of "cascading" together 1-bit **binary adders** to add large binary numbers is that if inputs **A** and **B** change, the sum at its output will not be valid until any carry-input has "rippled" through every full adder in the chain. Consequently, there will be a finite delay before the output of an adder responds to a change in its inputs resulting in the accumulated delay especially in large multi-bit binary adders becoming prohibitively large. This delay is called **Propagation delay**. Also "overflow" occurs when an **n-bit** adder adds two numbers together whose sum is greater than or equal to 2^n

One solution is to generate the carry-input signals directly from the **A** and **B** inputs rather than using the ripple arrangement above. This then produces another type of binary adder circuit called a **Carry Look Ahead Binary Adder** where the speed of the parallel adder can be greatly improved using carry-look ahead logic.

The 4-bit Binary Subtractor

Now that we know how to "ADD" together two 4-bit binary numbers how would we subtract two 4-bit binary numbers, for example, **A - B** using the circuit above. The answer is to use 2's-complement notation on all the bits in **B** must be complemented (inverted) and an extra one added using the carry-input. This can be achieved by inverting each **B** input bit using an inverter or **NOT-gate**.



Also, in the above circuit for the 4-bit binary adder, the first carry-in input is held LOW at logic "0", for the circuit to perform subtraction this input needs to be held HIGH at "1". With this in mind a ripple carry adder can with a small modification be used to perform half subtraction, full subtraction and/or comparison.

There are a number of 4-bit full-adder ICs available such as the 74LS283 and CD4008. which will add two 4-bit binary number and provide an additional input carry bit, as well as an output carry bit, so you can cascade them together to produce 8-bit, 12-bit, 16-bit, etc. adders.

The Binary Adder

Another common and very useful combinational logic circuit which can be constructed using just a few basic logic gates and adds together binary numbers is the **Binary Adder** circuit. The Binary Adder is made up from standard **AND** and **Ex-OR** gates and allow us to "add" together single bit binary numbers, **a** and **b** to produce two outputs, the **SUM** of the addition and a **CARRY** called the **Carry-out**, (**Cout**) bit. One of the main uses for the **Binary Adder** is in arithmetic and counting circuits.

Consider the addition of two denary (base 10) numbers below.

$$\begin{array}{rcl}
 123 & A & \text{(Augend)} \\
 + 789 & B & \text{(Addend)} \\
 \hline
 912 & \text{SUM} &
 \end{array}$$

Each column is added together starting from the right hand side and each digit has a weighted value depending upon its position in the columns. As each column is added together a carry is generated if the result is greater or equal to ten, the base number. This carry is then added to the result of the addition of the next column to the left and so on, simple school math's addition. The adding of binary numbers is basically the same as that of adding decimal numbers but this time a carry is only generated when the result in any column is greater or equal to "2", the base number of binary.

Binary Addition

Binary Addition follows the same basic rules as for the denary addition above except in binary there are only two digits and the largest digit is "1", so any "SUM" greater than 1 will result in a "CARRY". This carry 1 is passed over to the next column for addition and so on. Consider the single bit addition below.

$$\begin{array}{rclcl}
 0 & 0 & 1 & 1 \\
 + 0 & + 1 & + 0 & + 1 \\
 \hline
 0 & 1 & 1 & 10
 \end{array}$$

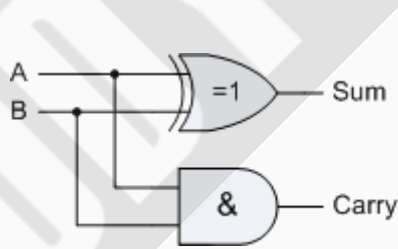
The single bits are added together and "0 + 0", "0 + 1", or "1 + 0" results in a sum of "0" or "1" until you get to "1 + 1" then the sum is equal to "2". For a simple 1-bit addition problem like this, the resulting carry bit could be ignored which would result in an output truth table resembling that of an **Ex-OR Gate** as seen in the Logic Gates section and whose result is the sum of the two bits but without the carry. An **Ex-OR** gate only produces an output "1" when either input is at logic "1", but not both. However, all microprocessors and electronic calculators require the carry bit to correctly calculate the equations so we need to rewrite them to include 2 bits of output data as shown below.

00	00	01	01
<u>+ 00</u>	<u>+ 01</u>	<u>+ 00</u>	<u>+ 01</u>
00	01	01	10

From the above equations we know that an **Ex-OR** gate will only produce an output "1" when "EITHER" input is at logic "1", so we need an additional output to produce a carry output, "1" when "BOTH" inputs "A" and "B" are at logic "1" and a standard **AND Gate** fits the bill nicely. By combining the **Ex-OR** gate with the **AND** gate results in a simple digital binary adder circuit known commonly as the "**Half Adder**" circuit.

The Half Adder Circuit

1-bit Adder with Carry-Out

Symbol	Truth Table			
	A	B	SUM	CARRY
	0	0	0	0
	0	1	1	0
	1	0	1	0
	1	1	0	1
Boolean Expression: Sum = $A \oplus B$ Carry = $A \cdot B$				

From the truth table we can see that the SUM (S) output is the result of the **Ex-OR** gate and the Carry-out (**Cout**) is the result of the **AND** gate. One major disadvantage of the Half Adder circuit when used as a binary adder, is that there is no provision for a "Carry-in" from the previous circuit when adding together multiple data bits. For example, suppose we want to add together two 8-bit bytes of data, any resulting carry bit would need to be able to "ripple" or move across the bit

patterns starting from the least significant bit (LSB). The most complicated operation the half adder can do is "1 + 1" but as the half adder has no carry input the resultant added value would be incorrect. One simple way to overcome this problem is to use a **Full Adder** type binary adder circuit.

The Full Adder Circuit

The main difference between the **Full Adder** and the previous seen **Half Adder** is that a full adder has three inputs, the same two single bit binary inputs **A** and **B** as before plus an additional *Carry-In* (**C-in**) input as shown below.

Full Adder with Carry-In

Symbol	Truth Table				
	A	B	C-in	Sum	C-out
	0	0	0	0	0
	0	1	0	1	0
	1	0	0	1	0
	1	1	0	0	1
	0	0	1	1	0
	0	1	1	0	1
	1	0	1	0	1
	1	1	1	1	1
Boolean Expression: $\text{Sum} = A \oplus B \oplus \text{C-in}$					

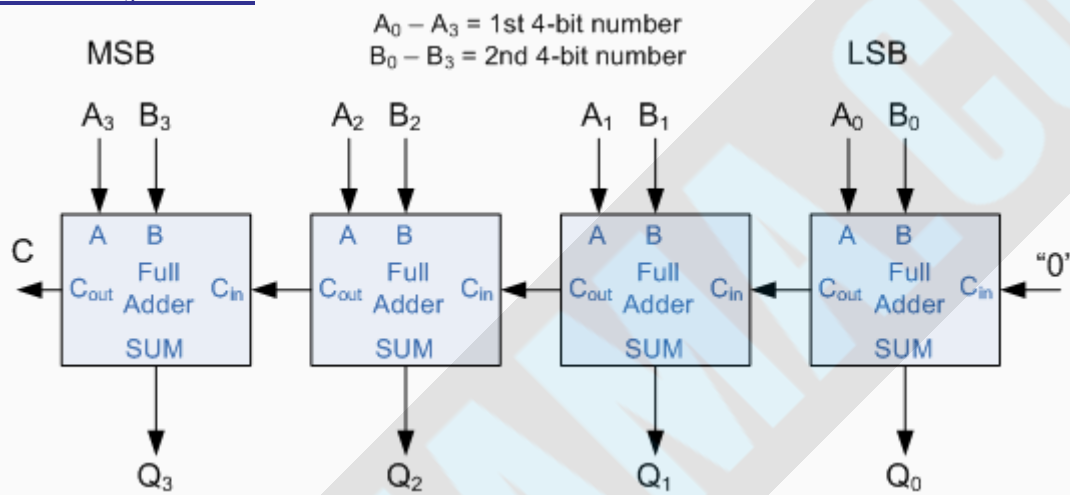
The 1-bit **Full Adder** circuit above is basically two half adders connected together and consists of three **Ex-OR** gates, two **AND** gates and an **OR** gate, six logic gates in total. The truth table for the full adder includes an additional column to take into account the Carry-in input as well as the summed output and carry-output. 4-bit full adder circuits are available as standard IC packages in the form of the TTL 74LS83 or the 74LS283 which can add together two 4-bit binary numbers and generate a **SUM** and a **CARRY** output. But what if we wanted to add together two **n-bit** numbers, then **n** 1-bit full adders need to be connected together to produce what is known as the **Ripple Carry Adder**.

The 4-bit Binary Adder

The **Ripple Carry Binary Adder** is simply **n**, full adders cascaded together with each full adder represents a single weighted column in the long addition with the

carry signals producing a "ripple" effect through the binary adder from right to left. For example, suppose we want to "add" together two 4-bit numbers, the two outputs of the first full adder will provide the first place digit sum of the addition plus a carry-out bit that acts as the carry-in digit of the next binary adder. The second binary adder in the chain also produces a summed output (the 2nd bit) plus another carry-out bit and we can keep adding more full adders to the combination to add larger numbers, linking the carry bit output from the first full binary adder to the next full adder, and so forth. An example of a 4-bit adder is given below.

A 4-bit Binary Adder

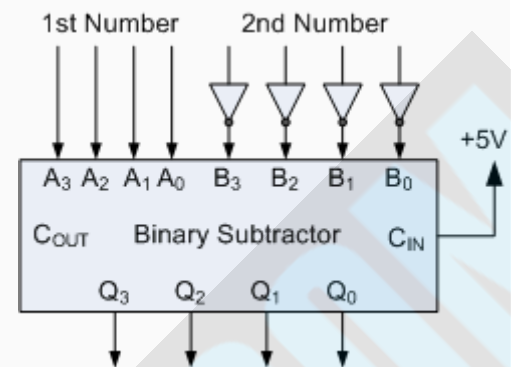


One main disadvantage of "cascading" together 1-bit **binary adders** to add large binary numbers is that if inputs **A** and **B** change, the sum at its output will not be valid until any carry-input has "rippled" through every full adder in the chain. Consequently, there will be a finite delay before the output of an adder responds to a change in its inputs resulting in the accumulated delay especially in large multi-bit binary adders becoming prohibitively large. This delay is called **Propagation delay**. Also "overflow" occurs when an **n-bit** adder adds two numbers together whose sum is greater than or equal to 2^n

One solution is to generate the carry-input signals directly from the **A** and **B** inputs rather than using the ripple arrangement above. This then produces another type of binary adder circuit called a **Carry Look Ahead Binary Adder** where the speed of the parallel adder can be greatly improved using carry-look ahead logic.

The 4-bit Binary Subtractor

Now that we know how to "ADD" together two 4-bit binary numbers how would we subtract two 4-bit binary numbers, for example, $A - B$ using the circuit above. The answer is to use 2's-complement notation on all the bits in B must be complemented (inverted) and an extra one added using the carry-input. This can be achieved by inverting each B input bit using an inverter or NOT-gate.



Also, in the above circuit for the 4-bit binary adder, the first carry-in input is held LOW at logic "0", for the circuit to perform subtraction this input needs to be held HIGH at "1". With this in mind a ripple carry adder can with a small modification be used to perform half subtraction, full subtraction and/or comparison.

There are a number of 4-bit full-adder ICs available such as the 74LS283 and CD4008. which will add two 4-bit binary number and provide an additional input carry bit, as well as an output carry bit, so you can cascade them together to produce 8-bit, 12-bit, 16-bit, etc. adders.

The Digital Comparator

Another common and very useful combinational logic circuit is that of the **Digital Comparator** circuit. Digital or Binary Comparators are made up from standard AND, NOR and NOT gates that compare the digital signals present at their input terminals and produce an output depending upon the condition of those inputs. For example, along with being able to add and subtract binary numbers we need to be able to compare them and determine whether the value of input A is greater than, smaller than or equal to the value at input B etc. The digital comparator accomplishes this using several logic gates that operate on the principles of Boolean algebra. There are two main types of digital comparator available and these are.

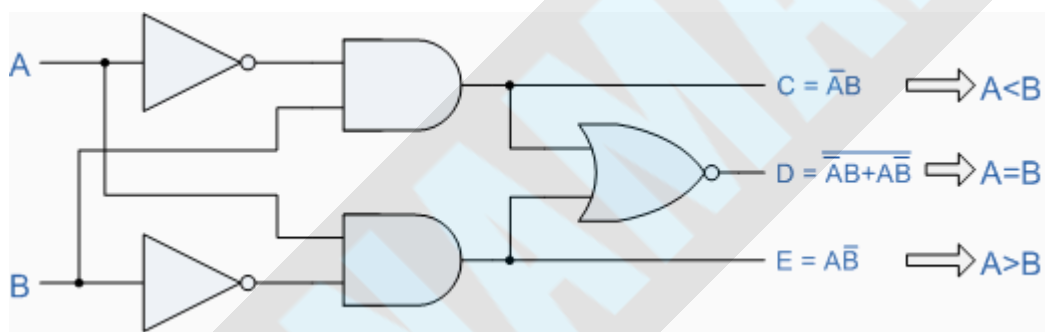
- **Identity Comparator** - is a digital comparator that has only one output terminal for when $A = B$ either "HIGH" $A = B = 1$ or "LOW" $A = B = 0$
- **Magnitude Comparator** - is a type of digital comparator that has three output terminals, one each for equality, $A = B$ greater than, $A > B$ and less than $A < B$

The purpose of a **Digital Comparator** is to compare a set of variables or unknown numbers, for example **A** (A1, A2, A3, An, etc) against that of a constant or unknown value such as **B** (B1, B2, B3, Bn, etc) and produce an output condition or flag depending upon the result of the comparison. For example, a magnitude comparator of two 1-bits, (**A** and **B**) inputs would produce the following three output conditions when compared to each other.

$$A > B, \quad A = B, \quad A < B$$

This is useful if we want to compare two variables and want to produce an output when any of the above three conditions are achieved. For example, produce an output from a counter when a certain count number is reached. Consider the simple 1-bit comparator below.

1-bit Comparator



Then the operation of a 1-bit digital comparator is given in the following Truth Table.

Truth Table

Inputs		Outputs		
B	A	$A > B$	$A = B$	$A < B$
0	0	0	1	0
0	1	1	0	0
1	0	0	0	1
1	1	0	1	0

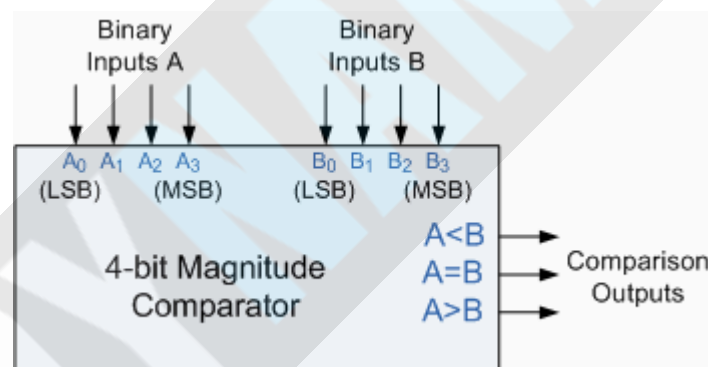
You may notice two distinct features about the comparator from the above truth table. Firstly, the circuit does not distinguish between either two "0" or two "1"s as an output $A = B$ is produced when they are both equal, either $A = B = "0"$ or $A = B = "1"$. Secondly, the output condition for $A = B$ resembles that of a

commonly available logic gate, the **Exclusive-NOR** or **Ex-NOR** function (equivalence) on each of the **n-bits** giving: $Q = A \oplus B$

Digital comparators actually use **Exclusive-NOR** gates within their design for comparing their respective pairs of bits. When we are comparing two binary or BCD values or variables against each other, we are comparing the "magnitude" of these values, a logic "0" against a logic "1" which is where the term **Magnitude Comparator** comes from.

As well as comparing individual bits, we can design larger bit comparators by cascading together **n** of these and produce a **n-bit** comparator just as we did for the **n-bit** adder in the previous tutorial. Multi-bit comparators can be constructed to compare whole binary or BCD words to produce an output if one word is larger, equal to or less than the other. A very good example of this is the 4-bit **Magnitude Comparator**. Here, two 4-bit words ("nibbles") are compared to each other to produce the relevant output with one word connected to inputs **A** and the other to be compared against connected to input **B** as shown below.

4-bit Magnitude Comparator



Some commercially available digital comparators such as the TTL 7485 or CMOS 4063 4-bit magnitude comparator have additional input terminals that allow more individual comparators to be "cascaded" together to compare words larger than 4-bits with magnitude comparators of "n"-bits being produced. These cascading inputs are connected directly to the corresponding outputs of the previous comparator as shown to compare 8, 16 or even 32-bit words.

Parity Generator and Checker:

Parity checkers are integrated circuits (ICs) used in digital systems to detect errors when streams of bits are sent from a transmitter to a receiver. Parity generators calculate the parity of data packets and add a parity amount to them. Both parity checkers and generators use parity memory, a basic form of error

detection which provides an extra bit for every byte stored. Whenever a byte is written to memory, the parity circuit examines the byte and determines whether it contains an even or odd number of ones. If the data byte contains an even number of ones, the extra (parity) bit is set to 1; otherwise, the parity bit is set to 0. When the data is read back from memory, the parity circuit examines all of the bits and determines if there are an odd or even number of ones. An even number of ones indicates that there is an error in one of the bits because a parity circuit, when storing a byte, always sets an error-free parity bit to indicate an odd number of ones. When a parity error is detected, the parity circuit generates a non-maskable interrupt (NMI) that halts the processor, ensuring that the error does not corrupt other data.

There are several important performance specifications for parity checkers and generators. Number of bits is the number of words that devices can handle in parallel. Common configurations are 4, 5, 6, 8, 9, 12, or 16 bits. Supply voltages for parity checkers and generators range from - 5 V to 5 V and include intermediate voltages such as - 4.5 V, - 3.3 V, - 3 V, 1.2 V, 1.5 V, 1.8 V, 2.5 V, 3 V, 3.3 V, and 3.6 V. Propagation delay is the time interval between the occurrence of a change at the output and the application of a change at the inputs. Operating temperature is a full-required range.

Selecting parity checkers and generators requires an analysis of logic families. [Transistor-transistor logic \(TTL\)](#) and related technologies such as Fairchild advanced Schottky TTL (FAST) use transistors as digital switches. By contrast, emitter coupled logic (ECL) uses transistors to steer current through gates that compute logical functions. Another logic family, complementary metal-oxide-semiconductor (CMOS), uses a combination of p-type and n-type metal-oxide-semiconductor field effect transistors (MOSFETs) to implement logic gates and other digital circuits. Bipolar CMOS (BiCMOS) is a silicon-germanium technology that combines the high speed of bipolar TTL with the low power consumption of CMOS. Other logic families for parity checkers and generators include cross-bar switch technology (CBT), gallium arsenide (GaAs), integrated injection logic (I²L) and silicon on sapphire (SOS). Gunning with transceiver logic (GTL) and gunning with transceiver logic plus (GTLP) are also available.

Parity checkers and generators are available in a variety of IC package types and with different numbers of pins. Basic IC package types for ALUs include ball grid array (BGA), quad flat package (QFP), single in-line package (SIP), and dual in-line package (DIP). Many packaging variants are available. For example, BGA variants include plastic-ball grid array (PBGA) and tape-ball grid array (TBGA). QFP variants include low-profile quad flat package (LQFP) and thin quad flat package (TQFP). DIPs are available in either ceramic (CDIP) or plastic (PDIP). Other IC package types include small outline package (SOP), thin small outline package (TSOP), and shrink small outline package (SSOP).

Transistor–transistor logic (TTL) is a class of digital circuits built from bipolar junction transistors (BJT) and resistors. It is called *transistor–transistor logic* because both the logic gating function (e.g., AND) and the amplifying function are performed by transistors (contrast this with RTL and DTL).

TTL is notable for being a widespread integrated circuit (IC) family used in many applications such as computers, industrial controls, test equipment and instrumentation, consumer electronics, synthesizers, etc. The designation *TTL* is sometimes used to mean TTL-compatible logic levels, even when not associated directly with TTL integrated circuits, for example as a label on the inputs and outputs of electronic instruments.

Fundamental TTL gate

TTL is a natural successor of DTL since it is based on the same fundamental concept - implementing the logic gate function by using the base-emitter junctions of a multiple-emitter transistor as switching elements like DTL input diodes. This IC structure is functionally equivalent to multiple transistors where the bases and collectors are tied together. The output of the simple TTL gate is buffered, like DTL, by a common emitter amplifier.

Input logical ones. When all the inputs are held at high voltage the base-emitter junctions of the multiple-emitter transistor are backward-biased. In contrast with DTL, small (about 10 μA) "collector" currents are drawn by the inputs since the transistor is in a reverse-active mode (with swapped collector and emitter). The base resistor in combination with the supply voltage acts as a substantially constant current source. It passes current through the base–collector junction of the multiple-emitter transistor and the base-emitter junction of the output transistor thus turning it on; the output voltage becomes low (logical zero).

Input logical zero. If one input voltage becomes zero, the corresponding base-emitter junction of the multiple-emitter transistor connects in parallel to the two connected in series junctions (the base-collector junction of the multiple-emitter transistor and the base-emitter junction of the second transistor). The input base-emitter junction steers all the base current of the output transistor to the input source (the ground). The base of the output transistor is deprived of current causing it to go into cut-off and the output voltage becomes high (logical one). During the transition the input transistor is briefly in its active region; so it draws a large current away from the base of the output transistor and thus quickly discharges its base. This is a critical advantage of TTL over DTL that speeds up the transition over a diode input structure.

The main disadvantage of TTL with a simple output stage is the relatively high output resistance at output logical "1" that is completely determined by the output collector resistor. It limits the number of inputs that can be connected (the fan out). Some advantage of the simple output stage is the high voltage level (up to V_{CC}) of the output logical "1" when the output is not loaded.

Logic of this type is most frequently encountered with the collector resistor of the output transistor omitted, making an open collector output. This allows the designer to fabricate logic by connecting the open collector outputs of several logic gates together and providing a single external pull-up resistor. If any of the logic gates becomes logic low (transistor conducting), the combined output will be low. Examples of this type of gate are the 7401 and 7403 series.

TTL with a "totem-pole" output stage

To solve the problem with the high output resistance of the simple output stage the second schematic adds to this a "totem-pole" ("push-pull") output. It consists of the two n-p-n transistors V_3 and V_4 , the "lifting" diode V_5 and the current-limiting resistor R_3 (see the figure on the right). It is driven by applying the same *current steering* idea.

When V_2 is "off", V_4 is "off" as well and V_3 operates in active region as a voltage follower producing high output voltage (logical "1"). When V_2 is "on", it activates V_4 , driving low voltage (logical "0") to the output. V_2 and V_4 collector-emitter junctions connect V_4 base-emitter junction in parallel to the series-connected V_3 base-emitter and V_5 anode-cathode junctions. V_3 base current is deprived; the transistor turns "off" and it does not impact on the output. In the middle of the transition, the resistor R_3 limits the current flowing directly through the series connected transistor V_3 , diode V_5 and transistor V_4 that all are conducting. It also limits the output current in the case of output logical "1" and short connection to the ground. The strength of the gate may be increased without proportionally affecting the power consumption by removing the pull-up and pull-down resistors from the output stage.

The main advantage of TTL with a "totem-pole" output stage is the low output resistance at output logical "1". It is determined by the upper output transistor V_3 operating in active region as a voltage follower. The resistor R_3 does not increase the output resistance since it is connected in the V_3 collector and its influence is compensated by the negative feedback. A disadvantage of the "totem-pole" output stage is the decreased voltage level (no more than 3.5 V) of the output logical "1" (even, if the output is unloaded). The reason of this reduction are the voltage drops across the V_3 base-emitter and V_5 anode-cathode junctions.

Comparison with other logic families

TTL devices consume substantially more power than equivalent CMOS devices at rest, but power consumption does not increase with clock speed as rapidly as for CMOS devices.^[21] Compared to contemporary ECL circuits, TTL uses less power and has easier design rules but is substantially slower. Designers can combine ECL and TTL devices in the same system to achieve best overall performance and economy, but level-shifting devices are required between the two logic families. TTL is less sensitive to damage from electrostatic discharge than early CMOS devices.

Due to the output structure of TTL devices, the output impedance is asymmetrical between the high and low state, making them unsuitable for driving transmission lines. This drawback is usually overcome by buffering the outputs with special line-driver devices where signals need to be sent through cables. ECL, by virtue of its symmetric low-impedance output structure, does not have this drawback.

The TTL "totem-pole" output structure often has a momentary overlap when both the upper and lower transistors are conducting, resulting in a substantial pulse of current drawn from the supply. These pulses can couple in unexpected ways between multiple integrated circuit packages, resulting in reduced noise margin and lower performance. TTL systems usually have a decoupling capacitor for every one or two IC packages, so that a current pulse from one chip does not momentarily reduce the supply voltage to the others.

Several manufacturers now supply CMOS logic equivalents with TTL-compatible input and output levels, usually bearing part numbers similar to the equivalent TTL component and with the same pinouts. For example, the 74HCT00 series provides many drop-in replacements for bipolar 7400 series parts, but uses CMOS technology.

Sub-types

Successive generations of technology produced compatible parts with improved power consumption or switching speed, or both. Although vendors uniformly marketed these various product lines as TTL with Schottky diodes, some of the underlying circuits, such as used in the LS family, could rather be considered DTL.

Variations of and successors to the basic TTL family, which has a typical gate propagation delay of 10ns and a power dissipation of 10 mW per gate, for a power-delay product (PDP) or switching energy of about 100 pJ, include:

- Low-power TTL (L), which traded switching speed (33ns) for a reduction in power consumption (1 mW) (now essentially replaced by CMOS logic)

- High-speed TTL (H), with faster switching than standard TTL (6ns) but significantly higher power dissipation (22 mW)
- Schottky TTL (S), introduced in 1969, which used Schottky diode clamps at gate inputs to prevent charge storage and improve switching time. These gates operated more quickly (3ns) but had higher power dissipation (19 mW)
- Low-power Schottky TTL (LS) — used the higher resistance values of low-power TTL and the Schottky diodes to provide a good combination of speed (9.5ns) and reduced power consumption (2 mW), and PDP of about 20 pJ. Probably the most common type of TTL, these were used as glue logic in microcomputers, essentially replacing the former H, L, and S sub-families.
- Fast (F) and Advanced-Schottky (AS) variants of LS from Fairchild and TI, respectively, circa 1985, with "Miller-killer" circuits to speed up the low-to-high transition. These families achieved PDPs of 10 pJ and 4 pJ, respectively, the lowest of all the TTL families.
- Most manufacturers offer commercial and extended temperature ranges: for example Texas Instruments 7400 series parts are rated from 0 to 70°C, and 5400 series devices over the military-specification temperature range of -55 to +125°C.
- Radiation-hardened devices are offered for space applications
- Special quality levels and high-reliability parts are available for military and aerospace applications.
- Low-voltage TTL (LVTTTL) for 3.3-volt power supplies and memory interfacing.

Applications

Before the advent of VLSI devices, TTL integrated circuits were a standard method of construction for the processors of mini-computer and mainframe processors; such as the DEC VAX and Data General Eclipse, and for equipment such as machine tool numerical controls, printers and video display terminals. As microprocessors became more functional, TTL devices became important for "glue logic" applications, such as fast bus drivers on a motherboard, which tie together the function blocks realized in VLSI elements.

CMOS

Complementary metal–oxide–semiconductor (CMOS) (pronounced /'si:mɒs/) is a technology for constructing [integrated circuits](#). CMOS technology is used in [microprocessors](#), [microcontrollers](#), [static RAM](#), and other [digital logic](#) circuits. CMOS technology is also used for several analog circuits such as [image sensors](#), [data converters](#), and highly integrated [transceivers](#) for many types of communication. [Frank Wanlass](#) patented CMOS in 1967 ([US patent 3,356,858](#)).

CMOS is also sometimes referred to as **complementary-symmetry metal–oxide–semiconductor** (or COS-MOS). The words "complementary-symmetry" refer to the fact that the typical digital design style with CMOS uses complementary and symmetrical pairs of [p-type](#) and [n-type metal oxide semiconductor field effect transistors](#) (MOSFETs) for logic functions.

Two important characteristics of CMOS devices are high [noise immunity](#) and low static [power consumption](#). Significant power is only drawn while the [transistors](#) in the CMOS device are switching between on and off states. Consequently, CMOS devices do not produce as much [waste heat](#) as other forms of logic, for example [transistor-transistor logic](#) (TTL) or [NMOS logic](#), which uses all n-channel devices without p-channel devices. CMOS also allows a high density of logic functions on a chip. It was primarily this reason why CMOS won the race in the eighties and became the most used technology to be implemented in [VLSI](#) chips.

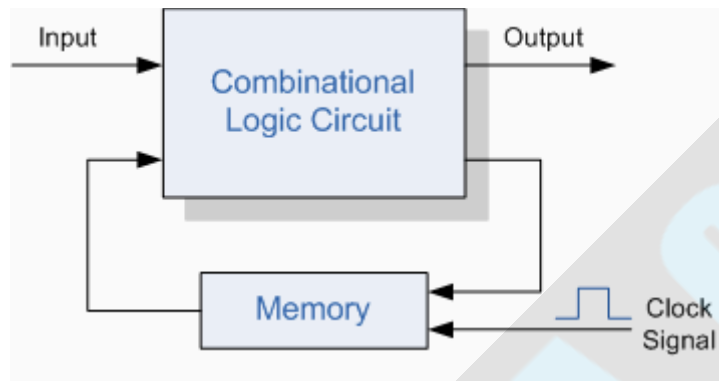
The phrase "metal–oxide–semiconductor" is a reference to the physical structure of certain [field-effect transistors](#), having a metal gate electrode placed on top of an oxide insulator, which in turn is on top of a [semiconductor material](#). [Aluminum](#) was once used but now the material is [polysilicon](#). Other [metal gates](#) have made a comeback with the advent of [high-k](#) dielectric materials in the CMOS process, as announced by IBM and Intel for the [45 nanometer](#) node and beyond.

Sequential Logic Basics

Unlike [Combinational Logic](#) circuits that change state depending upon the actual signals being applied to their inputs at that time, **Sequential Logic** circuits have some form of inherent "**Memory**" built in to them and they are able to take into account their previous input state as well as those actually present, a sort of "before" and "after" is involved. They are generally termed as **Two State** or [Bistable](#) devices which can have their output set in either of two

basic states, a logic level "1" or a logic level "0" and will remain "latched" indefinitely in this current state or condition until some other input trigger pulse or signal is applied which will cause it to change its state once again.

Sequential Logic Circuit

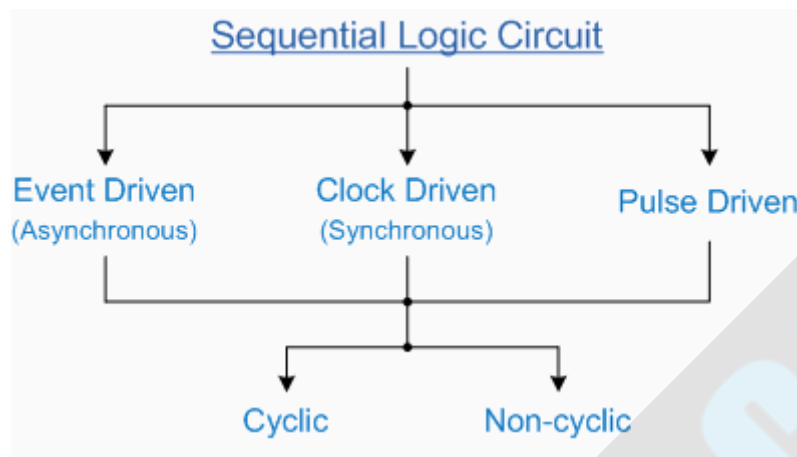


The word "Sequential" means that things happen in a "sequence", one after another and in **Sequential Logic** circuits, the actual clock signal determines when things will happen next. Simple sequential logic circuits can be constructed from standard **Bistable** circuits such as Flip-flops, Latches or Counters and which themselves can be made by simply connecting together [NAND Gates](#) and/or [NOR Gates](#) in a particular combinational way to produce the required sequential circuit.

Sequential Logic circuits can be divided into 3 main categories:

- 1. Clock Driven - Synchronous Circuits that are Synchronised to a specific clock signal.
- 2. Event Driven - Asynchronous Circuits that react or change state when an external event occurs.
- 3. Pulse Driven - Which is a Combination of Synchronous and Asynchronous.

Classification of Sequential Logic



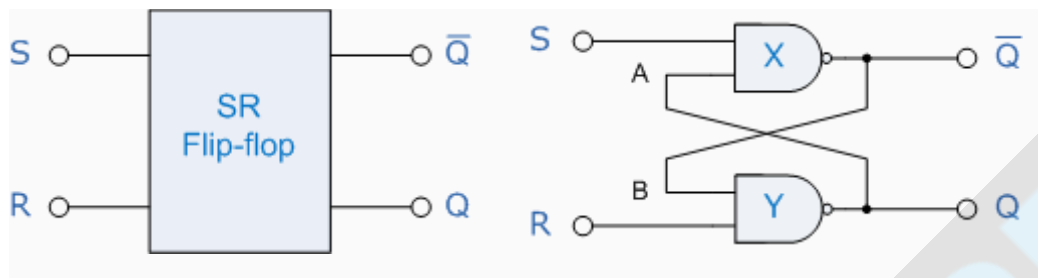
As well as the two logic states mentioned above logic level "1" and logic level "0", a third element is introduced that separates **Sequential Logic** circuits from their **Combinational Logic** counterparts, namely **TIME**. Sequential logic circuits that return back to their original state once reset, i.e. circuits with loops or feedback paths are said to be "Cyclic" in nature.

SR Flip-Flop

An **SR Flip-Flop** can be considered as a basic one-bit memory device that has two inputs, one which will "SET" the device and another which will "RESET" the device back to its original state and an output Q that will be either at a logic level "1" or logic "0" depending upon this Set/Reset condition. A basic NAND Gate SR flip flop circuit provides feedback from its outputs to its inputs and is commonly used in memory circuits to store data bits. The term "Flip-flop" relates to the actual operation of the device, as it can be "Flipped" into one logic state or "Flopped" back into another.

The simplest way to make any basic one-bit Set/Reset SR flip-flop is to connect together a pair of cross-coupled 2-input NAND Gates to form a Set-Reset Bistable or a SR NAND Gate Latch, so that there is feedback from each output to one of the other NAND Gate inputs. This device consists of two inputs, one called the Reset, R and the other called the Set, S with two corresponding outputs Q and its inverse or complement Q as shown below.

The SR NAND Gate Latch



The Set State

Consider the circuit shown above. If the input R is at logic level "0" ($R = 0$) and input S is at logic level "1" ($S = 1$), the NAND Gate Y has at least one of its inputs at logic "0" therefore, its output Q must be at a logic level "1" (NAND Gate principles). Output Q is also fed back to input A and so both inputs to the NAND Gate X are at logic level "1", and therefore its output $\bar{Q}$ must be at logic level "0". Again NAND gate principles. If the Reset input R changes state, and now becomes logic "1" with S remaining HIGH at logic level "1", NAND Gate Y inputs are now $R = "1"$ and $B = "0"$ and since one of its inputs is still at logic level "0" the output at Q remains at logic level "1" and the circuit is said to be "Latched" or "Set" with $Q = "1"$ and $\bar{Q} = "0"$.

Reset State

In this second stable state, Q is at logic level "0", $\bar{Q} = "1"$ its inverse output Q is at logic level "1", not $\bar{Q} = "1"$, and is given by $R = "1"$ and $S = "0"$. As gate X has one of its inputs at logic "0" its output $\bar{Q}$ must equal logic level "1" (again NAND gate principles). Output $\bar{Q}$ is fed back to input B, so both inputs to NAND gate Y are at logic "1", therefore, $Q = "0"$. If the set input, S now changes state to logic "1" with R remaining at logic "1", output Q still remains LOW at logic level "0" and the circuit's "Reset" state has been latched.

Truth Table for this Set-Reset Function

State	S	R	Q	$\bar{Q}$
Set	1	0	1	0
	1	1	1	0
Reset	0	1	0	1

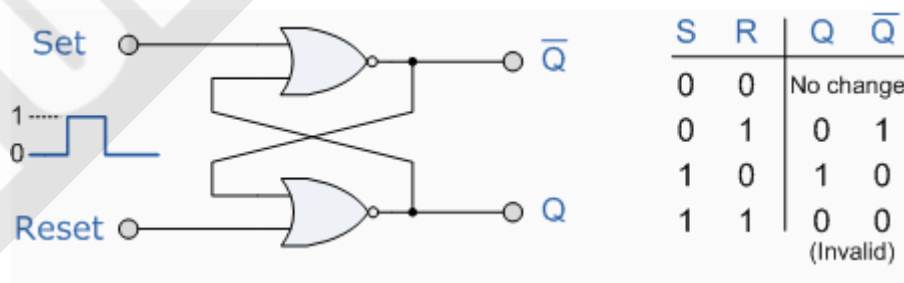
	1	1	0	1
Invalid	0	0	1	1

It can be seen that when both inputs $S = "1"$ and $R = "1"$ the outputs Q and $\bar{Q}$ can be at either logic level "1" or "0", depending upon the state of inputs S or R BEFORE this input condition existed. However, input state $R = "0"$ and $S = "0"$ is an undesirable or invalid condition and must be avoided because this will give both outputs Q and $\bar{Q}$ to be at logic level "1" at the same time and we would normally want Q to be the inverse of $\bar{Q}$. However, if the two inputs are now switched HIGH again after this condition to logic "1", both the outputs will go LOW resulting in the flip-flop becoming unstable and switch to an unknown data state based upon the unbalance. This unbalance can cause one of the outputs to switch faster than the other resulting in the flip-flop switching to one state or the other which may not be the required state and data corruption will exist. This unstable condition is known as its **Meta-stable** state.

Then, a bistable latch is activated or Set by a logic "1" applied to its S input and deactivated or Reset by a logic "1" applied to its R . The SR Latch is said to be in an "invalid" condition (Meta-stable) if both the Set and Reset inputs are activated simultaneously.

As well as using NAND Gates, it is also possible to construct simple 1-bit **SR Flip-flops** using two NOR Gates connected the same configuration. The circuit will work in a similar way to the NAND gate circuit above, except that the invalid condition exists when both its inputs are at logic level "1" and this is shown below.

The NOR Gate SR Flip-flop



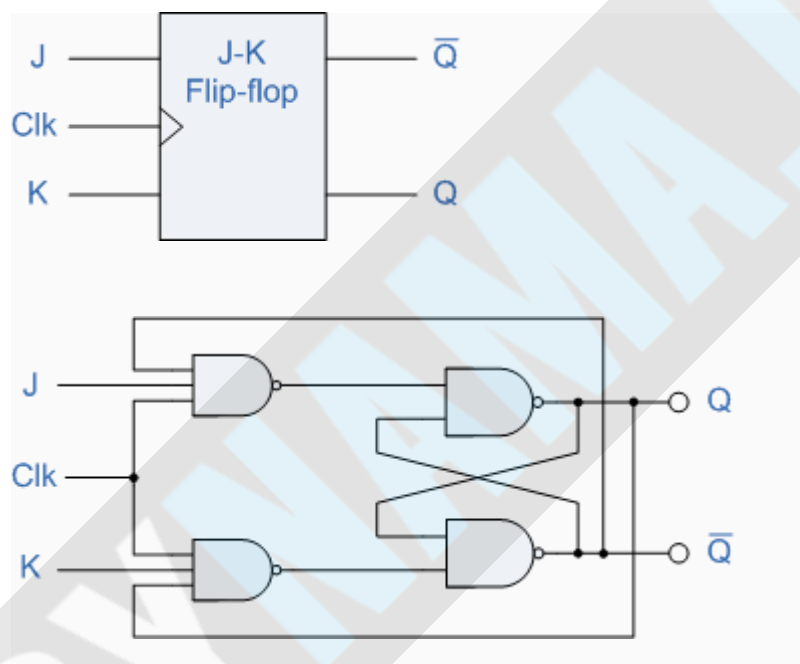
The JK Flip-Flop

From the previous tutorial we now know that the basic gated **SR NAND Flip-flop** suffers from two basic problems: Number 1, the $S = 0$ and $R = 0$ condition

or $S = R = 0$ must always be avoided, and number 2, if S or R change state while the enable input is high the correct latching action will not occur. Then to overcome these two problems the **JK Flip-Flop** was developed.

The **JK Flip-Flop** is basically a Gated SR Flip-Flop with the addition of clock input circuitry that prevents the illegal or invalid output that can occur when both input S equals logic level "1" and input R equals logic level "1". The symbol for a JK Flip-flop is similar to that of an [SR Bistable](#) as seen in the previous tutorial except for the addition of a clock input.

The JK Flip-flop



Both the S and the R inputs of the previous SR bistable have now been replaced by two inputs called the J and K inputs, respectively. The two 2-input NAND gates of the gated SR bistable have now been replaced by two 3-input AND gates with the third input of each gate connected to the outputs Q and $\bar{Q}$. This cross coupling of the SR Flip-flop allows the previously invalid condition of $S = "1"$ and $R = "1"$ state to be usefully used to turn it into a "Toggle action" as the two inputs are now interlocked. If the circuit is "**Set**" the J input is inhibited by the "0" status of the Q through the lower AND gate. If the circuit is "**Reset**" the K input is inhibited by the "0" status of Q through the upper AND gate. When both inputs J and K are equal to logic "1", the JK flip-flop changes state and the truth table for this is given below.

The Truth Table for the JK Function

J	K	Q	Q	
0	0	0	0	same as for the SR Latch
0	0	1	1	
0	1	0	0	
0	1	1	0	
1	0	0	1	
1	0	1	1	
1	1	0	1	toggle action
1	1	1	0	

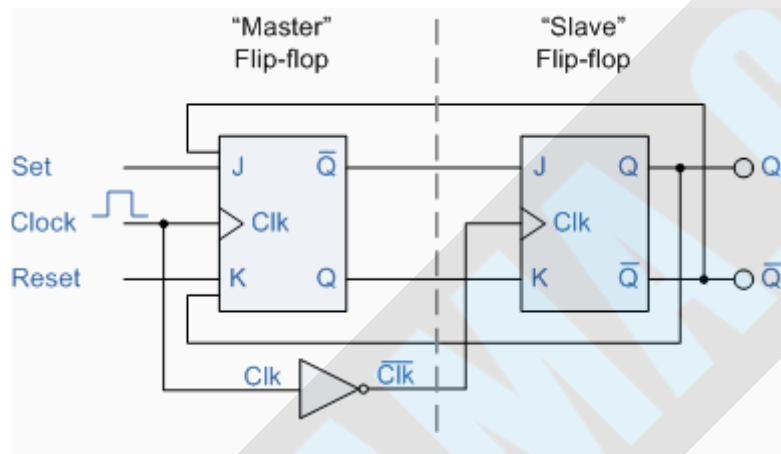
Then the JK Flip-flop is basically an SR Flip-flop with feedback and which enables only one of its two input terminals, either Set or Reset at any one time thereby eliminating the invalid condition seen previously in the SR Flip-flop circuit. Also when both the J and the K inputs are at logic level "1" at the same time, and the clock input is pulsed either "HIGH" or "LOW" the circuit will "Toggle" from a Set state to a Reset state, or visa-versa. This results in the JK Flip-flop acting more like a **T-type Flip-flop** when both terminals are "HIGH".

Although this circuit is an improvement on the clocked SR flip-flop it still suffers from timing problems called "race" if the output Q changes state before the timing pulse of the clock input has time to go "OFF". To avoid this the timing pulse period (T) must be kept as short as possible (high frequency). As this is sometimes is not possible with modern TTL IC's the much improved **Master-Slave JK Flip-flop** was developed. This eliminates all the timing problems by using two SR flip-flops connected together in series, one for the "Master" circuit, which triggers on the leading edge of the clock pulse and the other, the "Slave" circuit, which triggers on the falling edge of the clock pulse.

Master-Slave JK Flip-flop

The **Master-Slave Flip-Flop** is basically two JK bistable flip-flops connected together in a series configuration with the outputs from Q and $\bar{Q}$ from the "Slave" flip-flop being fed back to the inputs of the "Master" with the outputs of the "Master" flip-flop being connected to the two inputs of the "Slave" flip-flop as shown below.

Master-Slave JK Flip-Flops



The input signals J and K are connected to the "Master" flip-flop which "locks" the input while the clock (Clk) input is high at logic level "1". As the clock input of the "Slave" flip-flop is the inverse (complement) of the "Master" clock input, the outputs from the "Master" flip-flop are only "seen" by the "Slave" flip-flop when the clock input goes "LOW" to logic level "0". Therefore on the "High-to-Low" transition of the clock pulse the locked outputs of the "Master" flip-flop are fed through to the JK inputs of the "Slave" flip-flop making this type of flip-flop edge or pulse-triggered.

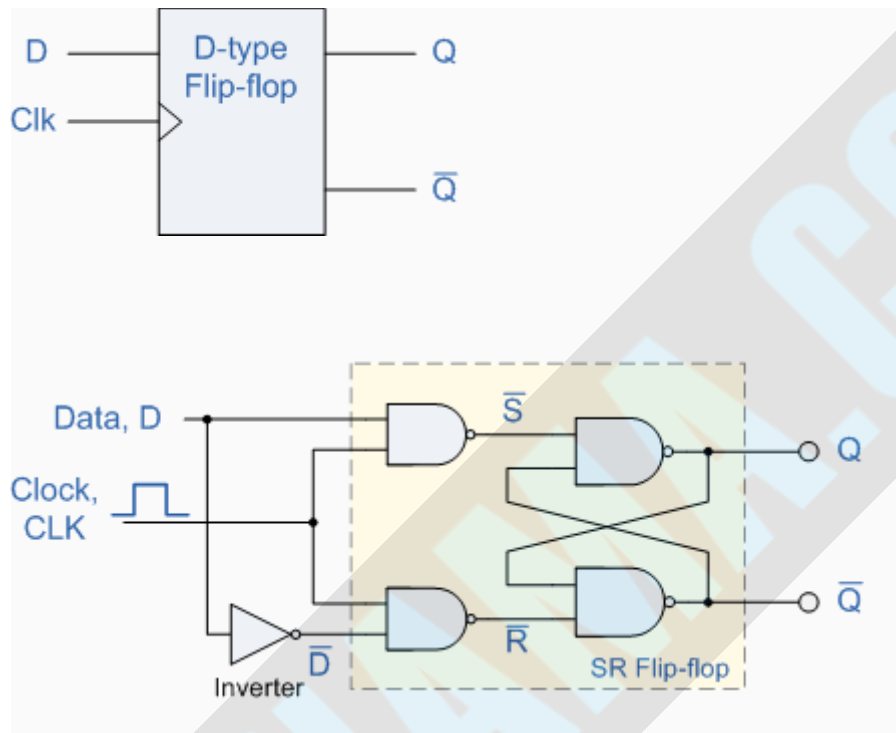
Then, the circuit accepts input data when the clock signal is "HIGH", and passes the data to the output on the falling-edge of the clock signal. In other words, the **Master-Slave JK Flip-flop** is a "Synchronous" device as it only passes data with the timing of the clock signal.

Data Latch

One of the main disadvantages of the basic [SR NAND Gate](#) Bistable circuit is that the indeterminate input condition of "SET" = logic "0" and "RESET" = logic "0" is forbidden. That state will force both outputs to be at logic "1", overriding the feedback latching action and whichever input goes to logic level "1" first will lose control, while the other input still at logic "0" controls the

resulting state of the latch. In order to prevent this from happening an inverter can be connected between the "SET" and the "RESET" inputs to produce a **D-Type Data Latch** or simply **Data Latch** as it is generally called.

Data Latch Circuit



We remember that the simple SR flip-flop requires two inputs, one to "SET" the output and one to "RESET" the output. By connecting an inverter (NOT gate) to the SR flip-flop we can "SET" and "RESET" the flip-flop using just one input as now the two latch inputs are complements of each other. This single input is called the "DATA" input. If this data input is HIGH the flip-flop would be "SET" and when it is LOW the flip-flop would be "RESET". However, this would be rather pointless since the flip-flop's output would always change on every data input. To avoid this an additional input called the "CLOCK" or "ENABLE" input is used to isolate the data input from the flip-flop after the desired data has been stored. This then forms the basis of a **Data Latch** or "D-Type latch".

The **D-Type Latch** will store and output whatever logic level is applied to its data terminal so long as the clock input is high. Once the clock input goes low the SET and RESET inputs of the flip-flop are both held at logic level "1" so it will not change state and store whatever data was present on its output before the clock transition occurred. In other words the output is "latched" at either logic "0" or logic "1".

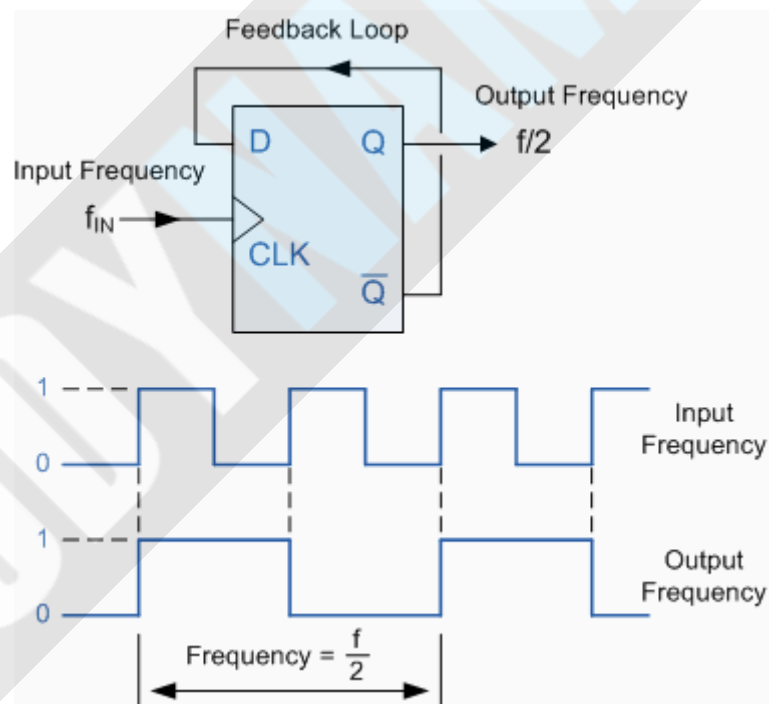
Truth Table for the D-type Flip-flop

Clk	D	Q	Q	OUTPUT
0	x	Q	Q	HOLD
1	0	0	1	RESET
1	1	1	0	SET

Frequency Division

One main use of a Data Latch is as a [Frequency Divider](#). In the Counters tutorials we saw how the **Data Latch** can be used as a "Binary Divider", or a "Frequency Divider" to produce a "divide-by-2" counter. Here the inverted output terminal Q (NOT-Q) is connected directly back to the Data input terminal D giving the device "feedback" as shown below.

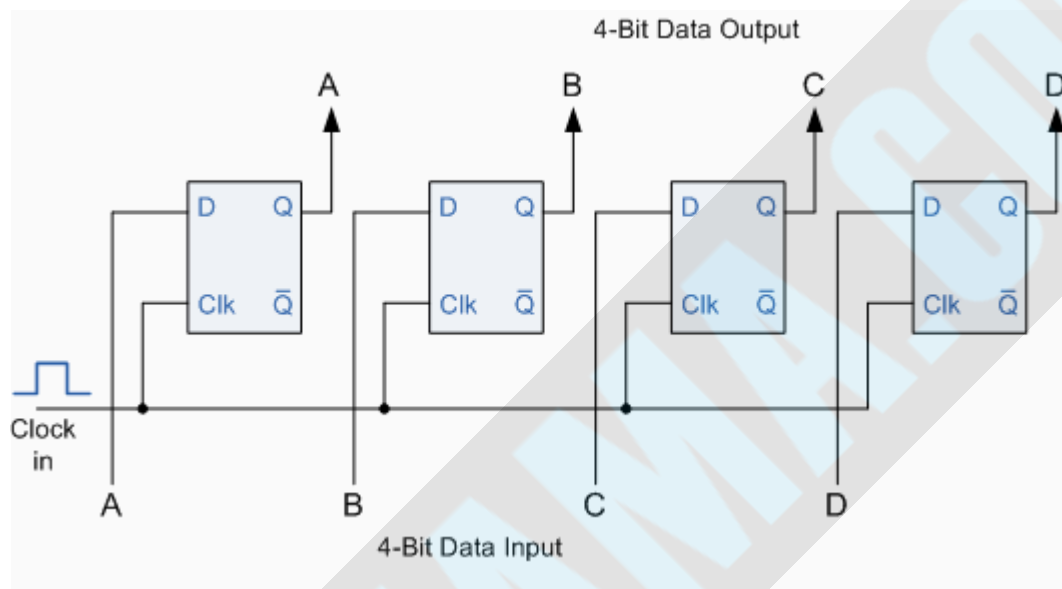
Divide-by-2 Counter



It can be seen from the frequency waveforms above, that by "feeding back" the output from Q to the input terminal D, the output pulses at Q have a frequency that are exactly one half ($f/2$) that of the input clock frequency, (F_{in}). In other words the circuit produces **Frequency Division** as it now divides the input frequency by a factor of two (an octave).

Another use of a Data Latch is to hold or remember its data, thereby acting as a single bit memory cell and IC's such as the TTL 74LS74 or the CMOS 4042 are available in Quad format for this purpose. By connecting together four, **1-bit** latches so that all their clock terminals are connected at the same time a simple "4-bit" Data latch can be made as shown below.

4-bit Data Latch

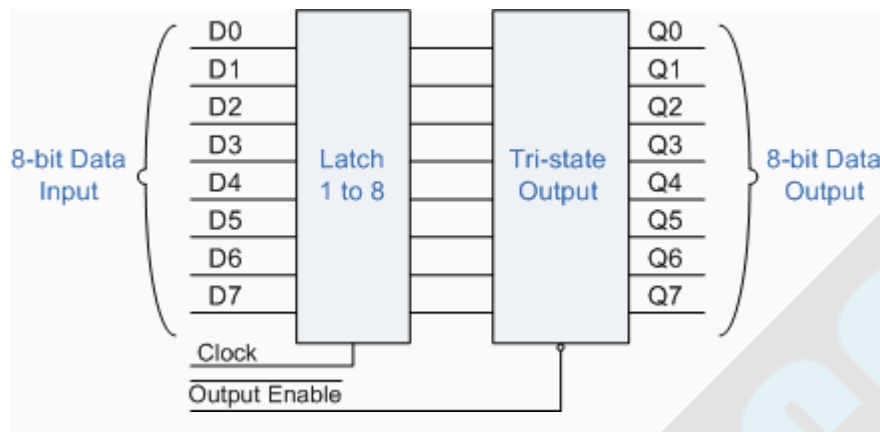


Transparent Data Latch

The **Data Latch** is a very useful devices in electronic and computer circuits. They can be designed to have very high output impedance at both outputs Q and its inverse $\bar{Q}$ to reduce the impedance effect on the connecting circuit when used as buffers, I/O ports, bi-directional bus drivers or even display drivers. But a single "1-bit" data latch is not very practical to use on its own and instead commercially available IC's incorporate 4, 8, 10, 16 or even 32 individual data latches into one single IC package, and one such IC device is the 74LS373 Octal D-type transparent latch.

The eight individual Data Latches of the 74LS373 are "transparent" D-type latches, meaning that when the clock (CLK) input is HIGH at logic level "1", the Q outputs follow the data D inputs and the latch appears "transparent" as the data flows through it. When the clock signal is LOW at logic level "0", the output is latched at the level of the data that was present before the clock input changed.

8-bit Data Latch



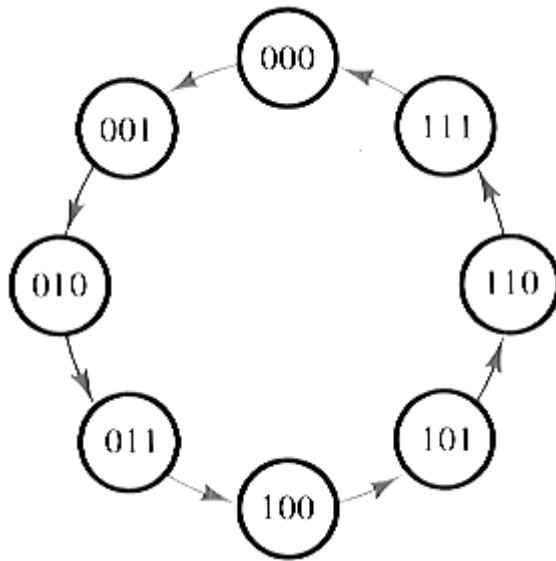
Functional diagram of the 74LS373 Octal Transparent Latch

synchronous (1) Pertaining to two or more processes that depend upon the occurrence of specific events such as common timing signals. (2) Occurring with a regular or predictable time relationship.

counter (1) A functional unit with a finite number of states each of which represents a number that can be, upon receipt of an appropriate signal, increased by unity or by a given constant. This device is usually capable of bringing the represented number to a specified value; for example zero.

So a "synchronous counter" is actually a functional unit with a certain number of states, each representing a number which can be increased or decreased upon receiving an appropriate signal (e.g. a rising edge pulse), and is usually used to count to, or count down to zero from, a specified number N.

.Basically, any sequential circuit that goes through a prescribed sequence of states upon the application of input pulses is called a counter. The input pulses, called count pulses, may be clock pulses or they may originate from an external source and may occur at prescribed intervals of time or at random. The sequence of states in a counter may follow a binary count or any other sequence.



Why do we need counters?

In a digital circuit, counters are used to do 3 main functions: timing, sequencing and counting.

A timing problem might require that a high-frequency pulse train, such as the output of a 10-MHz crystal oscillator, be divided to produce a pulse train of a much lower frequency, say 1 Hz. This application is required in a precision digital clock, where it is not possible to build a crystal oscillator whose natural frequency is 1 Hz.

A sequencing problem would arise if, for instance, it became necessary to apply power to various components of a large machine in a specific order. The starting of a rocket motor is an example where the energizing of fuel pumps, ignition, and possibly explosive bolts for staging must follow a critical order.

Measuring the flow of auto traffic on roadway is an application in which an event (the passage of a vehicle) must increment a tally. This can be done automatically with an electronic counter triggered by a photocell or road sensor. In this way, the total number of vehicles passing a certain point can be counted.

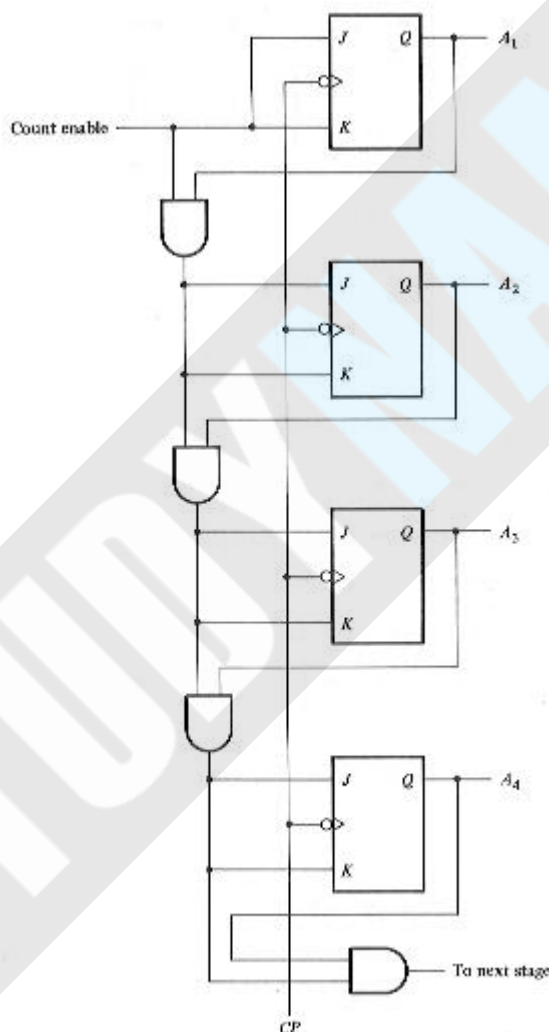
How are counters made?

Counters are generally made up of flip-flops and logic gates. Like flip-flops, counters can retain an output state after the input condition which brought about that state has been removed. Consequently, digital counters are classified as sequential circuits. While a flip-flop can occupy one of only two possible states, a counter can have many more than two states. In the case of a counter, the value of a state is expressed as a multidigit binary number, whose '1's and '0's are usually derived from the outputs of internal flip-flops that make up the counter. The number of states a counter may have is limited only by the amount of electronic hardware

that is available. The main types of flip-flops used are J-K flip-flops or T flip-flops, which are J-K flip-flops with both J and K inputs tied together. Before that, here's a quick reminder of how a J-K flip-flop works:

J input	K input	Output, Q
0	0	Q
0	1	0
1	0	1
1	1	not Q

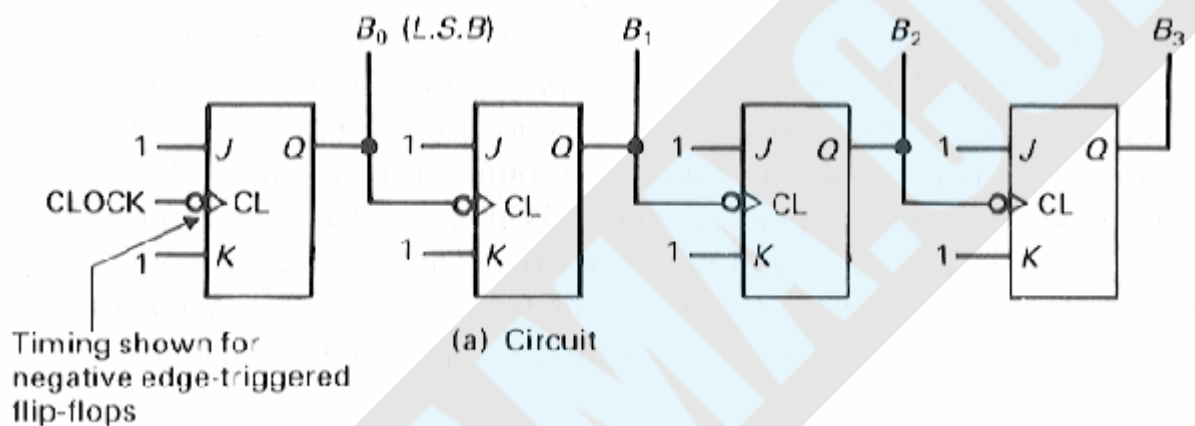
T flip-flops are used because set/reset ([1,0] [0,1]) functions are seldom used. Only the "do nothing" and toggle ([0,0] [1,1]) functions are used. Logic gates are used to decide when to toggle which outputs. Below is an example of a synchronous binary counter, implemented using J-K flip-flops and AND gates.



Why "synchronous"?

The difference between asynchronous and synchronous counters.

In an asynchronous counter, an external event is used to directly SET or CLEAR a flip-flop when it occurs. In a synchronous counter however, the external event is used to produce a pulse that is synchronised with the internal clock. An example of an asynchronous counter is a ripple counter. Each flip-flop in the ripple counter is clocked by the output from the previous flip-flop. Only the first flip-flop is clocked by an external clock. Below is an example of a 4-bit ripple counter:



So what's wrong with asynchronous counters?

☹️Dangers of asynchronous counters.

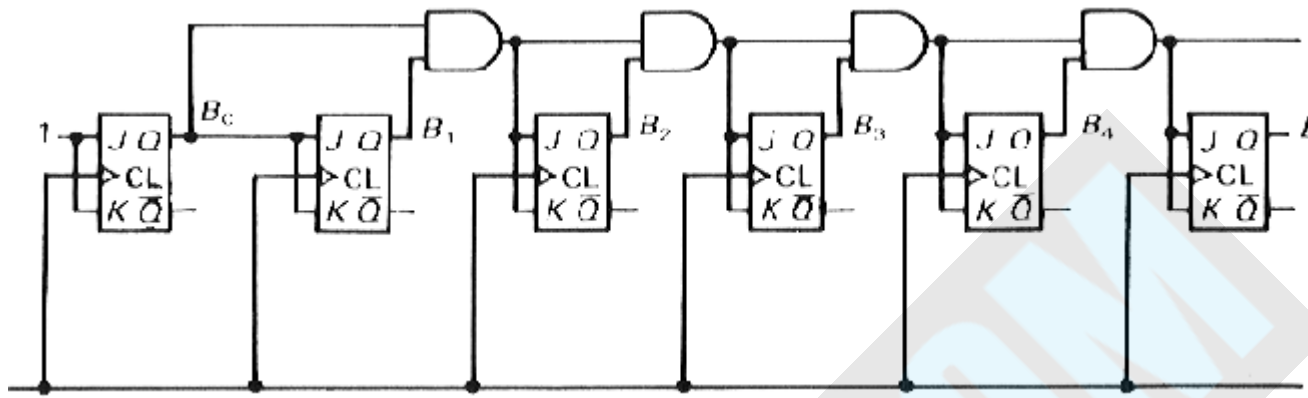
Although the asynchronous counter is easier to implement, it is more "dangerous" than the synchronous counter. In a complex system, there are many state changes on each clock edge, and some IC's (integrated circuits) respond faster than others. If an external event is allowed to affect a system whenever it occurs, a small percentage of the time it will occur near a clock transition, after some IC's have responded, but before others have. This intermingling of transitions often causes erroneous operations. What is worse, these problems are difficult to test for and difficult to foresee because of the random time difference between the events.



Different types of synchronous counters

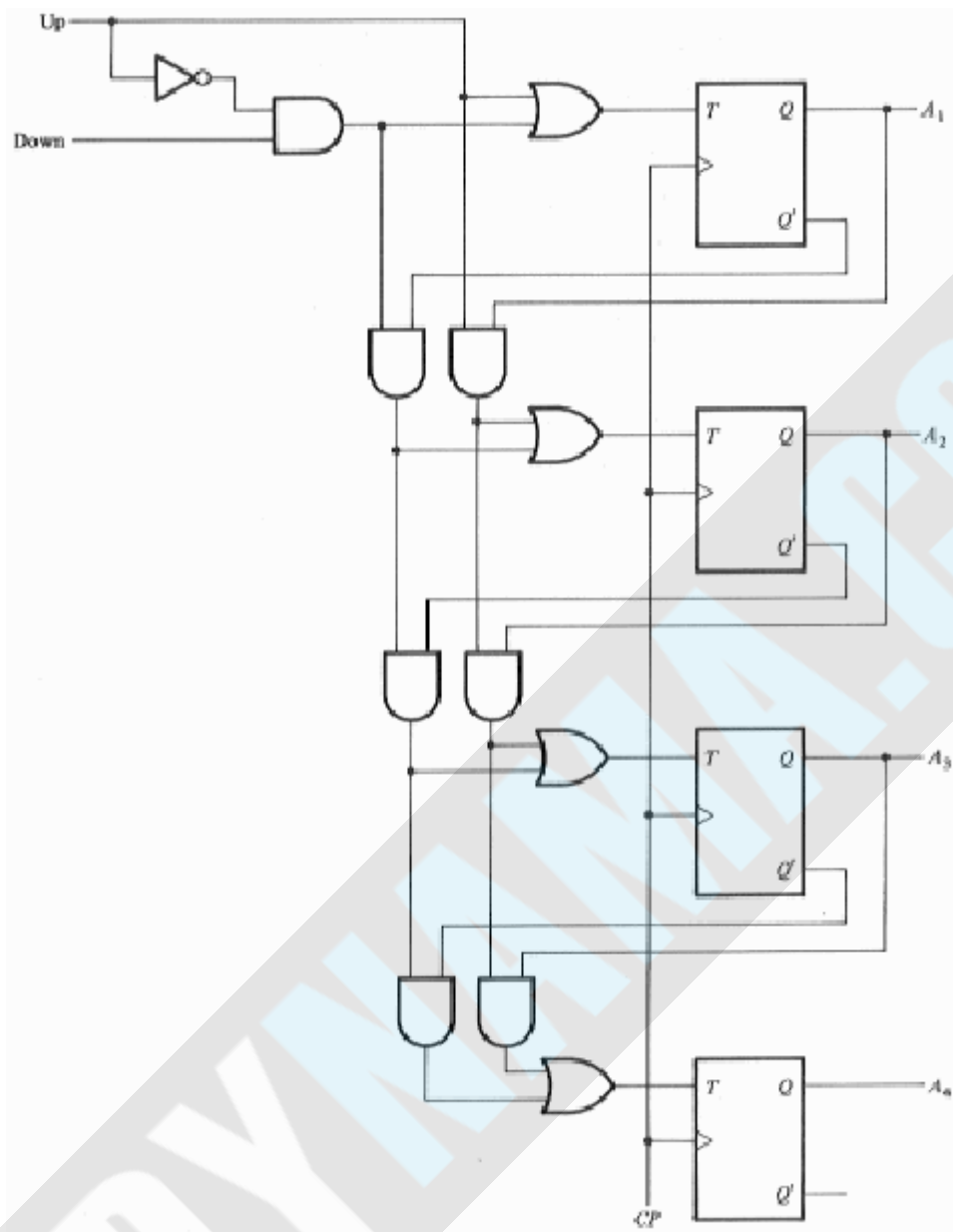
☹️Binary counters.

Binary counters are the simplest form of counters. An N-bit binary counter counts from 0 to $(2^N - 1)$ and back to 0 again.



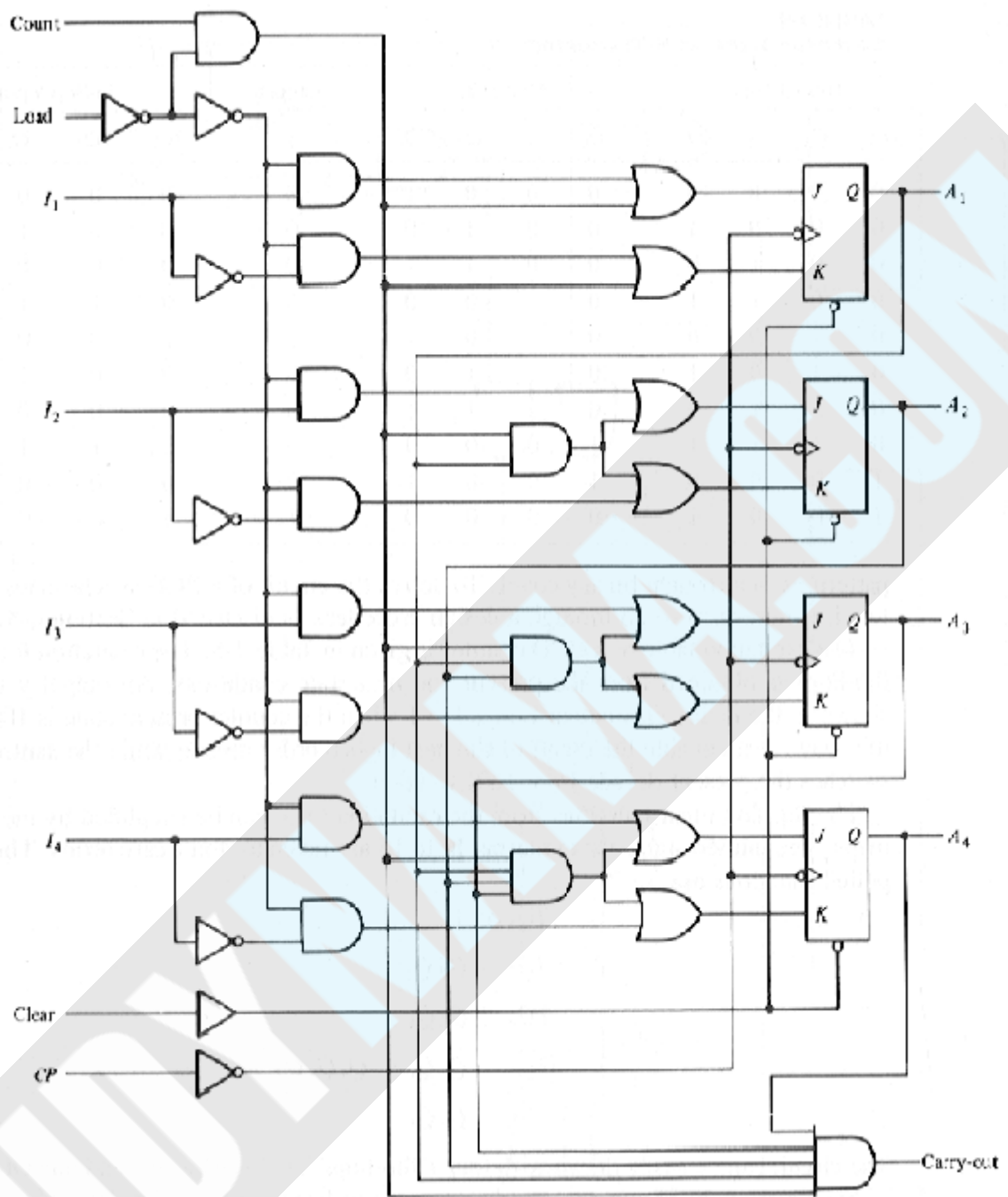
😊Up/down counters.

Instead of just counting up (up counter), counters can be made to count down (down counter) or both up and down (up-down counter). The diagram below shows an up-down counter. The counter counts up or down depending on which of the "up" and "down" inputs are high.



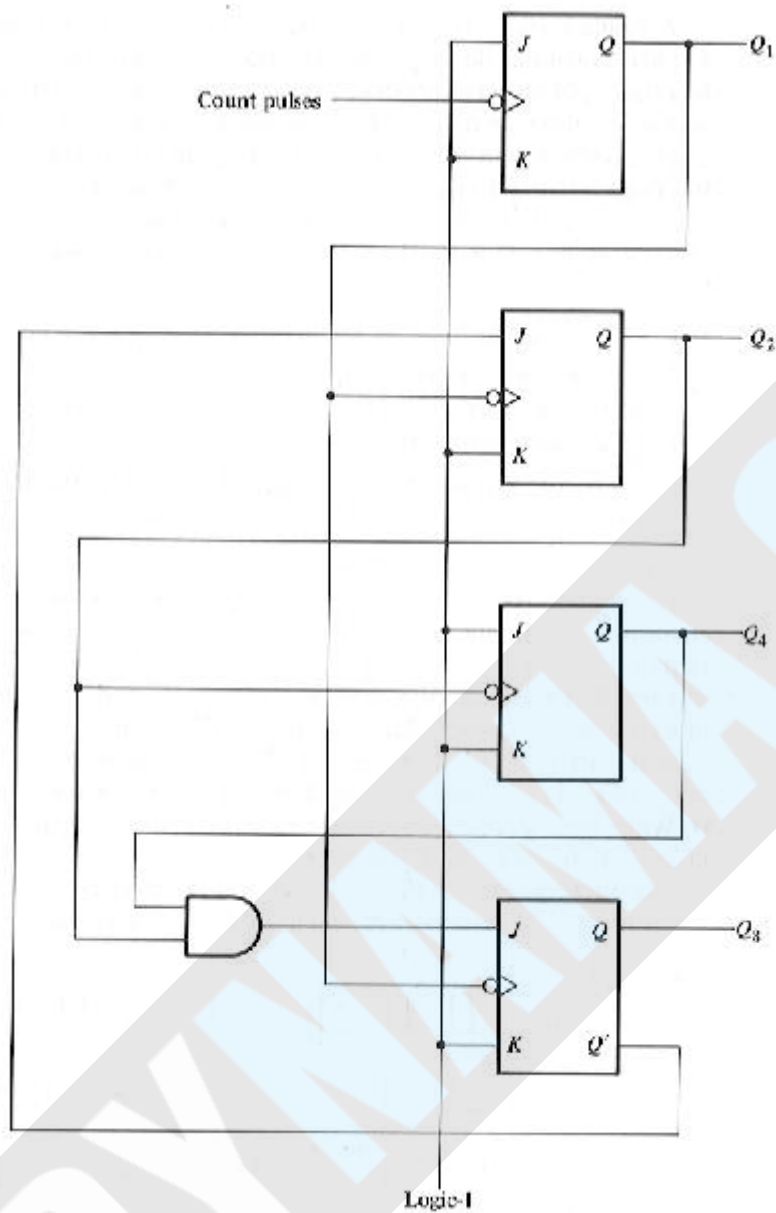
😊 Loadable counters.

And instead of counting from 0, a counter can be made to count from a given initial value. This type of counter is called a loadable counter.



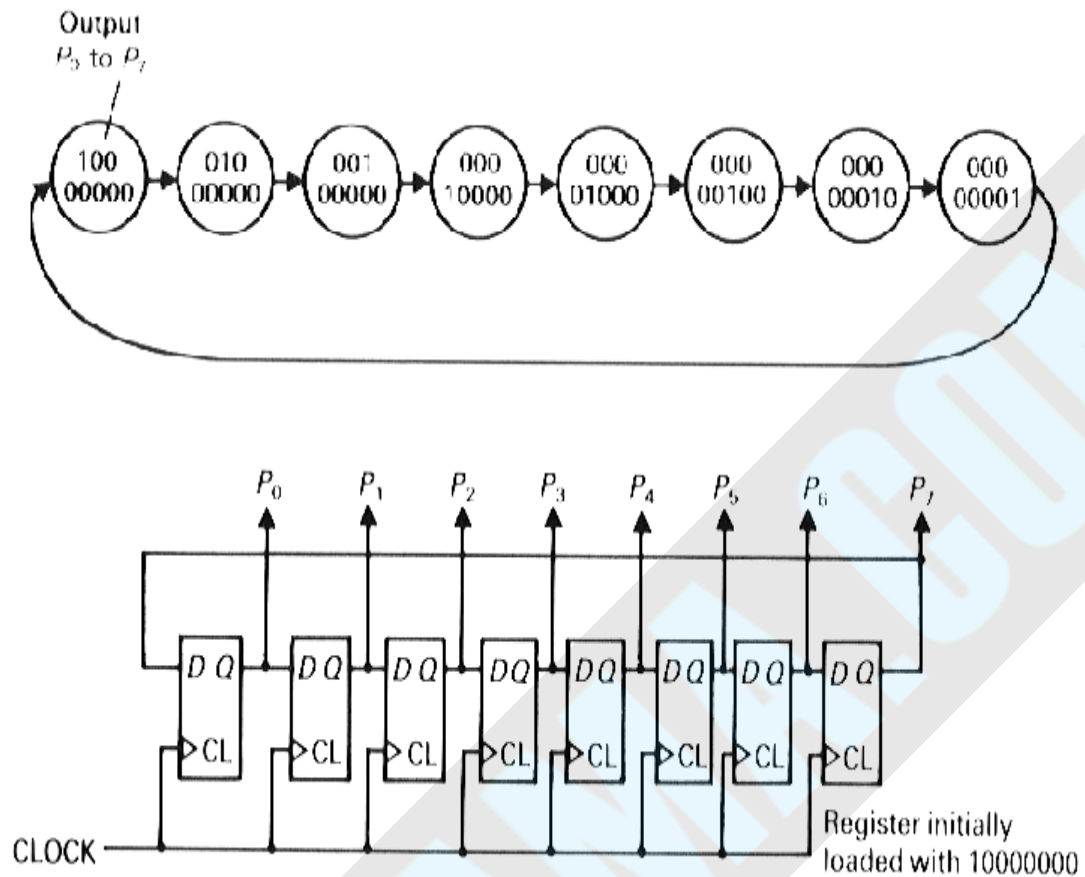
😊BCD counters.

A BCD counter counts in binary-coded decimal from 0000 to 1001 and back to 0000. Because of the return to 0 after a count of 9, a BCD counter does not have a regular pattern as in a straight binary count.



😊 Ring counters.

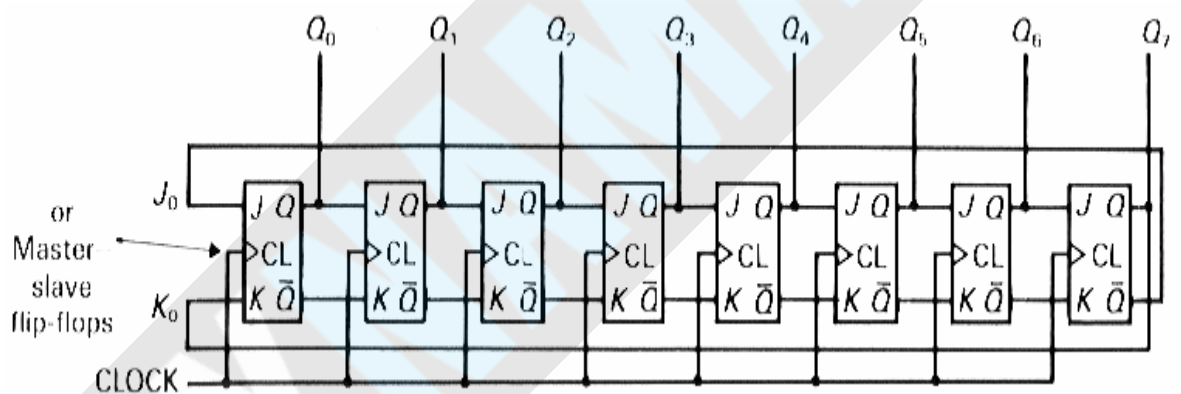
A ring counter is a circular shift register with only one flip-flop being set at any particular time; all others are cleared. The single bit is shifted from one flip-flop to the other to produce the sequence of timing signals.



😊 Johnson counters.

The Johnson counter, also called the twisted ring counter, is a variation of the ring counter, with the inverse output of the most significant flip-flop passed to the input of the least significant flip-flop. The sequence followed begins with all 0's in the register. The final 0 will cause 1's to be shifted into the register from the left-hand side when clock pulses are applied. When the first 1 reaches the most significant flip-flop, 0's will be inserted into the first flip-flop because of the cross-coupling between the output and the input of the counter.

Clock pulse	Q_0	Q_1	Q_2	Q_3	Q_4	Q_5	Q_6	Q_7
Φ_0	0	0	0	0	0	0	0	0
Φ_1	1	0	0	0	0	0	0	0
Φ_2	1	1	0	0	0	0	0	0
Φ_3	1	1	1	0	0	0	0	0
Φ_4	1	1	1	1	0	0	0	0
Φ_5	1	1	1	1	1	0	0	0
Φ_6	1	1	1	1	1	1	0	0
Φ_7	1	1	1	1	1	1	1	0
Φ_8	1	1	1	1	1	1	1	1
Φ_9	0	1	1	1	1	1	1	1
Φ_{10}	0	0	1	1	1	1	1	1
Φ_{11}	0	0	0	1	1	1	1	1
Φ_{12}	0	0	0	0	1	1	1	1
Φ_{13}	0	0	0	0	0	1	1	1
Φ_{14}	0	0	0	0	0	0	1	1
Φ_{15}	0	0	0	0	0	0	0	1



Shift Registers

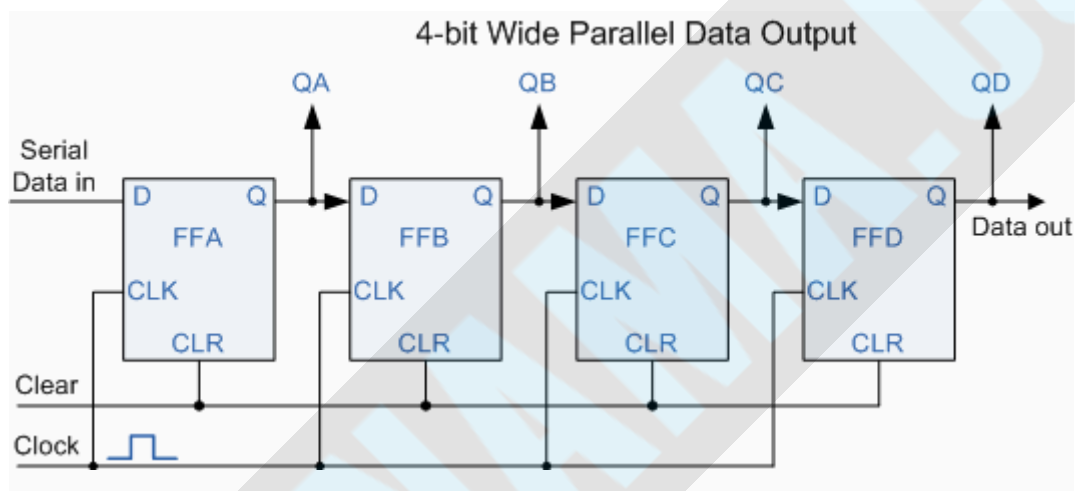
Shift Registers consists of a number of single bit "D-Type Data Latches" connected together in a chain arrangement so that the output from one data latch becomes the input of the next latch and so on, thereby moving the stored data serially from either the left or the right direction. The number of individual Data Latches used to make up **Shift Registers** are determined by the number of bits to be stored with the most common being 8-bits wide. Shift Registers are mainly used to store data and to convert data from either a serial to parallel or parallel to serial format with all the latches being driven by a common clock (Clk) signal making them Synchronous devices. They are generally provided with a Clear or Reset connection so that they can be "SET" or "RESET" as required.

Generally, Shift Registers operate in one of four different modes:

- Serial-in to Parallel-out (SIPO)
- Serial-in to Serial-out (SISO)
- Parallel-in to Parallel-out (PIPO)
- Parallel-in to Serial-out (PISO)

Serial-in to Parallel-out.

4-bit Serial-in to Parallel-out (SIPO) Shift Register



Lets assume that all the flip-flops (FFA to FFD) have just been RESET (CLEAR input) and that all the outputs QA to QD are at logic level "0" ie, no parallel data output. If a logic "1" is connected to the DATA input pin of FFA then on the first clock pulse the output of FFA and the resulting QA will be set HIGH to logic "1" with all the other outputs remaining LOW at logic "0".

Assume now that the DATA input pin of FFA has returned LOW to logic "0". The next clock pulse will change the output of FFA to logic "0" and the output of FFB and QB HIGH to logic "1". The logic "1" has now moved or been "Shifted" one place along the register to the right. When the third clock pulse arrives this logic "1" value moves to the output of FFC (QC) and so on until the arrival of the fifth clock pulse which sets all the outputs QA to QD back again to logic level "0" because the input has remained at a constant logic level "0".

The effect of each clock pulse is to shift the DATA contents of each stage one place to the right, and this is shown in the following table until the complete DATA is stored, which can now be read directly from the outputs of QA to QD.

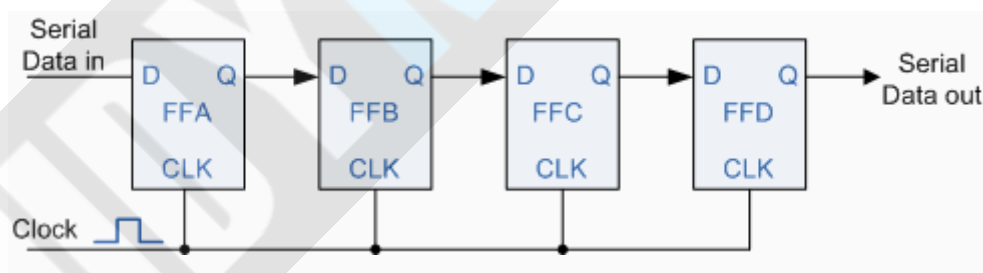
Then the DATA has been converted from a Serial Data signal to a Parallel Data word.

Clock Pulse No	QA	QB	QC	QD
0	0	0	0	0
1	1	0	0	0
2	0	1	0	0
3	0	0	1	0
4	0	0	0	1
5	0	0	0	0

Serial-in to Serial-out

This Shift Register is very similar to the one above except where as the data was read directly in a parallel form from the outputs QA to QD, this time the DATA is allowed to flow straight through the register. Since there is only one output the DATA leaves the shift register one bit at a time in a serial pattern and hence the name **Serial-in to Serial-Out Shift Register**.

4-bit Serial-in to Serial-out (SISO) Shift Register



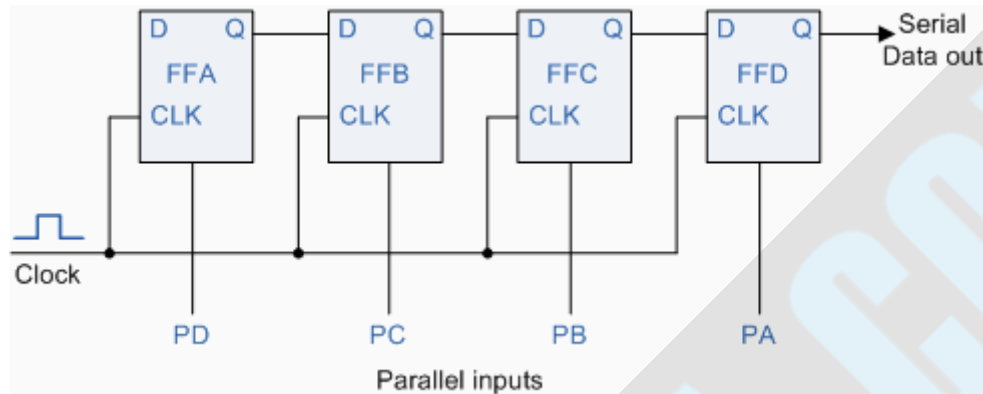
This type of **Shift Register** also acts as a temporary storage device or as a time delay device, with the amount of time delay being controlled by the number of stages in the register, 4, 8, 16 etc or by varying the application of the clock pulses. Commonly available IC's include the 74HC595 8-bit Serial-in/Serial-out Shift Register with 3-state outputs.

Parallel-in to Serial-out

Parallel-in to Serial-out Shift Registers act in the opposite way to the Serial-in to Parallel-out one above. The DATA is applied in parallel form to the parallel

input pins PA to PD of the register and is then read out sequentially from the register one bit at a time from PA to PD on each clock cycle in a serial format.

4-bit Parallel-in to Serial-out (PISO) Shift Register

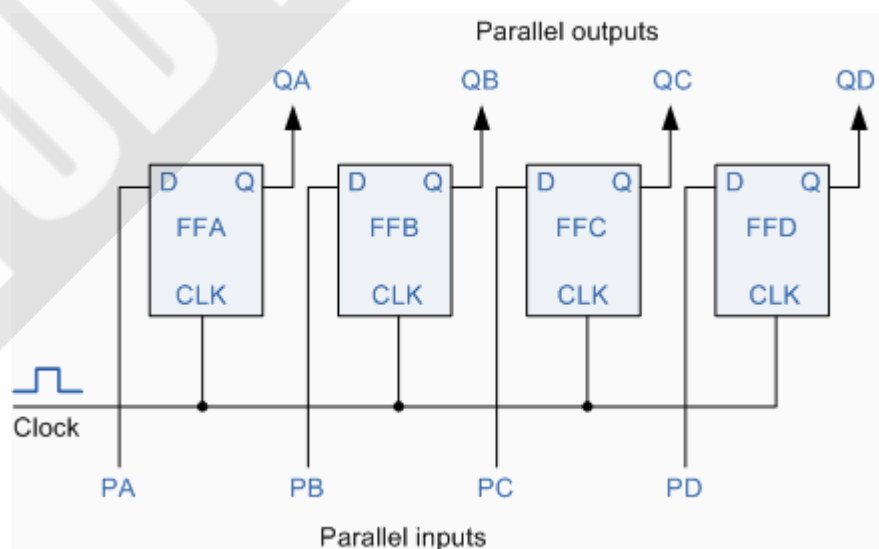


As this type of Shift Register converts parallel data, such as an 8-bit data word into serial data it can be used to multiplex many different input lines into a single serial DATA stream which can be sent directly to a computer or transmitted over a communications line. Commonly available IC's include the 74HC165 8-bit Parallel-in/Serial-out Shift Registers.

Parallel-in to Parallel-out

Parallel-in to Parallel-out Shift Registers also act as a temporary storage device or as a time delay device. The DATA is presented in a parallel format to the parallel input pins PA to PD and then shifts it to the corresponding output pins QA to QD when the registers are clocked.

4-bit Parallel-in/Parallel-out (PIPO) Shift Register



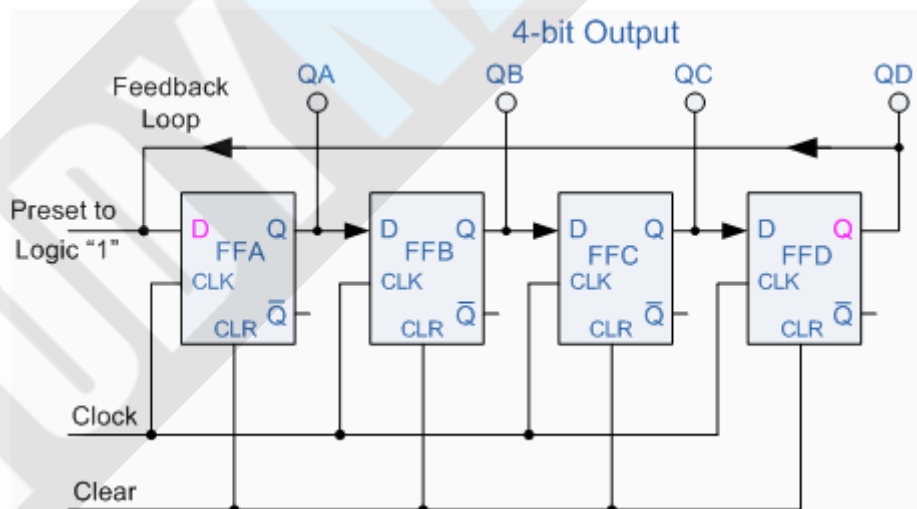
As with the Serial-in to Serial-out shift register, this type of register also acts as a temporary storage device or as a time delay device, with the amount of time delay being varied by the frequency of the clock pulses.

Today, high speed bi-directional universal type **Shift Registers** such as the TTL 74LS194, 74LS195 or the CMOS 4035 are available as a 4-bit multi-function devices that can be used in serial-serial, shift left, shift right, serial-parallel, parallel-serial, and as a parallel-parallel Data Registers, hence the name "**Universal**".

Ring Counters

we apply a serial data signal to the input of a Serial-in to Serial-out Shift Register, the same sequence of data will exit from the last flip-flop in the register chain after a preset number of clock cycles thereby acting as a time delay to the original signal. But what if we were to connect the output of the Shift Register back to its input, we then have a closed loop circuit that "Recirculates" the DATA around a loop, and this is the principal operation of **Ring Counters** or **Walking Ring Counter**. Consider the circuit below.

4-bit Ring Counter



The synchronous **Ring Counter** example above, will re-circulate the same DATA pattern between the 4 Flip-flops over and over again every 4th clock cycle, as long as the clock pulses are applied to it. But in order to cycle the DATA we must first "Load" the counter with a suitable DATA pattern for it to

work correctly as all logic "0"'s or all logic "1"'s outputted at each clock cycle would make the ring counter invalid.

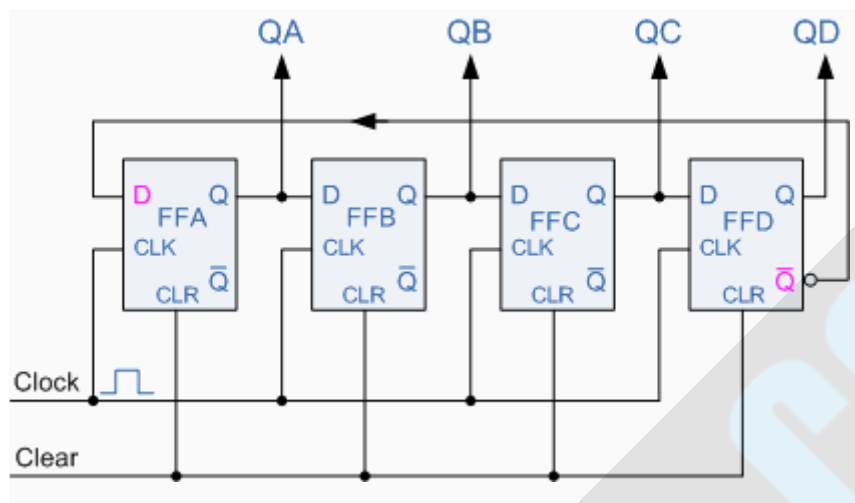
For **Ring Counters** to operate correctly they must start with the first flip-flop (FFA) in the logic "1" state and all the others at logic "0". To achieve this, a "CLEAR" signal is firstly applied to all the Flip-flops in order to "RESET" their outputs to a logic "0" level and then a "PRESET" pulse is applied to the input of the first Flip-flop (FFA) before the clock pulses are applied. This then places a single logic "1" value into the circuit of the Ring Counters.

The ring counter example shown above is also known as a "**MODULO-4**" or "MOD-4" counter since it has 4 distinct stages and each Flip-flop output has a frequency equal to one-fourth or a quarter ($1/4$) that of the main clock frequency. The "MODULO" or "MODULUS" of a counter is the number of states the counter counts or sequences through before repeating itself and a ring counter can be made to output any MODULO number and a "MOD-N" Ring Counter will require "N" number of Flip-flops connected together. For example, a MOD-8 Ring Counter requires 8 Flip-flops and a MOD-16 Ring Counter would require 16 Flip-flops.

Johnson Ring Counters

Johnson Ring Counters or "Twisted Ring Counters", are exactly the same idea as the *Walking Ring Counter* above, except that the inverted output Q of the last Flip-flop is connected back to the input D of the first Flip-flop as shown below. The main advantage of this type of ring counter is that it only needs half the number of Flip-flops compared to the standard walking ring counter then its Modulo number is halved.

4-bit Johnson Ring Counter



This inversion of Q before it is fed back to input D causes the counter to "count" in a different way. Instead of counting through a fixed set of patterns like the walking ring counter such as for a 4-bit counter, "1000"(1), "0100"(2), "0010"(4), "0001"(8) etc, the Johnson counter counts up and then down as the initial logic "1" passes through it to the right replacing the preceding logic "0". A 4-bit Johnson ring counter passes blocks of four logic "0" and then four logic "1" thereby producing an 8-bit pattern. As the inverted output Q is connected to the input D this 8-bit pattern continually repeats. For example, "1000", "1100", "1110", "1111", "0111", "0011", "0001", "0000" and this is demonstrated in the table below.

FFA	FFB	FFC	FFD
0	0	0	0
1	0	0	0
1	1	0	0
1	1	1	0
1	1	1	1
0	1	1	1
0	0	1	1
0	0	0	1

As well as counting, Ring Counters can be used to detect or recognise various patterns or number values. By connecting simple logic gates such as **AND** or **OR** gates to the outputs of the Flip-flops the circuit can be made to detect a set number or value. Standard 2, 3 or 4-stage Johnson Ring Counters can also be used to divide the frequency of the clock signal by varying their feedback connections and divide-by-3 or divide-by-5 outputs are also available.

A 3-stage Johnson Ring Counter can also be used as a 3-phase, 120 degree phase shift square wave generator by connecting to the outputs from A, B and NOT-B. The standard 5-stage Johnson counter such as the commonly available CD4017 is generally used as a Synchronous Decade Counter/Divider circuit. The smaller 2-stage circuit is also called a "Quadrature" (sine/cosine) Oscillator/Generator and is used to produce 4 outputs that are each "phase shifted" by 90 degrees with respect to each other, and this is shown below.

Part A

1. Define CMRR of op.amp.
2. What is virtual ground?
3. Convert binary 0110112 to Hexadecimal.
4. State logic equation for EX-OR gate
5. Define a multiplexer
6. State difference between Half adder and full adder
7. What is a D-Type Flip-Flop?
8. State difference between Synchronous and asynchronous counter
9. How many comparators are required for a 4 bit parallel comparator (simultaneous) A/D converter?
10. State difference between static and dynamic memory.
11. What is virtual Ground of an op.amp? and explain Op amp as inverter.
12. Give the Truth Table of 2 input Ex-OR gate and NOR gate.
13. Define fan in and fan out of a logic gate.

14. Give the logic diagram and Truth Table of JKMS FF.

PART - B

1. Derive the expression for the CE short circuit current gain of transistor at high frequency
2. i) What is the effect of $C_{b'e}$ on the input circuit of a BJT amplifier at High frequencies?
ii) Derive the equation for g_m which gives the relation between g_m , I_c and temperature.
3. Draw the high frequency hybrid $-h$ model for a transistor in the CE configuration and explain the significance of each component. Define alpha cut off frequency.
4. Draw the high frequency equivalent circuit of FET amplifier and derive all the parameters related to its frequency response.
5. Using hybrid $-h$ model for CE amplifier derive the expression for its short circuit current gain.
6. i) Define the frequency response of multistage amplifier and derive its upper and lower cut-off frequencies.
ii) How does Rise and Sag time related to cut-off frequencies and prove the same.

UNIT V

FUNDAMENTALS OF COMMUNICATION ENGINEERING

Networks:

Symmetrical and asymmetrical networks. characteristic impedance and propagation constant Derivation of characteristic impedance for T and Pi networks using Z_{oc} and Z_{sc} , image and iterative impedances - Derivation of Z_{i1} and Z_{i2} for asymmetrical T and L networks using Z_{oc} and Z_{sc} , Derivation of iterative impedances for asymmetrical T network. Equaliser: types, applications: constant resistance equalizer. (No derivations)

Symmetrical Networks:

A network in which all devices can send and receive data at the same rates. Symmetric networks support more bandwidth in one direction as compared to the other, and symmetric DSL offers clients the same bandwidth for both downloads and uploads. A lesser used definition for symmetric network involves resource access—in particular, the equal sharing of resource access.

Antenna:

Basic antenna principle, directive gain, directivity, radiation pattern, broad-side and end-fire array, Yagi antenna, Parabolic antenna.

Antenna Directivity:

Directivity is an important quality of an antenna. It describes how well an antenna concentrates, or bunches, radio waves in a given direction. A dipole transmits or receives most of its energy at right angles to the lengths of metal, while little energy is transferred along them.

If the dipole is mounted vertically, as is common, it will radiate waves away from the center of the antenna in all directions. However, for a commercial radio or television station, a transmitting antenna is often designed to concentrate the radiated energy in certain directions and suppress it in others.

For instance, several dipoles can be used together if placed close to one another. Such an arrangement is called a multiple-element antenna, which is also known as an array.

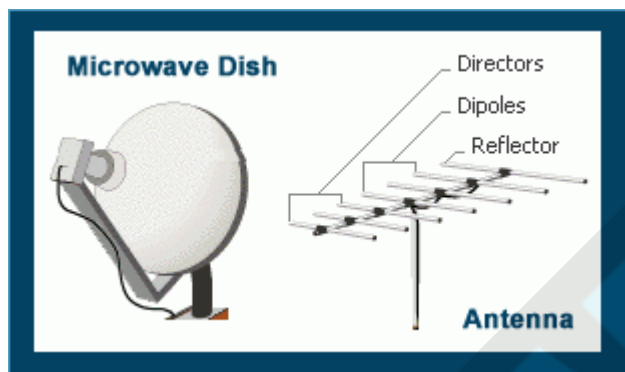
By properly arranging the separate elements and by properly feeding signals to the elements, the broadcast waves can be more efficiently concentrated toward an intended audience, without, for example, wasting broadcast signals over uninhabited areas.

Basic Antenna principle:

Antenna:

Antenna, also referred to as an aerial, device used to radiate and receive radio waves through the air or through space. Antennas are used to send radio waves to distant sites and to receive radio waves from distant sources. Many wireless communications devices, such as radios, broadcast television sets, radar, and cellular radio telephones, use antennas. Receiving antennas come in many different shapes, depending on the frequency and wavelength of the intended signal.

How Antenna works?



A transmitting antenna takes waves that are generated by electrical signals inside a device such as a radio and converts them to waves that travel in an open space. The waves that are generated by the electrical signals inside radios and other devices are known as guided waves, since they travel through transmission lines such as wires or cables.

The waves that travel in an open space are usually referred to as free-space waves, since they travel through the air or outer space without the need for a transmission line. A receiving antenna takes free-space waves and converts them to guided waves.

Radio waves are a type of electromagnetic radiation, a form of rapidly changing, or oscillating, energy. Radio waves have two related properties known as **frequency** and **wavelength**.

Frequency refers to the number of times per second that a wave oscillates, or varies in strength.

The wavelength is equal to the speed of a wave (the speed of light, or 300 million m/sec) divided by the frequency. Low-frequency radio waves have long wavelengths (measured in hundreds of meters), whereas high-frequency radio waves have short wavelengths (measured in centimeters).

An antenna can radiate radio waves into free space from a transmitter, or it can receive radio waves and guide them to a receiver, where they are reconstructed into the original message. For example, in sending an AM radio transmission, the radio first generates a carrier wave of

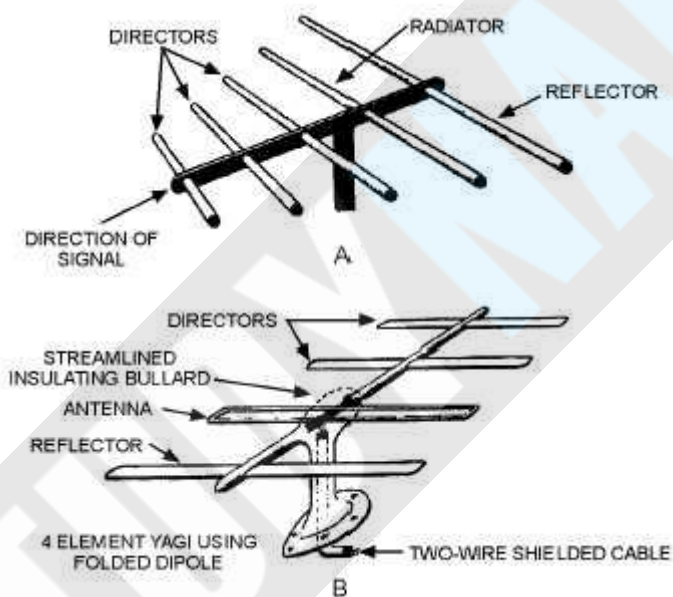
energy at a particular frequency. The carrier wave is modified to carry a message, such as music or a person's voice.

The modified radio waves then travel along a transmission line within the radio, such as a wire or cable, to the antenna. The transmission line is often known as a feed element. When the waves reach the antenna, they oscillate along the length of the antenna and back. Each oscillation pushes electromagnetic energy from the antenna, emitting the energy through free space as radio waves.

The antenna on a radio receiver behaves in much the same way. As radio waves traveling through free space reach the receiver's antenna, they set up, or induce, a weak electric current within the antenna. The current pushes the oscillating energy of the radio waves along the antenna, which is connected to the radio receiver by a transmission line. The radio receiver amplifies the radio waves and sends them to a loudspeaker, reproducing the original message.

Yagi antenna:

The Yagi antenna or more correctly, the Yagi - Uda antenna was developed by Japanese scientists in the 1930's. It consists of a half wave dipole (sometimes a folded one, sometimes not), a rear "reflector" and may or may not have one or more forward "directors". These are collectively referred to as the "elements".



The Yagi antenna consists of 2 parts:

- the antenna elements
- the antenna boom

There are three types of elements:

- the Reflector (REFL)
- the Driven Element (DE)

- the Directors (DIR)

Yagi antenna components:

Each Yagi antenna consists of dipoles, reflectors and directors. A dipole antenna receives radio frequency energy in a circular field ending at the center of the dipole. The Yagi antenna uses a series of dipoles in order to allow for a wider range of signal to reach the antenna.

With a Yagi antenna all parts of the antenna usually lay on the same plane. This can be extremely useful, especially with more modern Yagi antennas. The more dipoles that the Yagi antenna has on the same plane, the more bands of signal it can pick up at the same time

The Reflector is at the back of the antenna furthest away from the transmitting station. In other words the boom of the antenna is pointed towards the radio station over the horizon with the Reflector furthest away from the station.

The Driven Element is where the signal is intercepted by the receiving equipment and has the cable attached that takes the received signal to the receiver.

Amplitude Modulation:

MODULATION/DEMODULATION:

Modulation is the process of varying some characteristic of a periodic wave with an external signals. Modulation is the modifying of a signal to carry intelligent data over the communications channel. Several types of modulation are available, depending on the system requirement and equipment. The most frequently used types of modulation are amplitude modulation, frequency modulation, and phase modulation.

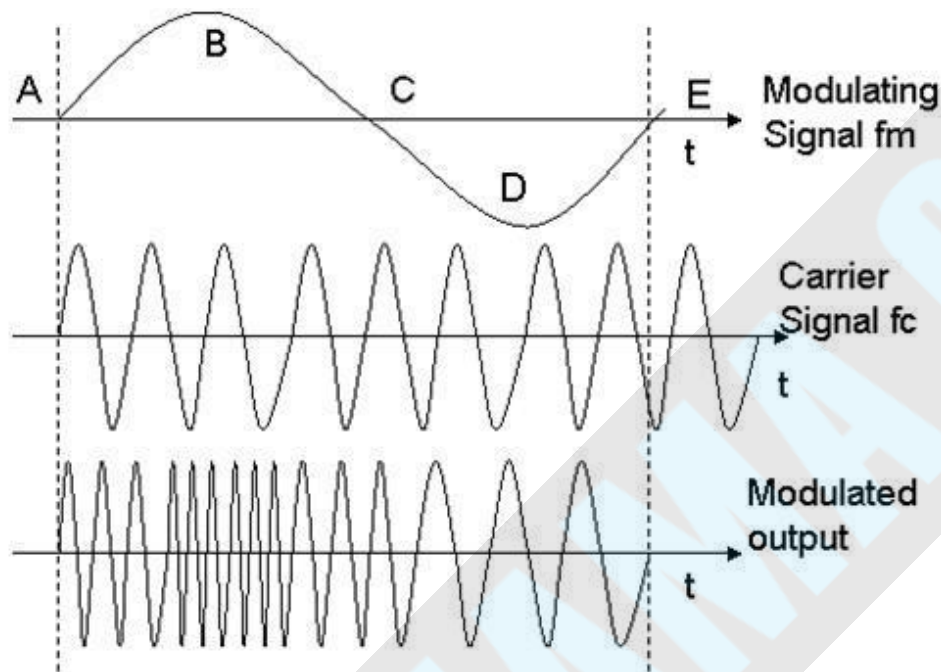
Demodulation is the act of returning modulated data signals to their original form.

1. **Amplitude modulation(AM):**
Amplitude modulation refers to modifying the amplitude of a sine wave to store data.
2. **Frequency Modulation (FM):**
Frequency modulation refers to changing the frequency of a signal to indicate a logic 1 or a logic 0. One frequency indicates a logic 1, and the other frequency indicates a logic 0.
3. **Phase Modulation (PM or Indirect FM):**
Phase modulation is more complex than amplitude modulation or frequency modulation. Phase modulation uses a signal frequency sine wave and performs phase shifts of the sine wave to store data. A modification of phase modulation involves the use of several discrete phase shifts to indicate the state of two or more data bits.

Frequency Modulation:

Frequency Modulation (FM) With frequency modulation, the modulating signal and the carrier are combined in such a way that causes the carrier FREQUENCY(f_c) to

vary above and below its normal(idling) frequency. The amplitude of the carrier remains constant as shown in figure below.



Microphones:

Introduction:

A microphone is a transducer as it converts sound waves (acoustic energy) into electrical energy. The very first microphone was purely mechanical in nature. A metal diaphragm is connected to a needle, which “draws” a pattern on a metallic foil. When the air pressure changes due to a person’s voice, the diaphragm vibrates and moves the needle. The needle then scratches the foil with the vibration pattern. The sound is recreated when the needle is made to run over the foil again. The vibration pattern being followed by the needle makes the diaphragm move and reproduces the sound.

Microphones now work the same way but does the process electronically. Instead of a scratched foil with the vibration patterns, the change in air pressure is now converted to an electrical signal. The diaphragms can be of any material such as plastic, paper or aluminum. Diaphragms differ in producing sound which gave rise to different classifications of microphones

Types of microphones:

Carbon microphones:

Carbon microphones are amongst the oldest, simplest and most used types of microphones even to this day. They work by converting air pressure variations into electrical resistance. The membrane collecting the sound waves presses against a carbon dust material that varies its electrical resistance in the process. By running electric current through the carbon dust, one can obtain an electrical current variation that is amplified and recorded.

Condenser Micorphones:

Condenser microphones rely on the properties of capacitors. However, the plates of the capacitor are no longer immobile and are free to move in relation to each other according to the air pressure changes. This generates a variation in the capacity of the device, which can be converted into electric signals.

Dynamic microphones:

Dynamic microphones on the other hand harness the electromagnetic effects determined by the movement of a magnet inside a conductive wire coil. The vibrations of the magnet are basically converted into tiny electrical currents that are amplified and recorded.

Ribbon microphones:

Ribbon microphones work on a principle rather similar to that of the dynamic microphones, but instead of vibrating a microphone inside a coil, a thin ribbon is suspended in a magnetic field. The vibration of the ribbon translates into inductance variations inside the coil generating the magnetic field.

Piezo-electric microphones:

Crystal microphones are based on the piezoelectric effect. Piezoelectric materials have the ability of directly converting electric energy into mechanical movement and vice versa. The most common piezoelectric material occurring naturally on Earth is quartz, which is often used to make crystal microphones.

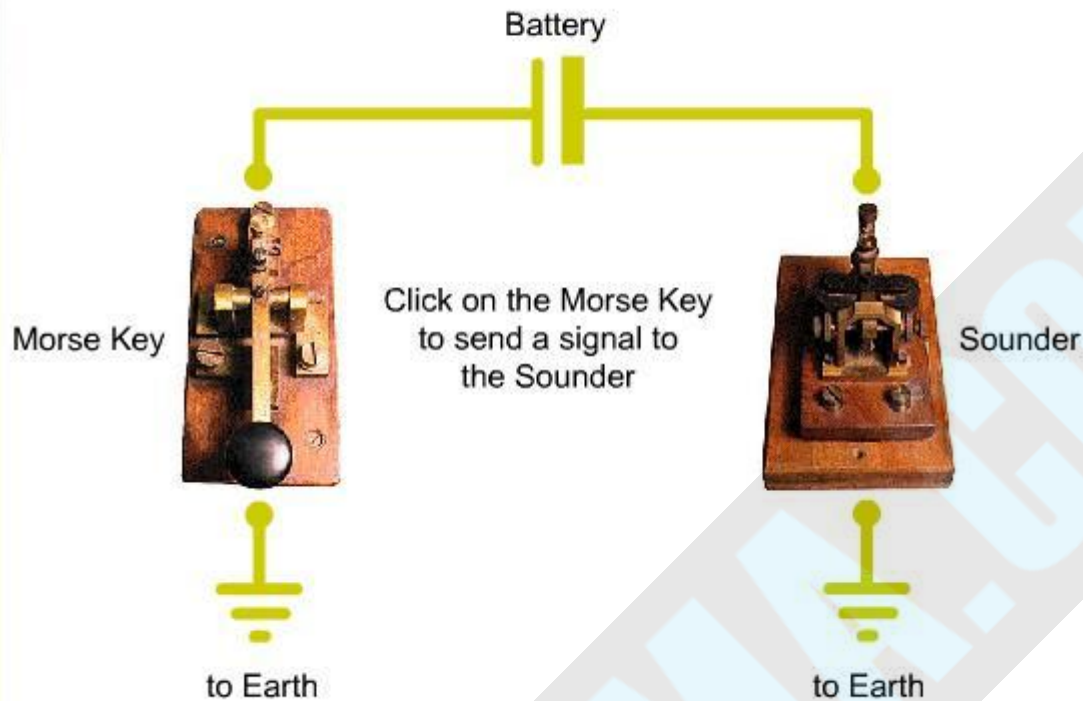
Telegraphy

Telegraphy is the long-distance transmission of written messages without physical transport of letters. **Radiotelegraphy** or **wireless telegraphy** transmits messages using radio.

Telegraphy includes recent forms of data transmission such as fax, email, and computer networks in general.

A **telegraph** is a machine for transmitting and receiving messages over long distances. A telegraph message sent by a telegraph operator (or telegrapher) using Morse code was known as a **telegram** or **cablegram**, often shortened to a *cable* or a *wire* message. Later, a telegram sent by the Telex network, a switched network of teleprinters similar to the telephone network, was known as a **telex** message.

Basic Principles of Telegraphy



Morse Code:

Morse code is a type of character encoding that transmits telegraphic information using rhythm. Morse code uses a standardized sequence of short and long elements to represent the letters, numerals, punctuation and special characters of a given message. The short and long elements can be formed by [sounds](#), marks, or pulses, in on off keying and are commonly known as "dots" and "dashes" or "dits" and "dahs". The speed

of Morse code is measured in words per minute or characters per minute, while fixed-length data forms of telecommunication transmission are usually measured in baud or bps.

Television:

Television (TV) is a widely used telecommunication medium for transmitting and receiving moving images, either monochromatic ("black and white") or color, usually accompanied by sound. "Television" may also refer specifically to a television set, television programming or television transmission.



Charge-Coupled Device

Charge-coupled device (CCD) is an analog shift register that enables the transportation of analog signals (electric charges) through successive stages (capacitors), controlled by a clock signal. Charge-coupled devices can be used as a form of memory or for delaying samples of analog signals. Today, they are most widely used in arrays of photoelectric light sensors to serialize parallel analog signals. Not all image sensors use CCD technology; for example, CMOS chips are also commercially available.

"CCD" refers to the way that the image signal is read out from the chip. Under the control of an external circuit, each capacitor can transfer its electric charge to one or another of its neighbors. CCDs are used in digital photography, digital photogrammetry, astronomy (particularly in photometry), sensors, electron microscopy, medical fluoroscopy, optical and UV spectroscopy, and high speed techniques such as lucky imaging.

Television is certainly one of the most influential forces of our time. Through the device called a television set or TV, you are able to receive news, sports, entertainment, information and commercials. The average American spends between two and five hours a day glued to "the tube"!

Have you ever wondered about the technology that makes television possible? How is it that dozens or hundreds of channels of full-motion video arrive at your house, in many cases for free? How does your television decode the signals to produce the picture? How will the new digital television signals change things? If you have ever wondered about your television (or, for that matter, about your computer monitor), then read on! In this article, we'll answer all of these questions and more. See the next page to get started.

conversion of the vibrations of sound (for example, music) into a permanent record, and its later playback in its original form (see SOUND,). In the most common method of sound recording, the magnetic method, transformed sound waves may be amplified and made to magnetize a metaloxide coated plastic recording tape so that the magnetization varies with the frequency and intensity of the sound. Sound recording involves some form of mechanical movement of the recording medium at a constant speed past the point of recording so that the sound recording may later be reproduced as a replica of the original sound.

Components of Television

HIGH FIDELITY

High fidelity is the technique of recording, broadcasting, and reproducing sound to match as closely as possible the characteristics of the original sound. To achieve high-fidelity reproduction, the sound must be free of distortion and include the full frequency range of human hearing—20 Hz to 20 kilohertz (see FREQUENCY,).

Components.

A high-fidelity system consists of the following components: the turntable and tonearm or possibly a CD player, the amplifier, the speaker system, and the control unit, sometimes referred to as a preamplifier/control unit. Supplementary components include the tuner and the tape recorder.

The turntable and tonearm.

(For the basic operating principles of these elements, *see* PHONOGRAPH,.) The turntable and tonearm translate the engraved patterns on a phonograph record into electrical voltage variations. The turntable is rotated by a motor that turns at a constant speed, thus avoiding distortions called wow and rumble. Wow consists of a slow variation in pitch caused by variation in the speed of the turntable, and rumble is a low-frequency tremor caused by defects in the turntable.

The tonearm and the cartridge form one of the most critical parts of the high-fidelity installation. The finely balanced tonearm holds a cartridge, which in turn holds a stylus, preferably tipped with long-wearing diamond. To reproduce recorded sound accurately and with minimum wear on the record, the cartridge must provide maximum compliance, that is, an easy lateral and vertical motion of the stylus. The stylus, moreover, must contact the record at a precise angle with the proper pressure.

The compact disc (CD) player.

CD players are increasingly replacing the conventional turntable and tonearm in high-fidelity systems. Offering more uniform frequency response, lower distortion, and inaudible background noise levels, compact discs have the additional advantage of longer life. Since CDs are never physically in contact with any pickup mechanism—digital codes embedded beneath the surface of the disc are read by a laser beam of light—these discs can last indefinitely if handled with care. Specially built CD players can also be used for data retrieval using CD-ROM (Read-Only Memory) discs, while interactive compact discs (CD-I), as well as interactive video discs (VD-I), can be used for a wide variety of educational and training purposes. In addition to their audio content, some CDs contain digitally driven graphics that can be displayed on a television screen. Such discs are referred to as CD-G.

The amplifier.

The amplifier converts the relatively weak electrical impulses received from the cartridge into power sufficient to drive the speakers. The amount of power that an amplifier can produce is rated in watts. Depending on the requirements of the speaker system, an amplifier may deliver from 10 to 125 watts of electrical power. The amplifier is controlled, as a rule, by a device called the preamplifier, which amplifies minute sound-signal voltages too small for the amplifier to handle. Preamplifiers also boost the bass and attenuate the treble to compensate for the poor bass and strong treble response of phonograph records. Most modern

amplifiers are equipped with so-called solid-state or integrated circuits. See INTEGRATED CIRCUIT.

The speaker system.

Loudspeakers, electromechanical devices that produce audible sound from amplified audio voltages, are extensively employed in radio receivers, motion picture sound systems, public-address systems, and other apparatus in which sound must be produced from a recording, a communications system, or a sound source of low intensity.

Several types of loudspeaker exist, but almost all loudspeakers now in use are dynamic speakers. These speakers include an extremely light coil of wire, called the voice coil, mounted within the magnetic field of a powerful permanent magnet or electromagnet. The coil of the electromagnet, if one is used, is called the field coil. A varying electric current from the amplifier passes through the voice coil and alters the magnetic force between the voice coil and the speaker's magnetic field. As a result, the coil vibrates with the changes in the current. A diaphragm or a large paper cone mechanically attached to the voice coil generates sound waves in the air when the coil moves.

The loudness and sound quality of such speakers can be increased by the use of properly designed enclosures or cabinets. Such cabinets may hold several loudspeakers of different sizes, small so-called tweeters for high notes, and large woofers for low notes.

The control unit.

As the nerve center of the high-fidelity system, the control unit performs a number of critical functions. For example, the surface noises of old records are attenuated by means of a device called the scratch filter; another device, the rumble filter, cuts down low-pitched noises, such as vibration from the phonograph motor; the loudness control compensates for the inability of the ear to hear high and low notes as clearly as it hears the middle range by increasing the relative level of treble and bass tones when the record is played at a reduced volume. The control unit also adjusts sound signals from the record player, the tape recorder, or the tuner.

The tuner.

The AM/FM tuner allows the listener to receive broadcasts from stations in the broadest band of the radio spectrum, from 500 to 1650 kilohertz (AM), 88 to 108 megahertz (FM). From the broadcast signals reaching the antenna, the tuner selects the frequency of the desired station to the exclusion of other stations in the broadcast range. It then extracts the audio voltage representing the program being transmitted and amplifies this voltage to activate the speakers of the high-fidelity system. See RADIO,.

The tape recorder.

This device records and reproduces sound by preserving electrical signals as magnetic patterns on thin plastic tape coated with magnetic oxide. In recording, the tape is drawn past a recording head, leaving a magnetic imprint. The tape is then drawn past a reproducing head that turns the magnetic pattern into an electrical signal; this signal, in turn, is amplified and reproduced as sound. The reproducing, or playback, head may be the same device as the recording head, or they may be separate devices. Tapes are readily erased for reuse and are virtually immune to the damage that eventually mars phonograph records.

The first magnetic-reading instrument, called a telegraphone, was invented in 1898 by the Danish electrical engineer Valdemar Poulsen (1869–1942), who used a magnetized steel tape to carry messages. Currently, the most popular form of tape recording is the so-called compact cassette, which uses a tape with two or four tracks. Cassette-tape recorders and players are available in a wide variety of sizes, from the tiny portable types used with stereo headphones to elaborate units incorporated in home high-fidelity systems.

STEREOPHONIC SOUND

Stereophonic sound re-creates for listeners the conditions that would exist near an actual sound source, such as an orchestra. The sound is picked up separately from the left and the right sides of the orchestra, and, through the use of two or more carefully placed speakers, a stereophonic recording is directed toward listeners in such a way that they seem to hear music from the left, the right, and the center. More importantly, they become aware of a veil of sound that seems to have depth and solidity as well as direction.

TYPES OF RECORDING:

Mechanical Recording.

The operation of a sound-recording system may, however, be most easily understood by considering the process of recording sound by the now obsolete mechanical method. In this method, sound waves are used directly or indirectly to actuate a stylus or cutter that engraves on a disk or cylinder a wavy-line pattern corresponding to the pattern of sound waves. This process, with minor modifications, was used for many years in the production of phonograph records. In the direct method of mechanical recording, sound waves strike a very light diaphragm of metal or other substance and set it into motion. Attached to the diaphragm is a needle or cutting point that vibrates with the diaphragm. Under the point is a disk or cylinder of wax, metallic foil, shellac, or other suitable substance that is moved past the needle so that the needle cuts a groove in the form of a spiral on a disk, or a helix on a cylinder. As the needle vibrates it traces a wavy groove laterally or vertically in the record; this groove is a mechanical replica of the sound that struck the diaphragm of the recording machine. If, for example, the sound wave consists of the musical tone of A in the treble clef, which has a frequency of 440 Hz (hertz or cycles per second), the needle oscillates 440 times/sec. If the record is moving under the needle at the rate of 10 cm/sec, the groove will exhibit a pattern of 44 oscillations (44 sine waves, or 44 crests and 44 troughs)/1 cm (0.4 in). To reproduce the recorded sound, a needle attached to a diaphragm is set in the groove, and the record is turned at the rate of 10 cm (4 in)/sec. The vertical or, more commonly, lateral crests and troughs of

the groove then move the needle at the rate of 440 oscillations/sec, and the attached diaphragm oscillates, producing sound waves in the air of the same pitch as the original tone (see OSCILLATION,). In the making of modern phonograph records the sound is first converted to electrical impulses by a microphone; these impulses are amplified and used to actuate the cutting needle by electromagnetic means. The cutting needle engraves a disk, called the master, made of shellac, which is used to make the metal molds from which vinyl records are mass produced.

Optical Recording.

In the optical method, sound waves are transformed by a microphone into equivalent electrical impulses that are then amplified and made to operate a device that changes the intensity of a light beam (by means of an electromagnetically actuated gate or light valve) or the size of the beam (by means of an electromagnetically actuated vibrating mirror or a slit of variable width). The resulting varying light beam is focused on a moving film, which is then developed to provide a photographic track. The track recorded with varying intensity is of variable density and constant width. The track recorded with the vibrating mirror or varying slit has variable areas of darkened and clear film. To reproduce the sound track on either film, a light source is focused on the film, and a PHOTOELECTRIC CELL, (q.v.) is placed behind the film. The fluctuations in the relative amount of light passing through the film generate a fluctuating electric current in the photoelectric cell; this current is amplified and then transformed into sound by means of some form of loudspeaker. See MOTION PICTURE,.

Electromagnetic Recording.

In audiotape recording, sound waves are amplified and recorded on a magnetized plastic or paper tape. The information is first converted into electrical impulses, which are then impressed in the magnetized tape by an electromagnetic record head. A playback head, which is also an electromagnetic device, converts the magnetic fields on the tape into electrical impulses that are then amplified and reconverted into audible sound waves.

Digital Recording.

In the combined mechanical and electronic system of ordinary phonograph recording, waveforms of sound are inevitably distorted to some degree, and they also pick up noises from the recording process itself. In computer-based recording these problems are eliminated. The digital recorder measures the waveforms thousands of times each second and assigns a numerical value, or digit, to each of the measurements. These digits are then translated into a stream of electronic pulses that are placed in a memory bank for later retranslation and playback. Such techniques have been used limitedly in recent years for the production of otherwise conventional phonograph records, but direct-digital records are now available in which electronic pulses are instead placed on a small, aluminized disc called a COMPACT DISC (q.v.; CD), where they somewhat resemble a spiral of Morse-code signals when viewed through a microscope. The plastic-protected CD is placed in a machine where a laser beam reads the coded information, and circuitry converts it to analog signals for playback through conventional speaker systems.

Stereophonic Recording.

Stereophonic recording in its simplest form uses two separate microphones to produce two recorded tracks, or channels, on magnetic tape. Similarly, the sound component of motion pictures reproduces stereophonic sound by multiple tracks on film.

Phonograph disks can also record stereophonic sound, or stereo, on two independent channels, one on each wall of a single groove. The groove is cut with a 90° stylus in such a way that one groove wall slants 45° to the left, and the other wall slants 45° to the right. Two independent coils 90° apart energize the cutting stylus to give a different pattern on each wall for each of the two channels. During playback of a disk, two sensors in the cartridge are mounted at a 90° angle to each other to pick up the two tracks.

Quadraphonic Recording.

A quadraphonic sound playback system requires the use of four separate amplification channels driving four speakers located in the four corners of the listening room. Various systems achieving quadraphonic recording and playback were perfected in the early 1970s, some involving a method of encoding and decoding which required only two channels to be recorded on tape or disk.

The lack of standardization of these systems and the reluctance of many music lovers to place four loudspeakers in the listening room caused the popularity of quadraphonic recording to wane. With the advent of home video recorders, or VCRs, and large-screen television sets in the 1980s, a new type of multi-channel sound has replaced quadraphonic sound. Called surround sound, this system also involves the use of four or more loudspeakers and channels and is used to re-create the all-enveloping sound experienced when attending certain motion pictures in specially equipped theaters.

PART A

1. Define Large signal amplifier.
2. What are applications of power amplifier?
3. What are the features of large signal amplifiers?
4. What are the classification of large signal amplifiers?
5. What is class A amplifier?
6. What is class A amplifier?
7. What is class C amplifier?
8. What is class AB amplifier?
9. What is the construction of a class D amplifier?
10. What are the classification of Class A amplifier?
11. What are the advantages of directly coupled class A amplifier?
12. What are the advantages of transformer coupled class A amplifier?
13. What is frequency distortion?
14. Define heat sink.

PART B

1. With neat circuit diagram explain the working principle of complementary symmetry class-B amplifier and
2. Explain and obtain the efficiency of transformer – coupled class A power amplifier.
3. Prove that the maximum efficiency of Push Pull class B amplifier is 78.5%.

4. i) Compare class A, class B and class C power amplifier based on their performance characteristics

ii) Explain the significance of heat sinks for thermal stability.

5. What is the difference between a voltage amplifier and a power amplifier?

6. Differentiate Class S from Class D amplifier and derive the efficiency of Class D amplifier.

7. i) Explain the Concepts of Feed back in Amplifier.

ii) Give the properties of Negative feedback.

8. Discuss the differential voltage/current-series/shunt feedback connections with expression for gain, input resistance and output resistance.

9. Draw and explain various feedback amplifier topologies?

10. Determine the voltage gain, input and output impedance with feedback for Voltage series feedback having $A = -100$, $R_i = 10\text{ k}_\Omega$, $R_o = 20\text{ k}_\Omega$ and $\beta = -0.1$

11. With the topologies compare the four types of negative feedback amplifier?

BE/B.TECH EXAMINATION, NOVEMBER/DECEMBER 2004

Second Semestor

Electrical and Electronics Engineering

EC 141- BASIC ELECTRICAL AND ELECTRONICS

(Common to Civil, mechanical Engineering)

Time: 3 hrs

Maximum: 100 marks

Answer ALL questions

PART A-(10X2=20 marks)

1. Calculate the time taken by an electron which has been accelerated through a potential difference of 1000V to traverse a distance of 2cm. Given $q = 1.6 \times 10^{-19}$ and $m = 9.1 \times 10^{-31}$ kg.
2. State two applications of magnetic deflection.
3. Write down the expression for drift current density due to electrons.
4. Draw the resistance-temperature characteristics of thermistor and comment on it.
5. Define tunneling phenomenon.
6. Calculate the values of I_c and I_e for a transistor $\alpha_{dc} = 0.97$ and $I_{cbo} = 10 \mu A$ and I_b is measured as $50 \mu A$.
7. Depletion MOSFET is commonly known as "Normally-ON-MOSFET" why?
8. What are all internal capacitance in MOSFET?
9. What is "interbase resistance" of UJT?
10. What is ion implantation process?

PART B-(5X16=80 marks)

11. (i) The electron beam in a CRT is displayed vertically by a magnetic field of flux density $2 \times 10^{-4} \text{ wb/m}^2$. The length of the magnetic field along the tube axis is the same as that of the electrostatic deflection plates. The final anode voltage is 800V. Derive and calculate the voltage which should be applied to the Y-deflection plates 1cm apart, to return the spot back to the centre of the screen. (10)
(ii) Describe with neat diagram the principle of operation of Dynamic scattering type LCD. (6)
12. (a) Derive the continuity equation from the first principle and also derive 3 special cases of continuity equation. (16)

Or

- (b) (i) Derive the Ebers-Moll model for a PNP transistor and give equation for I_e and I_c . (8)
(ii) The diode current is 0.6mA when the applied voltage is 400mV and 20mA when the applied voltage is 500mV. Determine n (η). Assume $kT/q = 25 \text{ mV}$. (8)

13.(a) Explain Hall effect. How can Hall effect be used to determine some of the properties of a semiconductor and also discuss its applications. (16)

Or

(b) Describe with the help of a relevant diagram, the construction of a LED and explain its working. (16)

14.(a) With the help of suitable diagram explain the working of different types of MOSFET. (16)

Or

(b)(i) Draw and explain the energy band diagram for conductors, insulators and semiconductors. (4)
(ii) Sketch the energy band diagram for P-N junction under open circuit condition and obtain the expression for contact difference of potential E_0 . (12)

15.(a)(i) With volt-ampere characteristics describe the working principle of an SCR. Also explain its construction details. (14)

(ii) Draw the two transistor model of an SCR. (2)

Or

(b) With necessary diagrams explain the fabrication process of NMOS devices. (16)

Second Semestor

Electrical and Electronics Engineering

EC 141- BASIC ELECTRICAL AND ELECTRONICS

(Common to Civil, mechanical Engineering)

Time: 3hrs

Maximum: 100 marks

Answer ALL questions.

PART A-(10X2=20 marks)

1. What are the factors on which the electrostatic deflection sensitivity of a CRO depends?
2. What is mass action law?
3. What is the Fermi level in (a) in N type material and (b) a P type material?
4. A silicon diode has a saturation current of $7.5 \mu\text{A}$ at room temperature. Calculate the saturation current at 400K .
5. What happens to the capacitance of a varactor diode when the reverse bias voltage across it is increased? Mention two applications of varactor diode.
6. What is the wavelength of (a) ultra violet and (b) infrared?
7. Derive the expression for common emitter current gain in terms of common base one.
8. What is early effect? By what other name, it is known?
9. When is metal semiconductor contact preferred? Mention a familiar device having this type of contact.
10. Define intrinsic stand off ratio of an UJT. What is its range?

PART B-(5X16=80 marks)

11. (i) How is drain current related to gate to source voltage and drain to source voltage in a JFET? (8)

(ii) Compare BJT and FET. (8)

12. (a) (i) Describe the motion of electron in a magnetic field. (4)

(ii) Prove that the time taken by electron for one revolution is $2\pi/x/w$, where w is the angular velocity of the electron. (6)

(iii) Two plane parallel plates are placed 8mm apart and applied a voltage of 300V . An electron travels from the plate of lower potential to the higher potential one with an initial velocity of 10^6m/s . Calculate the time taken by the electron to reach the other plate. The charge of electron is $1.6 \times 10^{-19}\text{C}$ and its mass is $9.1 \times 10^{-31}\text{kg}$. (6)

Or

(b) What is Fermi-Dirac distribution function? What is its significance? Draw this function applied to

(i) intrinsic material

(ii) n type material and (iii) p type material and explain. (16)

13.(a) Discuss the electron and hole concentrations at thermal equilibrium with the help of schematic band diagram and carrier concentrations for (i) intrinsic (ii) p type and (iii) n type semiconductors. (16)

Or

(b)(i) What is hall effect? Discuss its applications. (8)
(ii) Write about drift current and diffusion current. (8)

14.(a) What is tunnel diode? Explain its operation with the help of band diagrams and I-V characteristics for various biasing conditions. (16)

Or

(b)(i) Discuss the switching characteristics of a junction diode for a pulse input. (10)
(ii) Write notes on photo diode and its characteristics. (6)

15.(a) Draw the basic structure and circuit symbol of an SCR. Draw the characteristic curve of this device and explain for its nature. Derive the equation for the anode current and mention its significance. (16)

Or

(b)(i) Write detailed note on: (1) rectifying contact and (2) ohmic contact. (8)
(ii) Mention the applications of triac and diac. (8)