

MAR GREGORIOS COLLEGE OF ARTS & SCIENCE

Block No.8, College Road, Mogappair West, Chennai – 37

Affiliated to the University of Madras
Approved by the Government of Tamil Nadu
An ISO 9001:2015 Certified Institution



DEPARTMENT OF ELECTRONICS & COMMUNICATION SCIENCE

SUBJECT NAME: ELECTRONIC DEVICES

SUBJECT CODE: SG22A

SEMESTER: II

PREPARED BY: PROF. S.FABBIYOLA / PROF. M. SATHIYA

ELECTRONIC DEVICES

UNIT I Semiconductor Basics : Conductor – Semiconductor – Introduction to Intrinsic and Extrinsic semiconductor – P type and N type semiconductor – PN junction diode – V-I characteristics - Half wave, Full wave & Bridge rectifier – expression for efficiency and ripple factor - Construction of Basic logic gates using Diodes.

UNIT II Special Purpose Diodes : Zener and Avalanche Break down, Zener diode - V-I characteristics regulated power supply using Zener diode- LED, Photodiode, PIN Diode, Varactor Diode, Tunnel Diode – Principle, Working& Applications.

UNIT III - Transistors : Transistor symbols NPN & PNP– Transistor biasing for active, saturation & cutoff –Operation of a BJT - Characteristics of a transistor in CE, CB & CC modes– Early effect – Punch-through– Transistor testing– Transistor as a switch – - Construction of Basic logic gates using Transistors (qualitative analysis)- Transistor as an amplifier-UJT– Basic construction and working-Characteristics.

UNIT IV - Field Effect Transistors : FET – Construction - Working - Static – Transfer characteristics – Parameters of FET– FET as an amplifier– MOSFET– Enhancement MOSFET– Depletion MOSFET – Construction & Working – Drain characteristics of MOSFET – Comparison of JFET & MOSFET.

UNIT V - Power Devices: Power Transistors-SCR–TRIAC–DIAC and IGBT–Characteristics and working

UNIT - I

Conductors

Conductors are materials that have very low values of resistivity, usually in the micro-ohms per metre. This low value allows them to easily pass an electrical current due to there being plenty of free electrons floating about within their basic atom structure. But these electrons will only flow through a conductor if there is something to spur their movement, and that something is an electrical voltage. When a positive voltage potential is applied to the material these “free electrons” leave their parent atom and travel together through the material forming an electron drift, more commonly known as a current. How “freely” these electrons can move through a conductor depends on how easily they can break free from their constituent atoms when a voltage is applied. Then the amount of electrons that flow depends on the amount of resistivity the conductor has. Examples of good conductors are generally metals such as Copper, Aluminium, Silver or non metals such as Carbon because these materials have very few electrons in their outer “Valence Shell” or ring, resulting in them being easily knocked out of the atom’s orbit.



An Electrical Cable uses Conductors and Insulators

This allows them to flow freely through the material until they join up with other atoms, producing a “Domino Effect” through the material thereby creating an electrical current. Copper and Aluminium is the main conductor used in electrical cables as shown. Generally speaking, most metals are good conductors of electricity, as they have very small resistance values, usually in the region of micro-ohms per metre, ($\mu\Omega.m$). While metals such as copper and aluminium are very good conductors of electricity, they still have some resistance to the flow of electrons and consequently do not conduct perfectly. The energy which is lost in the process of passing an electrical current, appears in the form of heat which is why conductors and especially resistors become hot as the resistivity of conductors increases with ambient temperature.

Insulators

Insulators on the other hand are the exact opposite of conductors. They are made of materials, generally non-metals, that have very few or no “free electrons” floating about within their basic atom structure because the electrons in the outer valence shell are strongly attracted by the positively charged inner

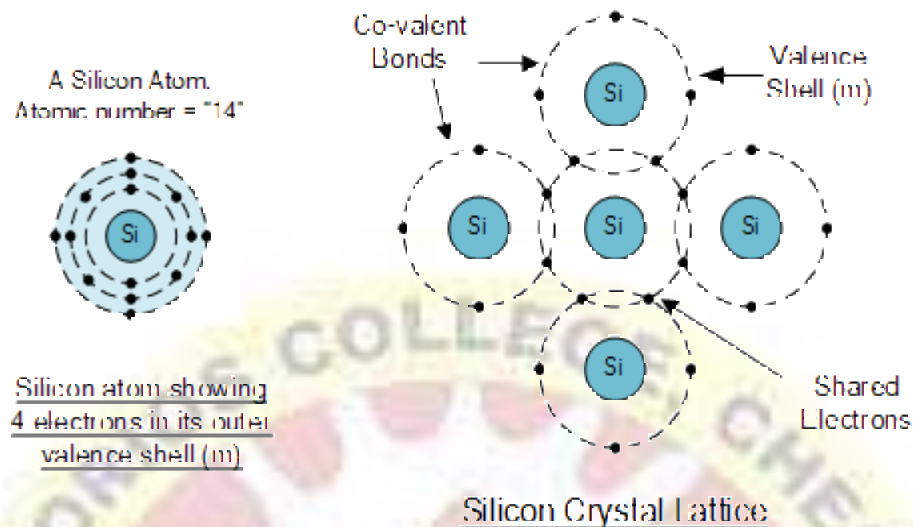
nucleus. In other words, the electrons are stuck to the parent atom and can not move around freely so if a potential voltage is applied to the material no current will flow as there are no “free electrons” available to move and which gives these materials their insulating properties. Insulators also have very high resistances, millions of ohms per metre, and are generally not affected by normal temperature changes (although at very high temperatures wood becomes charcoal and changes from an insulator to a conductor). Examples of good insulators are marble, fused quartz, PVC plastics, rubber etc. Insulators play a very important role within electrical and electronic circuits, because without them electrical circuits would short together and not work. For example, insulators made of glass or porcelain are used for insulating and supporting overhead transmission cables while epoxy-glass resin materials are used to make printed circuit boards, PCB's etc.

Semiconductor Basics

Semiconductors materials such as silicon (Si), germanium (Ge) and gallium arsenide (GaAs), have electrical properties somewhere in the middle, between those of a “conductor” and an “insulator”. They are not good conductors nor good insulators (hence their name “semi”-conductors). They have very few “free electrons” because their atoms are closely grouped together in a crystalline pattern called a “crystal lattice” but electrons are still able to flow, but only under special conditions. The ability of semiconductors to conduct electricity can be greatly improved by replacing or adding certain donor or acceptor atoms to this crystalline structure thereby, producing more free electrons than holes or vice versa. That is by adding a small percentage of another element to the base material, either silicon or germanium. On their own Silicon and Germanium are classed as intrinsic semiconductors, that is they are chemically pure, containing nothing but semi-conductive material. But by controlling the amount of impurities added to this intrinsic semiconductor material it is possible to control its conductivity. Various impurities called donors or acceptors can be added to this intrinsic material to produce free electrons or holes respectively. This process of adding donor or acceptor atoms to semiconductor atoms (the order of 1 impurity atom per 10 million (or more) atoms of the semiconductor) is called **Doping**. As the doped silicon is no longer pure, these donor and acceptor atoms are collectively referred to as “impurities”, and by doping these silicon material with a sufficient number of impurities, we can turn it into an N-type or P-type semi-conductor material. The most commonly used semiconductor basics material by far is **silicon**. Silicon has four valence electrons in its outermost shell which it shares with its neighbouring silicon atoms to form full orbital's of eight electrons. The structure of the bond between the two silicon atoms is such that each atom shares one electron with its neighbour making the bond very stable. As there are very few free electrons available to move around the silicon crystal, crystals of pure silicon (or germanium) are therefore good insulators, or at the very least very high value resistors. Silicon atoms are arranged in a definite symmetrical pattern making them a crystalline solid structure. A crystal of pure silica (silicon dioxide or glass) is generally said to be an intrinsic crystal (it has no impurities) and therefore has no free electrons. But simply connecting a silicon crystal to a battery supply is not enough to extract an electric current from it. To do that we need to create a “positive” and a

“negative” pole within the silicon allowing electrons and therefore electric current to flow out of the silicon. These poles are created by doping the silicon with certain impurities.

A Silicon Atom Structure



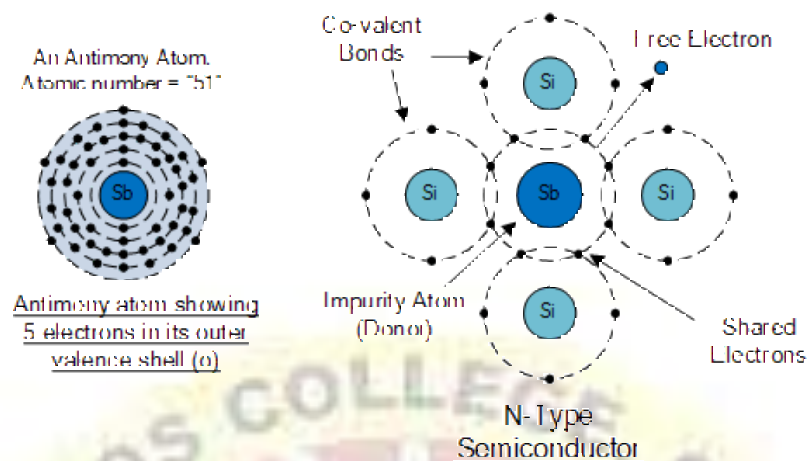
The diagram above shows the structure and lattice of a ‘normal’ pure crystal of Silicon.

N-type Semiconductor Basics

In order for our silicon crystal to conduct electricity, we need to introduce an impurity atom such as Arsenic, Antimony or Phosphorus into the crystalline structure making it extrinsic (impurities are added). These atoms have five outer electrons in their outermost orbital to share with neighbouring atoms and are commonly called “Pentavalent” impurities. This allows four out of the five orbital electrons to bond with its neighbouring silicon atoms leaving one “free electron” to become mobile when an electrical voltage is applied (electron flow). As each impurity atom “donates” one electron, pentavalent atoms are generally known as “donors”.

Antimony (symbol Sb) as well as **Phosphorus** (symbol P), are frequently used as a pentavalent additive to silicon. Antimony has 51 electrons arranged in five shells around its nucleus with the outermost orbital having five electrons. The resulting semiconductor basics material has an excess of current-carrying electrons, each with a negative charge, and is therefore referred to as an **N-type** material with the electrons called “Majority Carriers” while the resulting holes are called “Minority Carriers”. When stimulated by an external power source, the electrons freed from the silicon atoms by this stimulation are quickly replaced by the free electrons available from the doped Antimony atoms. But this action still leaves an extra electron (the freed electron) floating around the doped crystal making it negatively charged. Then a semiconductor material is classed as N-type when its donor density is greater than its acceptor density, in other words, it has more electrons than holes thereby creating a negative pole as shown.

Antimony Atom and Doping



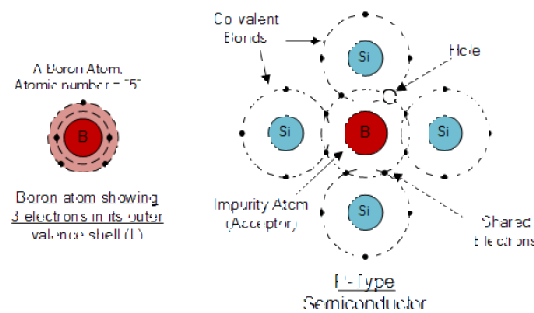
The diagram above shows the structure and lattice of the donor impurity atom Antimony.

P-Type Semiconductor Basics

If we introduce a “Trivalent” (3-electron) impurity into the crystalline structure, such as Aluminium, Boron or Indium, which have only three valence electrons available in their outermost orbital, the fourth closed bond cannot be formed. Therefore, a complete connection is not possible, giving the semiconductor material an abundance of positively charged carriers known as holes in the structure of the crystal where electrons are effectively missing. As there is now a hole in the silicon crystal, a neighbouring electron is attracted to it and will try to move into the hole to fill it. However, the electron filling the hole leaves another hole behind it as it moves. This in turn attracts another electron which in turn creates another hole behind it, and so forth giving the appearance that the holes are moving as a positive charge through the crystal structure (conventional current flow). This movement of holes results in a shortage of electrons in the silicon turning the entire doped crystal into a positive pole. As each impurity atom generates a hole, trivalent impurities are generally known as “**Acceptors**” as they are continually “accepting” extra or free electrons.

Boron (symbol B) is commonly used as a trivalent additive as it has only five electrons arranged in three shells around its nucleus with the outermost orbital having only three electrons. The doping of Boron atoms causes conduction to consist mainly of positive charge carriers resulting in a **P-type** material with the positive holes being called “Majority Carriers” while the free electrons are called “Minority Carriers”. Then a semiconductor basics material is classed as P-type when its acceptor density is greater than its donor density. Therefore, a P-type semiconductor has more holes than electrons.

Boron Atom and Doping



The diagram above shows the structure and lattice of the acceptor impurity atom Boron.

Semiconductor Basics Summary

N-type (e.g. doped with Antimony)

These are materials which have **Pentavalent** impurity atoms (Donors) added and conduct by “electron” movement and are therefore called, **N-type Semiconductors**.

In N-type semiconductors there are:

- 1. The Donors are positively charged.
- 2. There are a large number of free electrons.
- 3. A small number of holes in relation to the number of free electrons.
- 4. Doping gives:
 - positively charged donors.
 - negatively charged free electrons.
- 5. Supply of energy gives:
 - negatively charged free electrons.
 - positively charged holes.

P-type (e.g. doped with Boron)

These are materials which have **Trivalent** impurity atoms (Acceptors) added and conduct by “hole” movement and are therefore called, **P-type Semiconductors**.

In these types of materials are:

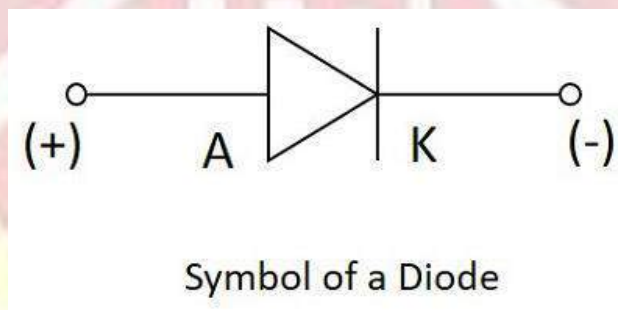
- 1. The Acceptors are negatively charged.
- 2. There are a large number of holes.
- 3. A small number of free electrons in relation to the number of holes.

- 4. Doping gives:
 - negatively charged acceptors.
 - positively charged holes.
- 5. Supply of energy gives:
 - positively charged holes.
 - negatively charged free electrons and both P and N-types as a whole, are electrically neutral on their own.

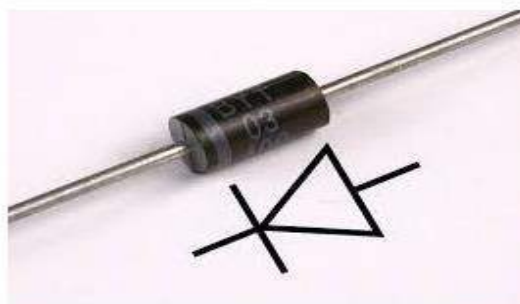
Antimony (Sb) and Boron (B) are two of the most commonly used doping agents as they are more feely available compared to other types of materials. They are also classed as “metalloids”. However, the periodic table groups together a number of other different chemical elements all with either three, or five electrons in their outermost orbital shell making them suitable as a doping material. These other chemical elements can also be used as doping agents to a base material of either Silicon (Si) or Germanium (Ge) to produce different types of basic semiconductor materials for use in electronic semiconductor components, microprocessor and solar cell applications.

PN – Junction Diode

A semiconductor diode is a two terminal electronic component with a PN junction. This is also called as a **Rectifier**.



The **anode** which is the **positive terminal** of a diode is represented with **A** and the **cathode**, which is the **negative terminal** is represented with **K**. To know the anode and cathode of a practical diode, a fine line is drawn on the diode which means cathode, while the other end represents anode.

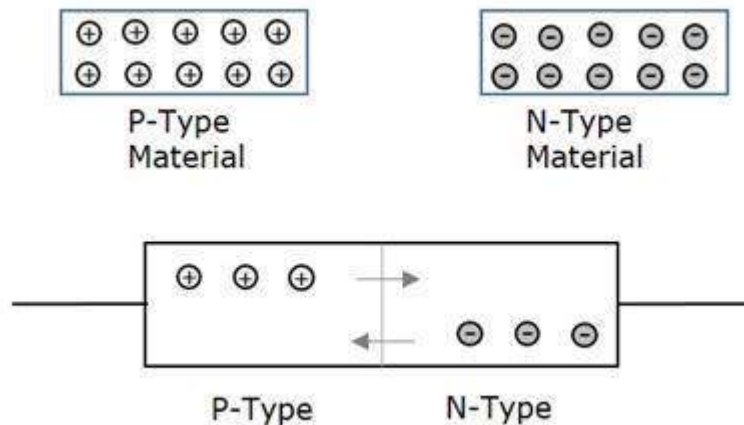


Representing anode and cathode of a practical diode through its symbol

As we had already discussed about the P-type and N-type semiconductors, and the behavior of their carriers, let us now try to join these materials together to see what happens.

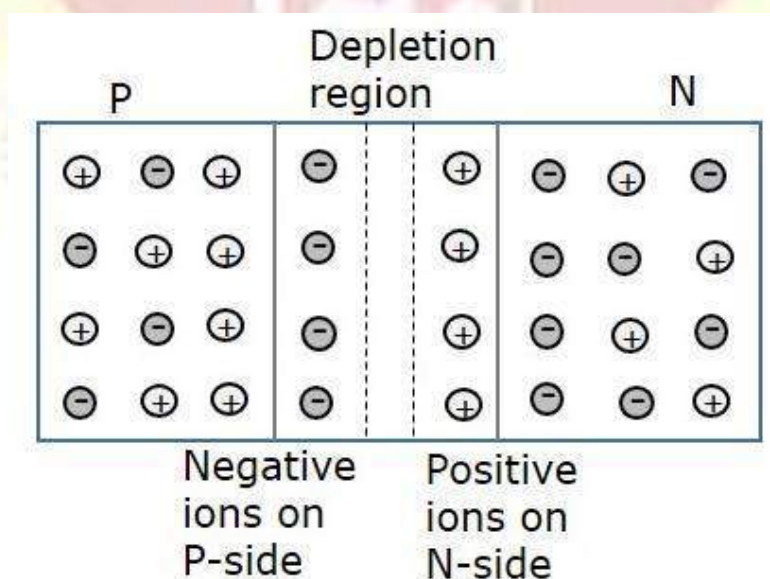
Formation of a Diode

If a P-type and an N-type material are brought close to each other, both of them join to form a junction, as shown in the figure below.



A P-type material has **holes** as the **majority carriers** and an N-type material has **electrons** as the **majority carriers**. As opposite charges attract, few holes in P-type tend to go to n-side, whereas few electrons in N-type tend to go to P-side.

As both of them travel towards the junction, holes and electrons recombine with each other to neutralize and forms ions. Now, in this junction, there exists a region where the positive and negative ions are formed, called as PN junction or junction barrier as shown in the figure.



The formation of negative ions on P-side and positive ions on N-side results in the formation of a narrow charged region on either side of the PN junction. This region is now free from movable charge carriers. The ions present here have been stationary and maintain a region of space between them without any charge

carriers. As this region acts as a barrier between P and N type materials, this is also called as **Barrier junction**. This has another name called as **Depletion region** meaning it depletes both the regions. There occurs a potential difference PD due to the formation of ions, across the junction called as **Potential Barrier** as it prevents further movement of holes and electrons through the junction.

Biassing of a Diode

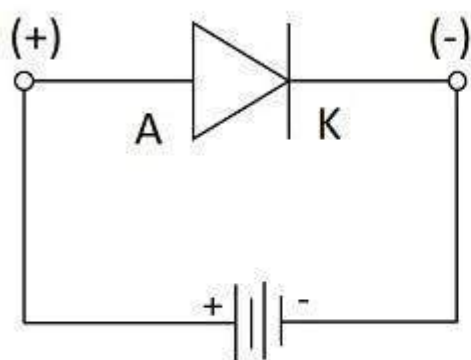
When a diode or any two-terminal component is connected in a circuit, it has two biased conditions with the given supply. They are **Forward biased** condition and **Reverse biased** condition.

Forward Biased Condition

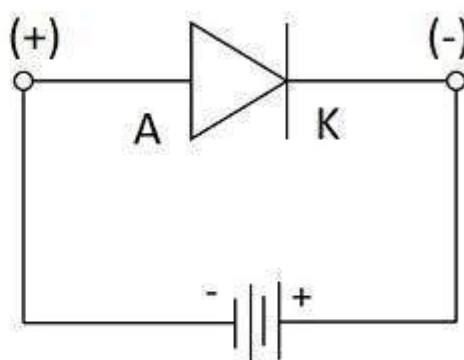
When a diode is connected in a circuit, with its **anode to the positive** terminal and **cathode to the negative** terminal of the supply, then such a connection is said to be **forward biased** condition. This kind of connection makes the circuit more and more forward biased and helps in more conduction. A diode conducts well in forward biased condition.

Reverse Biased Condition

When a diode is connected in a circuit, with its **anode to the negative** terminal and **cathode to the positive** terminal of the supply, then such a connection is said to be **Reverse biased** condition. This kind of connection makes the circuit more and more reverse biased and helps in minimizing and preventing the conduction. A diode cannot conduct in reverse biased condition.



Forward biased Connection

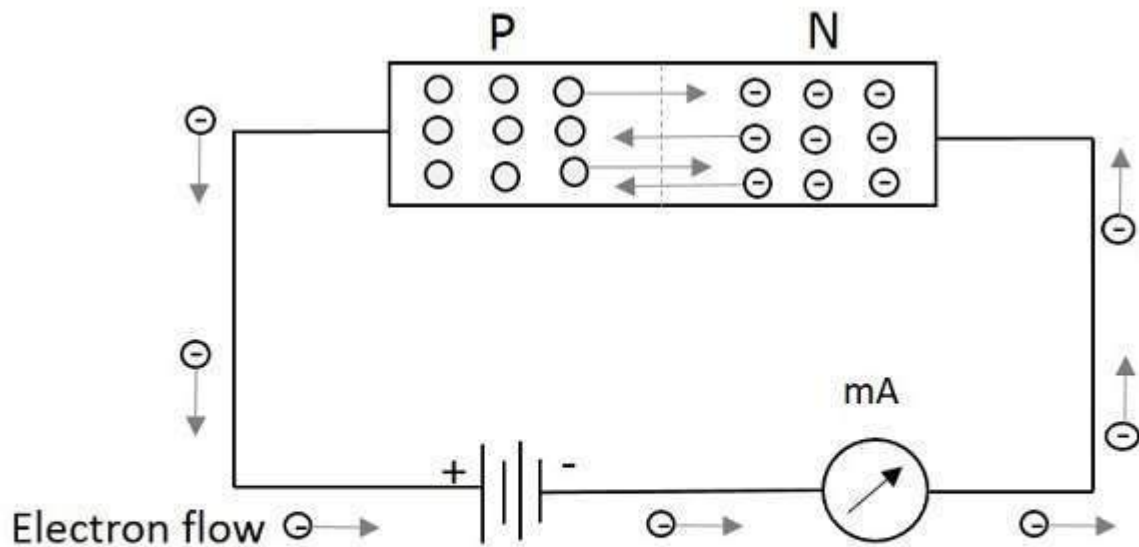


Reverse biased Connection

Let us now try to know what happens if a diode is connected in forward biased and in reverse biased conditions.

Working under Forward Biased

When an external voltage is applied to a diode such that it cancels the potential barrier and permits the flow of current is called as **forward bias**. When anode and cathode are connected to positive and negative terminals respectively, the holes in P-type and electrons in N-type tend to move across the junction, breaking the barrier. There exists a free flow of current with this, almost eliminating the barrier.



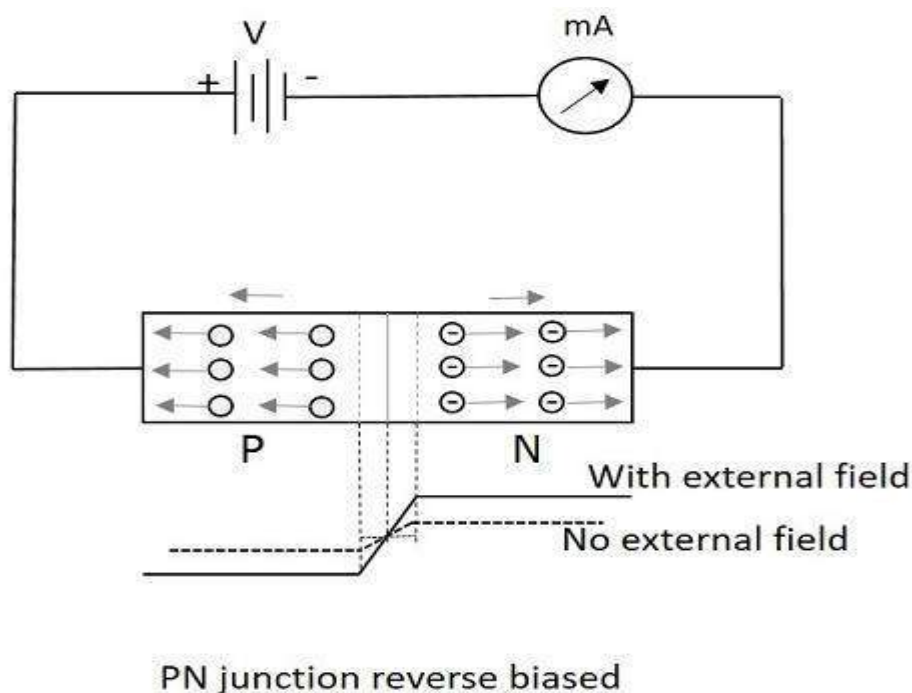
PN junction forward biased

With the repulsive force provided by positive terminal to holes and by negative terminal to electrons, the recombination takes place in the junction. The supply voltage should be such high that it forces the movement of electrons and holes through the barrier and to cross it to provide **forward current**. Forward Current is the current produced by the diode when operating in forward biased condition and it is indicated by I_f .

Working under Reverse Biased

When an external voltage is applied to a diode such that it increases the potential barrier and restricts the flow of current is called as **Reverse bias**. When anode and cathode are connected to negative and positive terminals respectively, the electrons are attracted towards the positive terminal and holes are attracted towards the negative terminal. Hence both will be away from the potential barrier **increasing the junction resistance** and preventing any electron to cross the junction.

The following figure explains this. The graph of conduction when no field is applied and when some external field is applied are also drawn.



With the increasing reverse bias, the junction has few minority carriers to cross the junction. This current is normally negligible. This reverse current is almost constant when the temperature is constant. But when this reverse voltage increases further, then a point called **reverse breakdown occurs**, where an avalanche of current flows through the junction. This high reverse current damages the device. **Reverse current** is the current produced by the diode when operating in reverse biased condition and it is indicated by I_r . Hence a diode provides high resistance path in reverse biased condition and doesn't conduct, where it provides a low resistance path in forward biased condition and conducts. Thus we can conclude that a diode is a one-way device which conducts in forward bias and acts as an insulator in reverse bias. This behavior makes it work as a rectifier, which converts AC to DC.

Peak Inverse Voltage

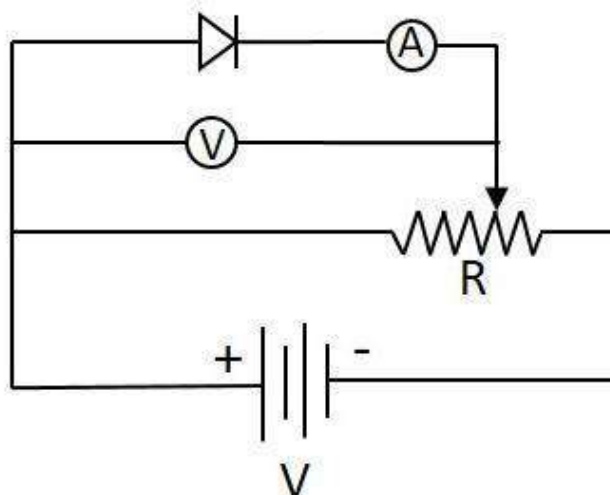
Peak Inverse Voltage is shortly called as **PIV**. It states the maximum voltage applied in reverse bias. The Peak Inverse Voltage can be defined as “**The maximum reverse voltage that a diode can withstand without being destroyed**”. Hence, this voltage is considered during reverse biased condition. It denotes how a diode can be safely operated in reverse bias.

Purpose of a Diode

A diode is used to block the electric current flow in one direction, i.e. in forward direction and to block in reverse direction. This principle of diode makes it work as a **Rectifier**. For a circuit to allow the current flow in one direction but to stop in the other direction, the rectifier diode is the best choice. Thus the **output** will be **DC** removing the AC components. The circuits such as half wave and full wave rectifiers are made using diodes. A diode is also used as a **Switch**. It helps a faster ON and OFF for the output that should occur in a quick rate.

V - I Characteristics of a Diode

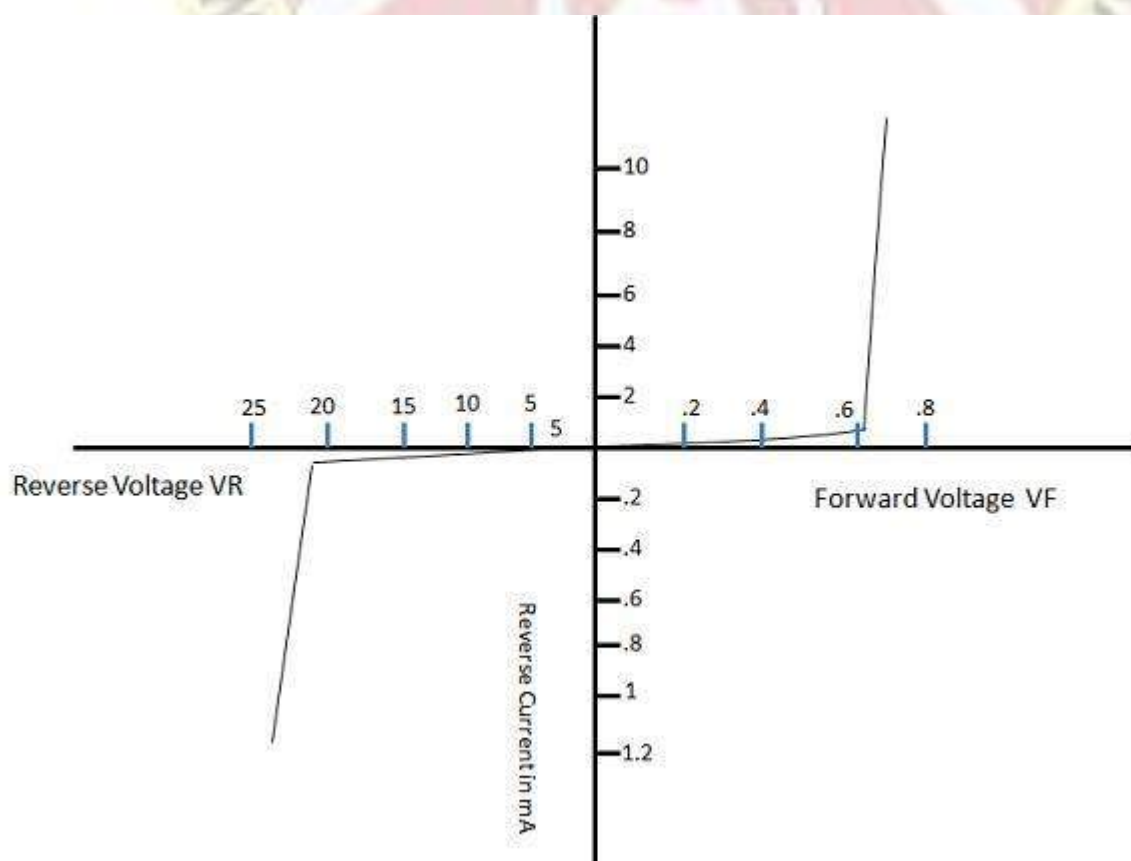
A Practical circuit arrangement for a PN junction diode is as shown in the following figure. An ammeter is connected in series and voltmeter in parallel, while the supply is controlled through a variable resistor.



A Practical diode circuit

During the operation, when the diode is in forward biased condition, at some particular voltage, the potential barrier gets eliminated. Such a voltage is called as **Cut-off Voltage** or **Knee Voltage**. If the forward voltage exceeds beyond the limit, the forward current rises up exponentially and if this is done further, the device is damaged due to overheating.

The following graph shows the state of diode conduction in forward and reverse biased conditions.



During the reverse bias, current produced through minority carriers exist known as “**Reverse current**”. As the reverse voltage increases, this reverse current increases and it suddenly breaks down at a point, resulting in the permanent destruction of the junction.

Forward Characteristic

When a diode is forward biased it conducts current (I_F) in forward direction. Following are the observations – Forward Biased

- Forward Voltage is measured across the diode and Forward Current is a measure of current through the diode.
- When the forward voltage across the diode equals 0V, forward current (I_F) equals 0 mA.
- When the value starts from the starting point (0) of the graph, if V_F is progressively increased in 0.1-V steps, I_F begins to rise.
- When the value of V_F is large enough to overcome the barrier potential of the P-N junction, a considerable increase in I_F occurs. The point at which this occurs is often called the knee voltage V_K . For germanium diodes, V_K is approximately 0.3 V, and 0.7 V for silicon.
- If the value of I_F increases much beyond V_K , the forward current becomes quite large.

This operation causes excessive heat to develop across the junction and can destroy a diode. To avoid this situation, a protective resistor is connected in series with the diode. This resistor limits the forward current to its maximum rated value. Normally, a currentlimiting resistor is used when diodes are operated in the forward direction.

Reverse Characteristic

When a diode is reverse biased, it conducts Reverse current that is usually quite small. A typical diode reverse IV characteristic is shown in the above figure.

The vertical reverse current line in this graph has current values expressed in microamperes. The amount of minority current carriers that take part in conduction of reverse current is quite small. In general, this means that reverse current remains constant over a large part of reverse voltage. When the reverse voltage of a diode is increased from the start, there is a very slight change in the reverse current. At the breakdown voltage (V_{BR}) point, current increases very rapidly. The voltage across the diode remains reasonably constant at this time.

This constant-voltage characteristic leads to a number of applications of diode under reverse bias condition. The processes which are responsible for current conduction in a reverse-biased diode are called as **Avalanche breakdown** and **Zener breakdown**.

Diode Specifications

Like any other selection, selection of a diode for a specific application must be considered. Manufacturer generally provides this type of information. Specifications like maximum voltage and current ratings, usual operating conditions, mechanical facts, lead identification, mounting procedures, etc.

Following are some of the important specifications –

- **Maximum forward current (IFM)** – The absolute maximum repetitive forward current that can pass through a diode.
- **Maximum reverse voltage (VRM)** – The absolute maximum or peak reverse bias voltage that can be applied to a diode.
- **Reverse breakdown voltage (VBR)** – The minimum steady-state reverse voltage at which breakdown will occur.
- **Maximum forward surge current (IFM-surge)** – The maximum current that can be tolerated for a short interval of time. This current value is much greater than IFM.
- **Maximum reverse current (IR)** – The absolute maximum reverse current that can be tolerated at device operating temperature.
- **Forward voltage (VF)** – Maximum forward voltage drop for a given forward current at device operating temperature.
- **Power dissipation (PD)** – The maximum power that the device can safely absorb on a continuous basis in free air at 25° C.
- **Reverse recovery time (Trr)** – The maximum time that it takes the device to switch from on to off state.

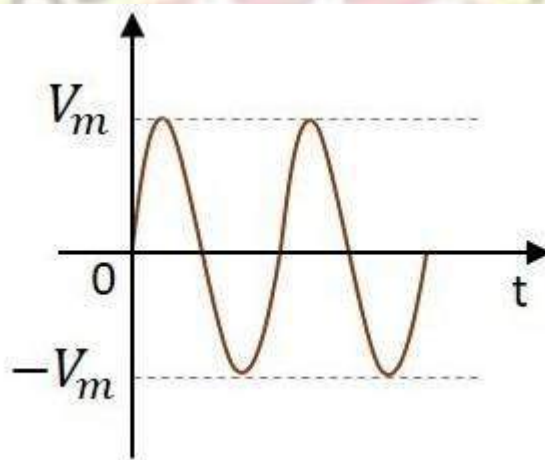
Important Terms

- **Breakdown Voltage** – It is the minimum reverse bias voltage at which PN junction breaks down with sudden rise in reverse current.
- **Knee Voltage** – It is the forward voltage at which the current through the junction starts to increase rapidly.
- **Peak Inverse Voltage** – It is the maximum reverse voltage that can be applied to the PN junction, without damaging it.

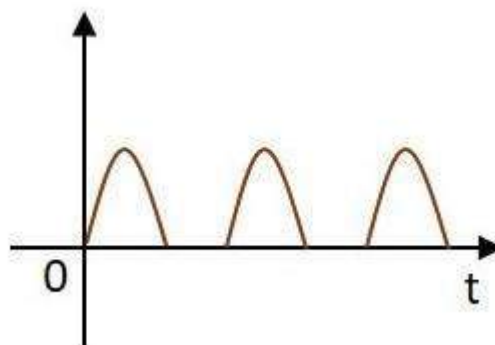
- **Maximum Forward Rating** – It is the highest instantaneous forward current that a PN junction can pass, without damaging it.
- **Maximum Power Rating** – It is the maximum power that can be dissipated from the junction, without damaging the junction.

Rectification

An alternating current has the property to change its state continuously. This is understood by observing the sine wave by which an alternating current is indicated. It raises in its positive direction goes to a peak positive value, reduces from there to normal and again goes to negative portion and reaches the negative peak and again gets back to normal and goes on.



During its journey in the formation of wave, we can observe that the wave goes in positive and negative directions. Actually it alters completely and hence the name alternating current. But during the process of rectification, this alternating current is changed into direct current DC. The wave which flows in both positive and negative direction till then, will get its direction restricted only to positive direction, when converted to DC. Hence the current is allowed to flow only in positive direction and resisted in negative direction, just as in the figure below.



The circuit which does rectification is called as a **Rectifier circuit**. A diode is used as a rectifier, to construct a rectifier circuit.

Types of Rectifier circuits

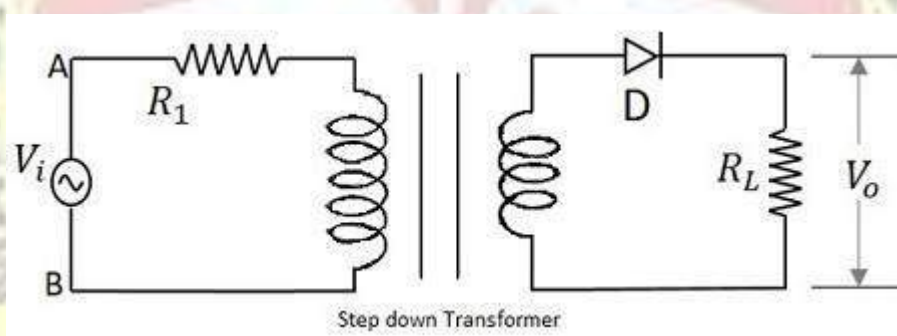
There are two main types of rectifier circuits, depending upon their output. They are

- Half-wave Rectifier
- Full-wave Rectifier

A Half-wave rectifier circuit rectifies only positive half cycles of the input supply whereas a Full-wave rectifier circuit rectifies both positive and negative half cycles of the input supply.

Half-Wave Rectifier

The name half-wave rectifier itself states that the **rectification** is done only for **half** of the cycle. The AC signal is given through an input transformer which steps up or down according to the usage. Mostly a step down transformer is used in rectifier circuits, so as to reduce the input voltage. The input signal given to the transformer is passed through a PN junction diode which acts as a rectifier. This diode converts the AC voltage into pulsating dc for only the positive half cycles of the input. A load resistor is connected at the end of the circuit. The figure below shows the circuit of a half wave rectifier.

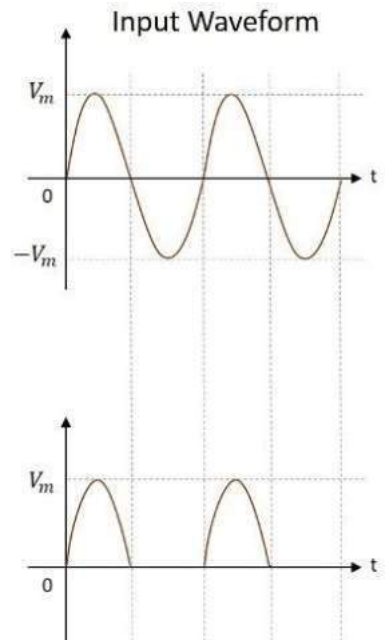


Working of a HWR

The input signal is given to the transformer which reduces the voltage levels. The output from the transformer is given to the diode which acts as a rectifier. This diode gets ON conducts for positive half cycles of input signal. Hence a current flows in the circuit and there will be a voltage drop across the load resistor. The diode gets OFF doesn't conduct for negative half cycles and hence the output for negative half cycles will be, $V_o=0$. Hence the output is present for positive half cycles of the input voltage only neglecting the reverse leakage current. This output will be pulsating which is taken across the load resistor.

Waveforms of a HWR

The input and output waveforms are as shown in the following figure.



Hence the output of a half wave rectifier is a pulsating dc. Let us try to analyze the above circuit by understanding few values which are obtained from the output of half wave rectifier.

Full-Wave Rectifier

A Rectifier circuit that rectifies both the positive and negative half cycles can be termed as a full wave rectifier as it rectifies the complete cycle. The construction of a full wave rectifier can be made in two types. They are

- Center-tapped Full wave rectifier
- Bridge full wave rectifier

Both of them have their advantages and disadvantages. Let us now go through both of their construction and working along with their waveforms to know which one is better and why.

Center-tapped Full-Wave Rectifier

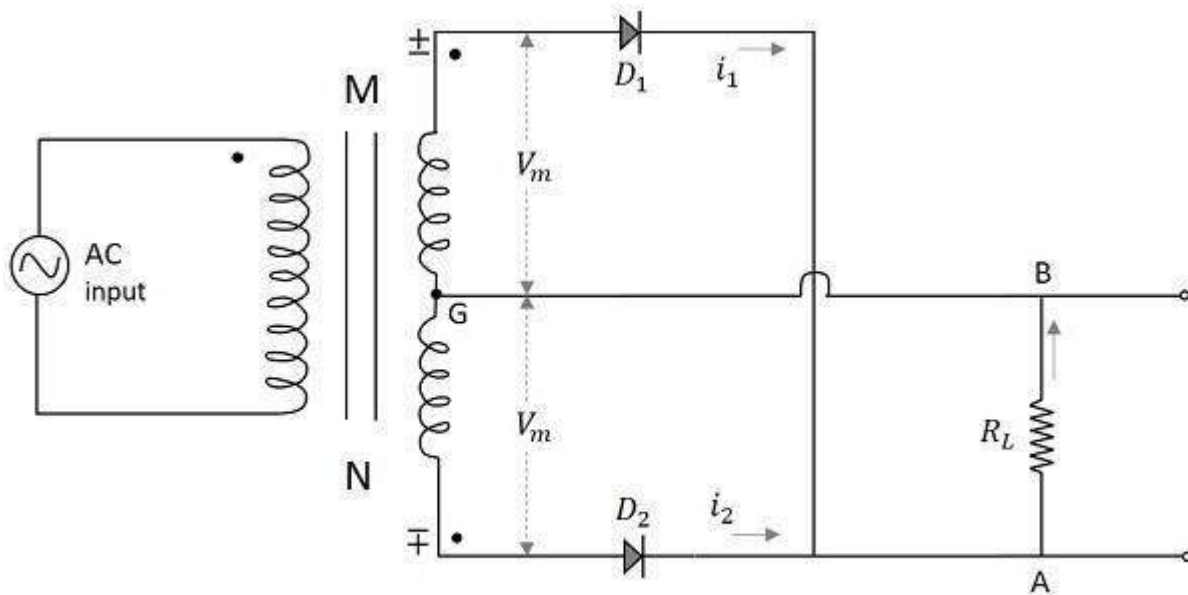
A rectifier circuit whose transformer secondary is tapped to get the desired output voltage, using two diodes alternatively, to rectify the complete cycle is called as a **Center-tapped Full wave rectifier circuit**. The transformer is center tapped here unlike the other cases.

The features of a center-tapping transformer are –

- The tapping is done by drawing a lead at the mid-point on the secondary winding. This winding is split into two equal halves by doing so.
- The voltage at the tapped mid-point is zero. This forms a neutral point.
- The center tapping provides two separate output voltages which are equal in magnitude but opposite in polarity to each other.

- A number of tapings can be drawn out to obtain different levels of voltages.

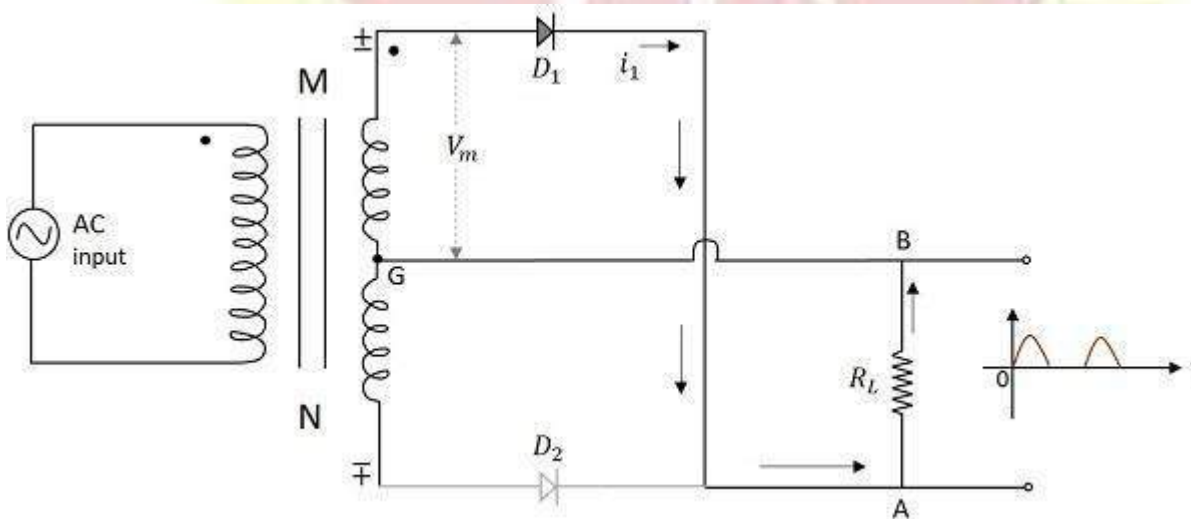
The center-tapped transformer with two rectifier diodes is used in the construction of a **Center-tapped full wave rectifier**. The circuit diagram of a center tapped full wave rectifier is as shown below.



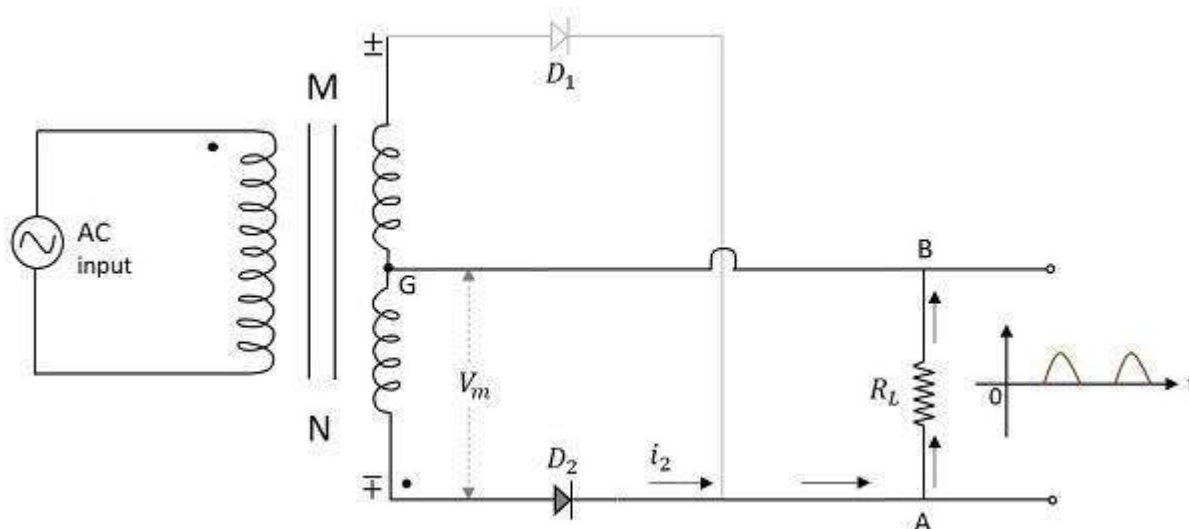
Circuit diagram of a center-tapped full wave rectifier

Working of a CT- FWR

The working of a center-tapped full wave rectifier can be understood by the above figure. When the positive half cycle of the input voltage is applied, the point M at the transformer secondary becomes positive with respect to the point N. This makes the diode D_1 forward biased. Hence current i_1 flows through the load resistor from A to B. We now have the positive half cycles in the output

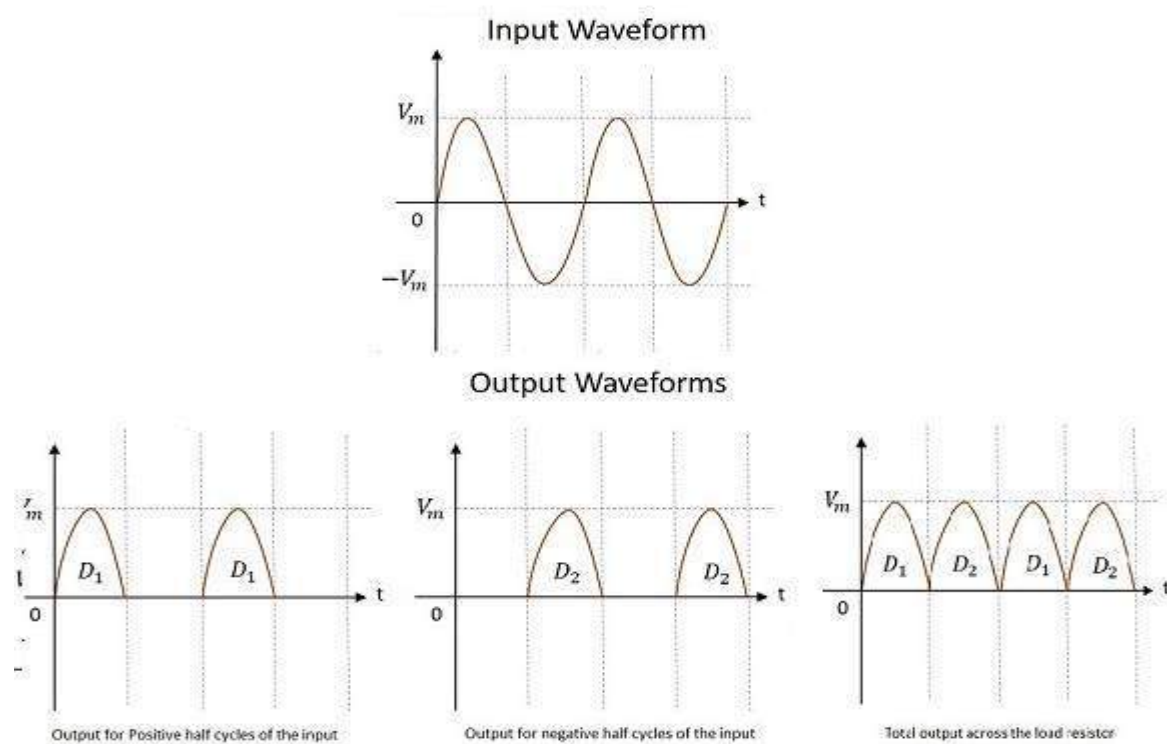


When the negative half cycle of the input voltage is applied, the point M at the transformer secondary becomes negative with respect to the point N. This makes the diode D_2 forward biased. Hence current i_2 flows through the load resistor from A to B. We now have the positive half cycles in the output, even during the negative half cycles of the input.



Waveforms of CT FWR

The input and output waveforms of the center-tapped full wave rectifier are as follows.



From the above figure it is evident that the output is obtained for both the positive and negative half cycles. It is also observed that the output across the load resistor is in the **same direction** for both the half cycles.

Peak Inverse Voltage

As the maximum voltage across half secondary winding is V_m , the whole of the secondary voltage appears across the non-conducting diode. Hence the **peak inverse voltage** is twice the maximum voltage across the half-secondary winding, i.e.

$$PIV = 2V_m$$

Disadvantages

There are few disadvantages for a center-tapped full wave rectifier such as –

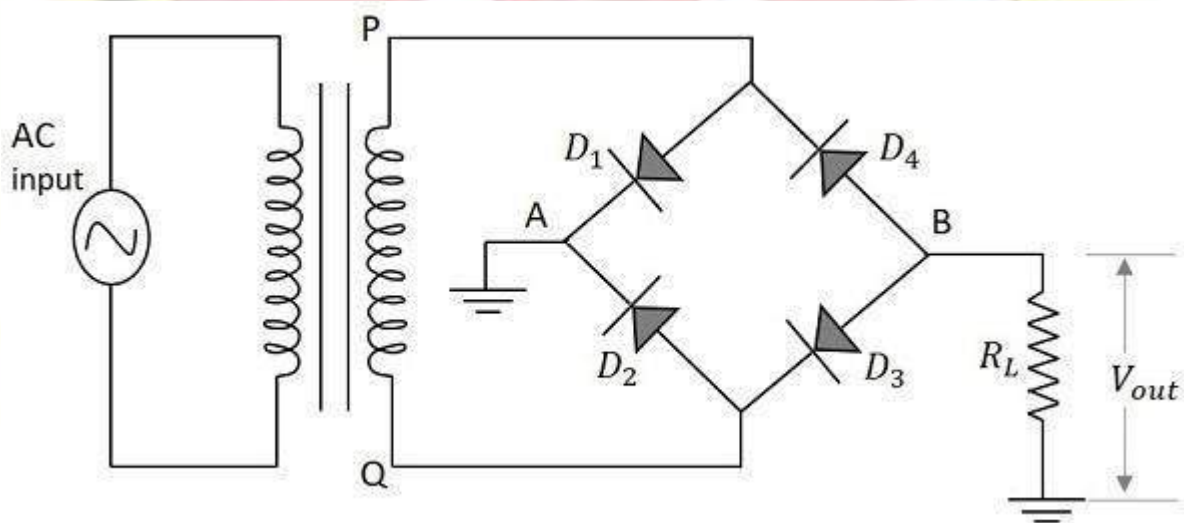
- Location of center-tapping is difficult
- The dc output voltage is small
- PIV of the diodes should be high

The next kind of full wave rectifier circuit is the **Bridge Full wave rectifier circuit**.

Bridge Full-Wave Rectifier

This is such a full wave rectifier circuit which utilizes four diodes connected in bridge form so as not only to produce the output during the full cycle of input, but also to eliminate the disadvantages of the center-tapped full wave rectifier circuit.

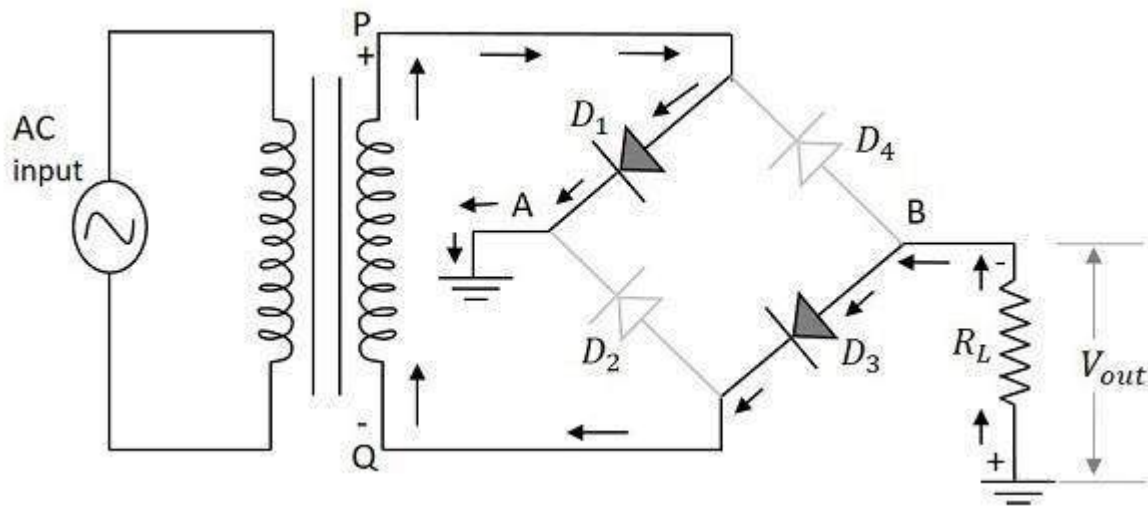
There is no need of any center-tapping of the transformer in this circuit. Four diodes called D_1 , D_2 , D_3 and D_4 are used in constructing a bridge type network so that two of the diodes conduct for one half cycle and two conduct for the other half cycle of the input supply. The circuit of a bridge full wave rectifier is as shown in the following figure.



Working of a Bridge Full-Wave Rectifier

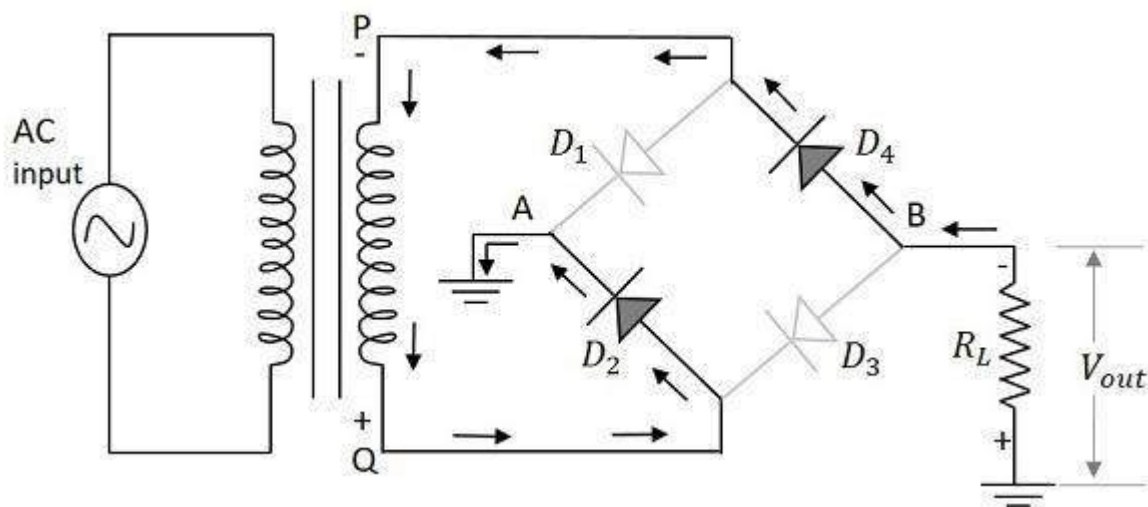
The full wave rectifier with four diodes connected in bridge circuit is employed to get a better full wave output response. When the positive half cycle of the input supply is given, point P becomes positive with respect to the point Q. This makes the diode D_1 and D_3 forward biased while D_2 and D_4 reverse biased. These two diodes will now be in series with the load resistor.

The following figure indicates this along with the conventional current flow in the circuit.



Hence the diodes D_1 and D_3 conduct during the positive half cycle of the input supply to produce the output along the load resistor. As two diodes work in order to produce the output, the voltage will be twice the output voltage of the center tapped full wave rectifier. When the negative half cycle of the input supply is given, point P becomes negative with respect to the point Q. This makes the diode D_1 and D_3 reverse biased while D_2 and D_4 forward biased. These two diodes will now be in series with the load resistor.

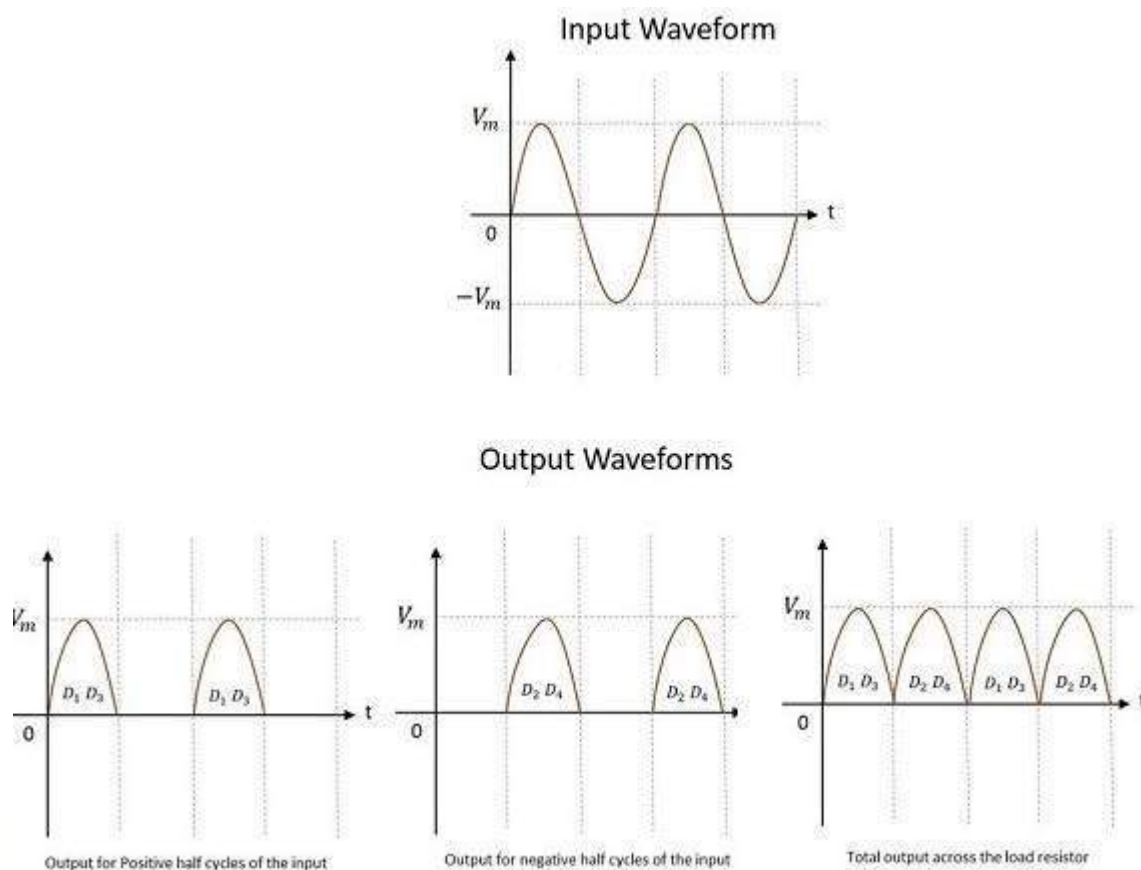
The following figure indicates this along with the conventional current flow in the circuit.



Hence the diodes D_2 and D_4 conduct during the negative half cycle of the input supply to produce the output along the load resistor. Here also two diodes work to produce the output voltage. The current flows in the same direction as during the positive half cycle of the input.

Waveforms of Bridge FWR

The input and output waveforms of the center-tapped full wave rectifier are as follows.



From the above figure, it is evident that the output is obtained for both the positive and negative half cycles. It is also observed that the output across the load resistor is in the **same direction** for both the half cycles.

Peak Inverse Voltage

Whenever two of the diodes are being in parallel to the secondary of the transformer, the maximum secondary voltage across the transformer appears at the non-conducting diodes which makes the PIV of the rectifier circuit. Hence the **peak inverse voltage** is the maximum voltage across the secondary winding, i.e.

$$PIV = V_m$$

Advantages

There are many advantages for a bridge full wave rectifier, such as –

- No need of center-tapping.
- The dc output voltage is twice that of the center-tapper FWR.
- PIV of the diodes is of the half value that of the center-tapper FWR.
- The design of the circuit is easier with better output.

Let us now analyze the characteristics of a full-wave rectifier.

Half-Wave vs. Full-Wave Rectifier

After having gone through all the values of different parameters of the full wave rectifier, let us just try to compare and contrast the features of half-wave and full-wave rectifiers.

Terms	Half Wave Rectifier	Center Tapped FWR	Bridge FWR
Number of Diodes	1	2	4
Transformer tapping	No	Yes	No
Peak Inverse Voltage	V_m	$2V_m$	V_m
Maximum Efficiency	40.6%	81.2%	81.2%
Average / dc current	I_m/π	$2I_m/\pi$	$2I_m/\pi$
DC voltage	V_m/π	$2V_m/\pi$	$2V_m/\pi$
RMS current	$I_m/2$	$I_m/2$	$I_m/2$
Ripple Factor	1.211	0.480	0.480
Output frequency	f_{in}	$2f_{in}$	$2f_{in}$

Ripple factor : The output of a rectifier consists of a d.c. component as well as an a.c. component which is also known as ripple. This a.c. component is undesirable and accounts for the pulsations in the rectifier output. The effectiveness of a rectifier depends upon the magnitude of a.c. component in the output. Hence the smaller this component, the more effective is the rectifier.

Formula and Derivation

Ripple factor is given in terms of RMS value of ac component to RMS value of dc component.

$$\begin{aligned}
 \text{Ripple Factor, } \gamma &= \frac{\sqrt{(I_{rms})^2 - (I_{dc})^2}}{I_{dc}} \\
 &= \frac{\sqrt{(V_{rms})^2 - (V_{dc})^2}}{V_{dc}}
 \end{aligned}$$

So now we derive the formula of ripple factor. Derivation of ripple factor can be easily derived by the definition of ripple factor. As per definition we know, ripple factor is the ratio of rms of ac component to rms of dc components in rectified output.

By I_{rms} and I_{dc} we can find the ripple factor of the rectifier.

$$I_{rms} = \sqrt{I_{dc}^2 + I_{ac}^2}$$

or

$$I_{ac} = \sqrt{I_{rms}^2 - I_{dc}^2}$$

Dividing throughout by I_{dc} , we get,

$$\frac{I_{ac}}{I_{dc}} = \frac{1}{I_{dc}} \sqrt{I_{rms}^2 - I_{dc}^2}$$

But I_{ac}/I_{dc} is the ripple factor.

$$\therefore \text{Ripple factor} = \frac{1}{I_{dc}} \sqrt{I_{rms}^2 - I_{dc}^2} = \sqrt{\left(\frac{I_{rms}}{I_{dc}}\right)^2 - 1}$$

Here now we find ripple factor for half wave and full wave rectifier.

Ripple factor for half wave rectifier

For half wave rectification,

$$I_{rms} = I_m/2$$

$$I_{dc} = I_m/\pi$$

$$\text{Ripple factor} = \sqrt{\left(\frac{I_m/2}{I_m/\pi}\right)^2 - 1} = 1.21$$

Ripple factor of half wave rectifier is about 1.21 by the derivation. As per you can see output voltage has much more AC component in DC output voltage so the half-wave rectifier is ineffective in the conversion of A.C to D.C.

Ripple factor for full wave rectifier

For full wave rectifier,

$$I_{rms} = I_m/\sqrt{2}$$

$$I_{dc} = 2I_m/\pi$$

$$I_{rms} = \frac{I_m}{\sqrt{2}} \quad ; \quad I_{dc} = \frac{2 I_m}{\pi}$$

$$\therefore \text{Ripple factor} = \sqrt{\left(\frac{I_m/\sqrt{2}}{2 I_m/\pi}\right)^2 - 1} = 0.48$$

$$i.e. \frac{\text{effective a.c. component}}{\text{d.c. component}} = 0.48$$

This shows that in the output of a full-wave rectifier, the d.c. component is more than the a.c. component. Consequently, the pulsations in the output will be less than in half-wave rectifier. For this reason, full-wave rectification is invariably used for conversion of a.c. into d.c.

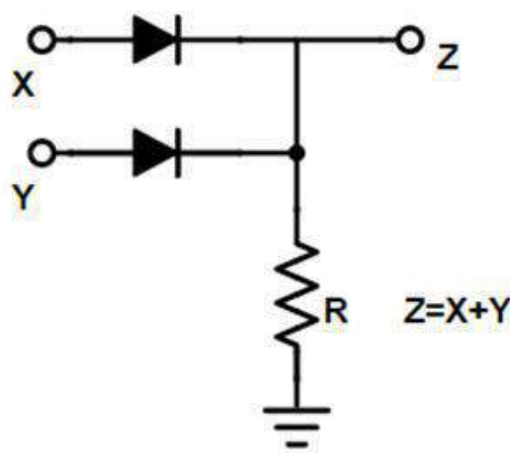
Conclusion

Ripple is the fluctuation in output of the rectifier and ripple factor is necessary for measuring the fluctuation rate in rectified output. By the uses of some capacitive filter and other filters, we can reduce the ripple in output voltage. In most of every rectifier circuit uses capacitor in parallel of diodes or thyristor which works as a filter in circuit. This capacitor helps to reduce the ripple in the output of the rectifier. Here we also saw the ripple factor of half wave and full wave rectifier.

Construction of gates using diodes

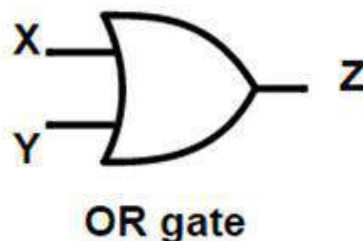
OR Gate

The circuit design of OR gate (by using diodes) is given below. The circuit uses two diodes at input side. In all logic circuits, +5 volts is represented as HIGH level logic and 0 volts or ground is represented as the LOW level logic. The two logic levels are represented as binary numbers 0 and 1. Here 0 means low (switch is at OFF position) and 1 means high (switch is at ON position). The OR gate produces low output only when both the inputs are low, in all other conditions or for all other combinations of inputs, the output of OR gate will be high. If we connect HIGH logic level voltage i.e. +5 volts to any one of the input, then the diode becomes forward biased and that makes the output to be high. Or if we connect high level voltage to both inputs, then both the inputs will be high making the diodes to be in forward bias condition. Then the OR circuit output will be HIGH. If the two inputs of the OR gate are low, then the inputs diodes become reverse biased, allowing no current to flow in the circuit. So the output will be LOW (0).



OR gate Logic Symbol and Boolean expression

The OR gate is logically represented as shown below with two inputs and one outputs.



Boolean expression

The mathematical functioning of OR gate is given as $Z = X + Y$. Here X and Y are the inputs and Z is the output of OR gate. The truth table for logical OR gate is given below.

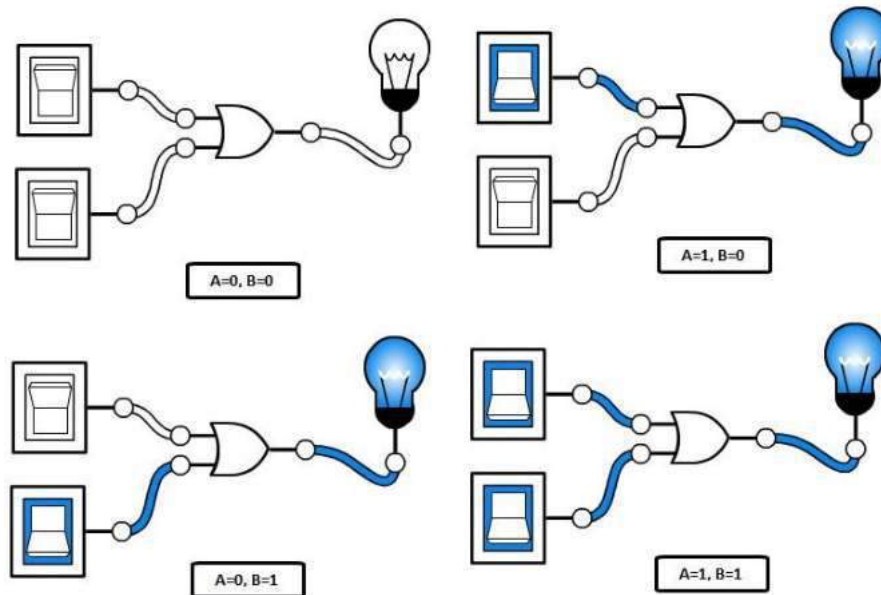
Inputs		Outputs
X	Y	Z
0	0	0
0	1	1
1	0	1
1	1	1

This shows us that the output of OR gate will be high with all combinations of inputs except for both low inputs condition.

Explanation of OR gate with light switch circuit

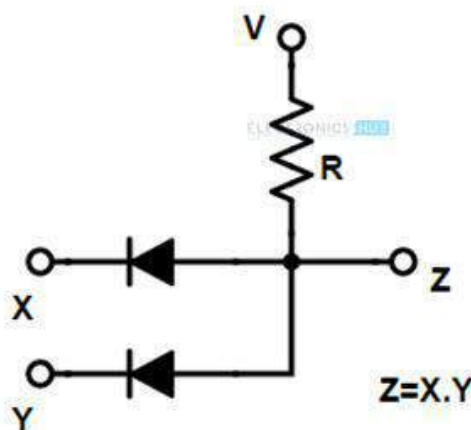
The OR gate switching circuit will have two inputs connected to two manually toggled switches. Let the two switches be A and B, then we can explain the switching operation of OR gate as

- When both switches A and B are open (switches are supplied with low level input signal) i.e. $A=0$, $B=0$, then the bulb will not glow.
- When switch A is close (supplied with high level input signal) and B is open (supplied with low level input signal) i.e. $A=1$, $B=0$, then the bulb will glow.
- When switch A is open (supplied with low level input signal) and B is close (supplied with high level input signal) i.e. $A=0$, $B=1$, then the bulb will glow.
- When both switches A and B are close (switches are supplied with high level input signal) i.e. $A=1$, $B=1$, then the bulb will glow.



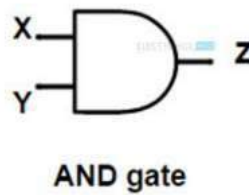
AND Gate

AND gate can be designed by using two simple diodes. The circuit driving voltage V is applied to the parallelly connected diodes and the output is collected as the voltage drop at the diodes. In logic gates, the terms high voltage level means +5 V and low logic level means 0 V or ground. When one of the inputs of the AND gate is connected to logic HIGH and other is connected to logic LOW then the diodes are at reverse bias condition and no voltage drop at the output. So the output is measured as LOW. If the two inputs are connected to LOW level input, then also the diodes will turn to reverse biased condition and allows no current. So again the output is measured as 0. But when the two inputs (two diodes) are connected to a HIGH voltage level, then the two diodes are in forward biased condition (diodes switches are ON) so the output of the AND gate is HIGH, and measured as logic 1.



AND gate Logic Symbol and Boolean expression

The AND gate is logically represented as shown below with two inputs and one output.



Boolean expression

If the inputs of the AND gate is X, Y and the output is Z then the operation of AND gate is mathematically expressed in the Boolean expression as $Z = X \cdot Y$. This means the AND gate produces the multiplication of its inputs.

Truth table

The truth table for logical AND gate is given below.

Inputs		Outputs
X	Y	Z
0	0	0
0	1	0
1	0	0
1	1	1

The truth table describes us that the output of AND gate will be LOW with all combinations of inputs except for both high inputs condition.

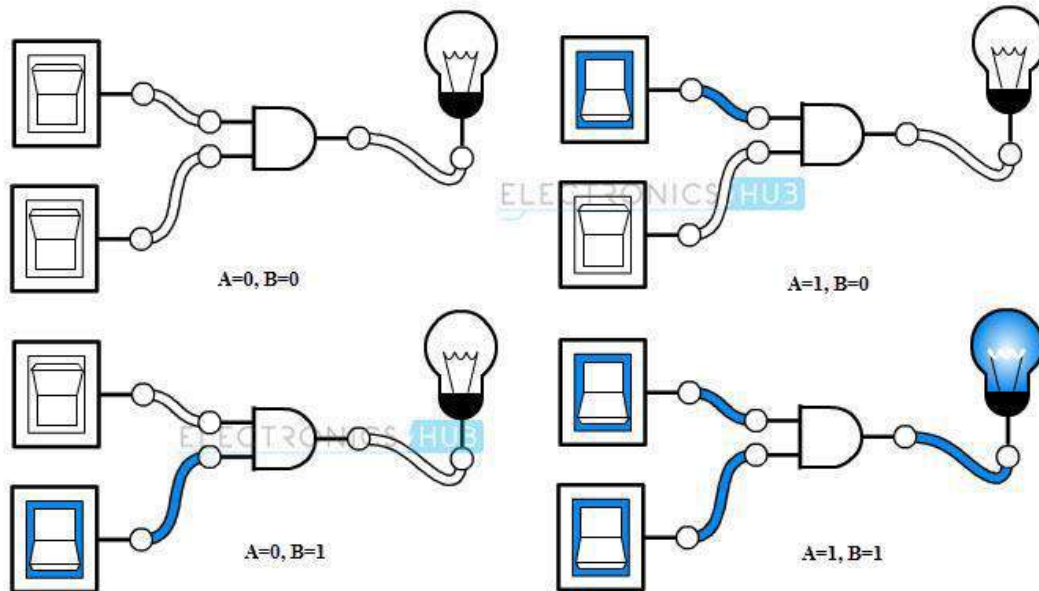
Explanation of AND gate with light switch circuit

The AND gate switching circuit will have two inputs with two manually toggled switches. Let the two switches be A and B, then the we can explain the switching operation of AND gate as

- When both switches A and B are open (switches are supplied with low level input signal) i.e. $A=0$, $B=0$, then the bulb will not glow.
- When switch A is close (supplied with high level input signal) and B is open (supplied with low level input signal) i.e. $A=1$, $B=0$, then the bulb will not glow.

- When switch A is open (supplied with low level input signal) and B is close (supplied with high level input signal) i.e. $A=0, B=1$, then the bulb will not glow.
- When both switches A and B are close (switches are supplied with high level input signal) i.e. $A=1, B=1$, then the bulb will glow.

This operation of AND gate as light switching circuit is described in below pictures

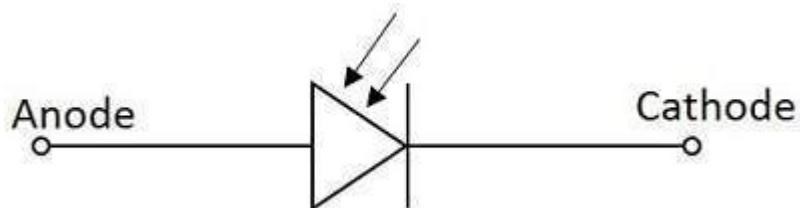


UNIT - 2

Photo Diode

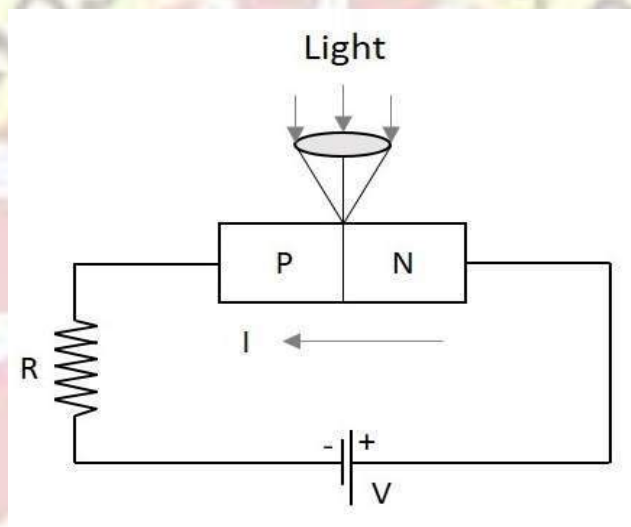
Photo diode, as the name implies, is a PN junction which works on light. The intensity of light affects the level of conduction in this diode. The photo diode has a P- type material and an N-type material with an **intrinsic** material or a **depletion region** in between. This diode is generally operated in **reverse bias** condition. The light when focused on the depletion region, electron-hole pairs are formed and flow of electron occurs. This conduction of electrons depends upon the intensity of light focused.

The figure below indicates the symbol for a photodiode.



Working of a Photo Diode

It is a **reverse-biased diode**. Reverse current increases as the intensity of incident light increases. This means that reverse current is directly proportional to the intensity of falling light. It consists of a PN junction mounted on a P-type substrate and sealed in a metallic case. The junction point is made of transparent lens and it is the window where the light is supposed to fall. As we know, when PN junction diode is reverse biased, a very small amount of reverse current flows. The reverse current is generated thermally by electron-hole pairs in the depletion region of the diode. When light falls on PN junction, it is absorbed by the junction. This will generate more electron-hole pairs. Or we can say, characteristically, the amount of reverse current increases. In other words, as the intensity of falling light increases, resistance of the PN junction diode decreases.



The Photo diode is encapsulated in a glass package to allow the light to fall onto it. In order to focus the light exactly on the depletion region of the diode, a lens is placed above the junction, just as illustrated above. Even when there is no light, a small amount of current flows which is termed as **Dark Current**. By changing the illumination level, reverse current can be changed.

Advantages of Photo diode

Photo diode has many advantages such as

- Low noise
- High gain
- High speed operation
- High sensitivity to light
- Low cost
- Small size
- Long lifetime

Applications of Photo diode

There are many applications for photo diode such as –

- Character detection
- Objects can be detected visible or invisible.
- Used in circuits that require high stability and speed.
- Used in Demodulation
- Used in switching circuits
- Used in Encoders
- Used in optical communication equipment

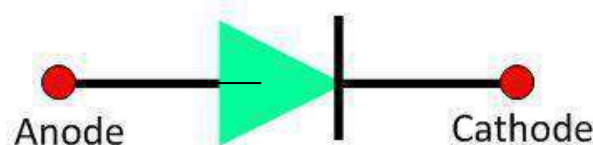
PIN Diode

The diode in which the intrinsic layer of high resistivity is sandwiched between the P and N-region of semiconductor material such type of diode is known as the PIN diode. The high resistive layer of the intrinsic region provides the large electric field between the P and N-region. The electric field is induced because of the movement of the holes and the electrons. The direction of the electric field is from n-region to p-region. The high electric field generates the large electron holes pairs due to which the diode process even for the small signals. The PIN diode is a type of photodetector used for converting the light energy into the electrical energy.

The intrinsic layer between the P and N-type regions increases the distance between them. The width of the region is inversely proportional to their capacitance. If the separation between the P and N region increases their capacitance decreases. This characteristic of diode increases their response time and makes the diode suitable for works like a microwaves applications.

Symbol of PIN Diode

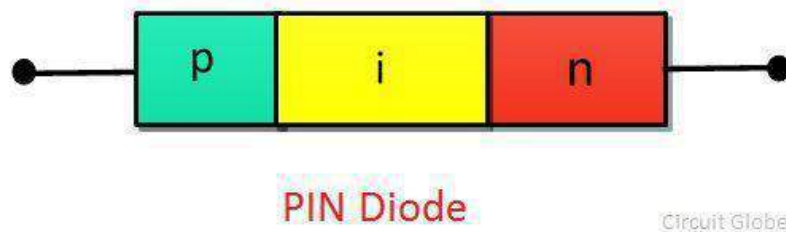
The symbolic representation of the PIN diode is shown in the figure below. The anode and cathode are the two terminal of the PIN diode. The anode is the positive terminal and cathode represent their negative terminals.



Symbol of PIN Diode

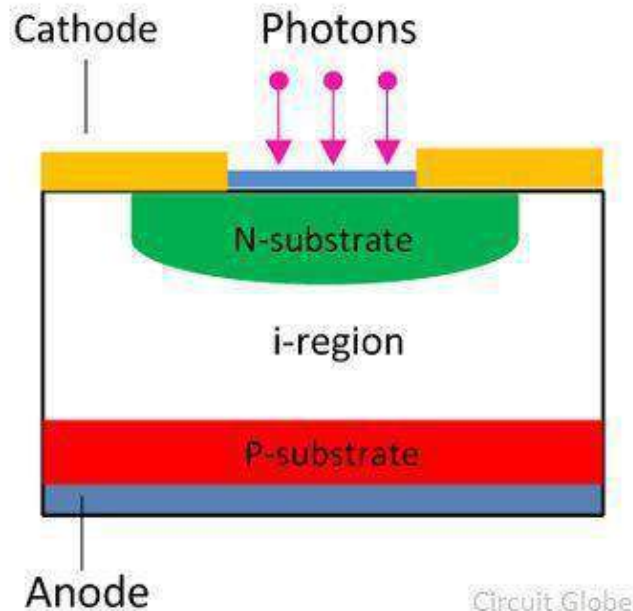
PIN Diode Structure

The diode consists the P-region and N-region which is separated by the intrinsic semiconductor material. In P-region the hole is the majority charge carrier while in n-region the electron is the majority charge carrier. The intrinsic region has no free charge carrier. It acts as an insulator between n and the p-type region. The i-region has the high resistance which obstructs the flow of electrons to pass through it.



Working of PIN Diode

The working of the PIN diode is similar to the ordinary diode. When the diode is unbiased, their charge carrier will diffuse. The word diffusion means the charge carriers of the depletion region try to move to their region. The process of diffusion occurs continue until the charges become equilibrium in the depletion region.



Let the N and I-layer make the depletion region. The diffusion of the hole and electron across the region generates the depletion layer across the NI-region. The thin depletion layer induces across n-region, and thick depletion region of opposite polarity induces across the I-region.

Forward Biased PIN Diode

When the diode is kept forward biased, the charges are continuously injected into the I-region from the P and N-region. This reduces the forward resistance of the diode, and it behaves like a variable resistance. The charge carrier which enters from P and N-region into the i-region are not immediately combined into the intrinsic region. The finite quantity of charge stored in the intrinsic region decreases their resistivity. Let Q be the quantity of charge stored in the depletion region. Let τ be the time used for the recombination of the charges. The quantity of the charges stored in the intrinsic region depends on their recombination time. The forward current starts flowing into the I region.

Reversed Biased PIN Diode

When the reverse voltage is applied across the diode, the width of the depletion region increases. The thickness of the region increases until the entire mobile charge carrier of the I-region swept away from it. **The reverse voltage requires for removing the complete charge carrier from the I-region is known as the swept voltage.**

In reverse bias, the diode behaves like a capacitor. The P and N region acts as the positive and negative plates of the capacitor, and the intrinsic region is the insulator between the plates.

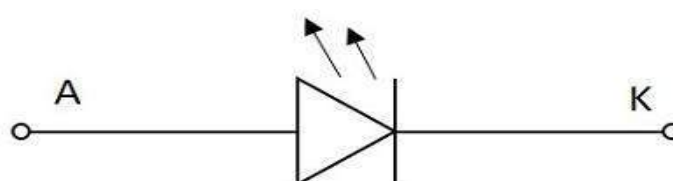
Applications of PIN Diode

- **High Voltage Rectifier** – It is used as a high voltage rectifier. The diode has a large intrinsic region between the N and P-region which can tolerate the high reverse voltage.
- **Photo-detector** – The PIN diode is used for converting the light energy into the electrical energy. The diode has large depletion region which improves their performance by increasing the volume of light conversion.

LED Light Emitting Diodes

It is the most popular diodes used in our daily life. This is also a normal PN junction diode except that instead of silicon and germanium, the materials like gallium arsenide, gallium arsenide phosphide are used in its construction.

The figure below shows the symbol of a Light emitting diode.



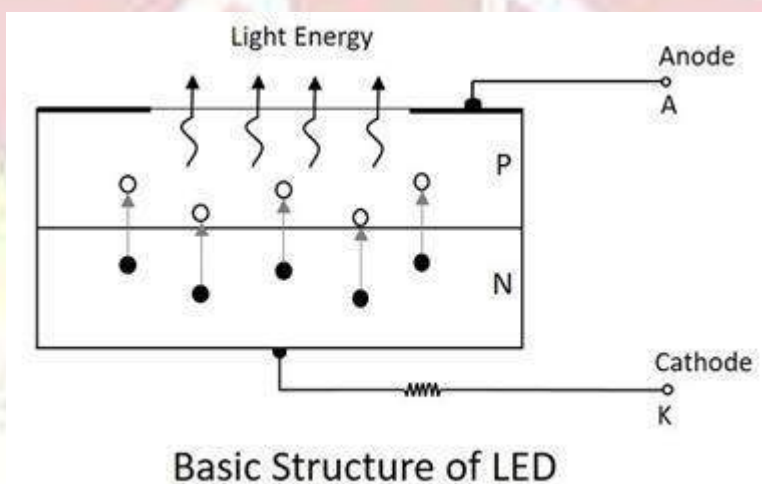
Symbol of LED

Like a normal PN junction diode, this is connected in forward bias condition so that the diode conducts. The conduction takes place in a LED when the free electrons in the conduction band combine with the holes in the valence band. This process of recombination emits **light**. This process is called as **Electroluminescence**. The color of the light emitted depends upon the gap between the energy bands. The materials used also effect the colors like, gallium arsenide phosphide emits either red or yellow, gallium phosphide emits either red or green and gallium nitrate emits blue light, whereas gallium arsenide emits infrared light. The LEDs for non-visible Infrared light are used mostly in remote controls.



LED in the above figure has a flat side and curved side, the lead at the flat side is made shorter than the other one, so as to indicate that the shorter one is **Cathode** or negative terminal and the other one is **Anode** or the Positive terminal.

The basic structure of LED is as shown in the figure below.



As shown in the above figure, as the electrons jump into the holes, the energy is dissipated spontaneously in the form of light. LED is a current dependent device. The output light intensity depends upon the current through the diode.

Advantages of LED

There are many advantages of LED such as –

- High efficiency
- High speed

- High reliability
- Low heat dissipation
- Larger life span
- Low cost
- Easily controlled and programmable
- High levels of brightness and intensity
- Low voltage and current requirements
- Less wiring required
- Low maintenance cost
- No UV radiation
- Instant Lighting effect

Applications of LED

There are many applications for LED such as –

In Displays

- Especially used for seven segment display
- Digital clocks
- Microwave ovens
- Traffic signaling
- Display boards in railways and public places
- Toys

In Electronic Appliances

- Stereo tuners
- Calculators
- DC power supplies
- On/Off indicators in amplifiers
- Power indicators

Commercial Use

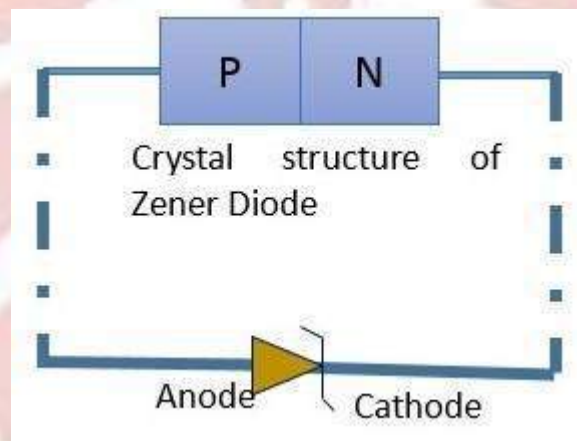
- Infrared readable machines
- Barcode readers
- Solid state video displays

Optical Communications

- In Optical switching applications
- For Optical coupling where manual help is unavailable
- Information transfer through FOC
- Image sensing circuits
- Burglar alarms
- In Railway signaling techniques
- Door and other security control systems

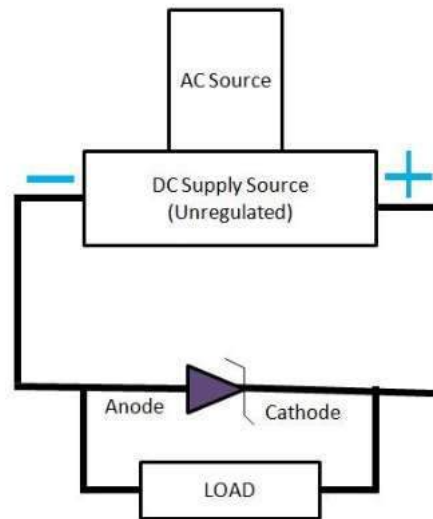
Zener Diode

It is a specific type of semiconductor diode, which is made to operate in the reverse breakdown region. The following figure depicts the crystal structure and the symbol of a Zener diode. It is mostly similar to that of a conventional diode. However, small modification is done to distinguish it from a symbol of a regular diode. The bent line indicates letter 'Z' of the Zener.



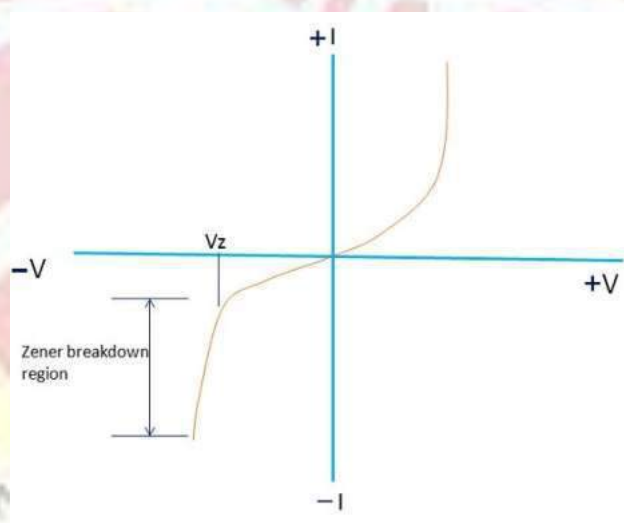
The most significant difference in Zener diodes and regular PN junction diodes is in the mode which they are used in circuits. These diodes are normally operated only in the reverse bias direction, which implies that the anode must be connected to the negative side of the voltage source and the cathode to the positive. If a regular diode is used in the same way as Zener diode, it will be destroyed due to excessive current. This property makes the Zener diode less significant.

The following illustration shows a regulator with a Zener diode.



The Zener diode is connected in reverse bias direction across unregulated DC supply source. It is heavily doped so that the reverse breakdown voltage is reduced. This results in a very thin depletion layer. Due to this, the Zener diode has sharp reverse breakdown voltage V_z .

As per the circuit action, breakdown occurs sharply with a sudden increase in current as shown in the following figure.



Voltage V_z remains constant with an increase in current. Due to this property, Zener diode is widely used in voltage regulation. It provides almost constant output voltage irrespective of the change in current through the Zener. Thus, the load voltage remains at a constant value. At a particular reverse voltage known as knee voltage, current increases sharply with constant voltage. Due to this property, Zener diodes are widely used in voltage stabilization.

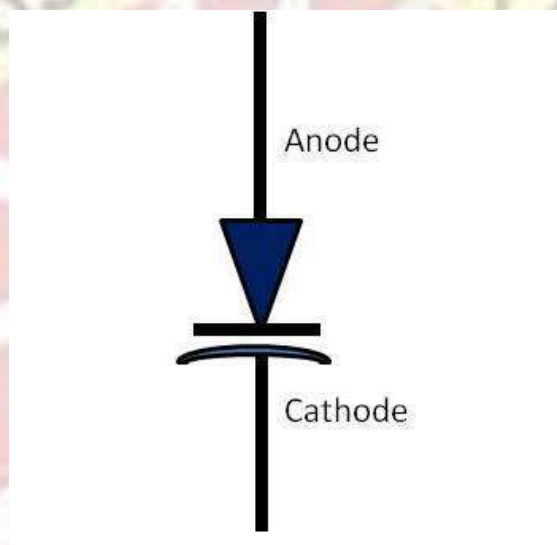
Varactor Diode

This is a special P-N junction diode with an inconsistent concentration of impurities in its P-N materials. In a normal PN junction diode, doping impurities are usually dispersed equally throughout the material. Varactor diode doped with a very small quantity of impurities near the junction and impurity concentration

increases moving away from the junction. In conventional junction diode, the depletion region is an area which separates the P and N material. The depletion region is developed in the beginning when the junction is initially formed. There are no current carriers in this region and thus the depletion region acts as a dielectric medium or insulator. The P-type material with holes as majority carriers and N type material with electrons as majority carriers now act as charged plates. Thus the diode can be considered as a capacitor with N- and P-type opposite charged plates and the depletion region acts as dielectric. As we know, P and N materials, being semiconductors, are separated by a depletion region insulator.

Diodes which are designed to respond to the capacitance effect under reverse bias are called **varactors**, **varicap diodes**, or **voltage-variable capacitors**.

The following figure shows the symbol of Varactor diode.



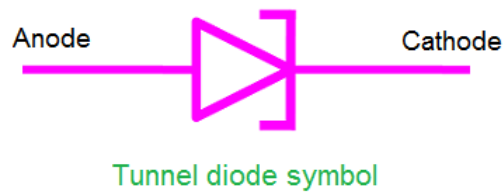
Varactor diodes are normally operated in the reverse bias condition. When the reverse bias increases, the width of the depletion region also increases resulting in less capacitance. This means when reverse bias decreases, a corresponding increase in capacitance can be seen. Thus, diode capacitance varies inversely proportional to the bias voltage. Usually this is not linear. It is operated between zero and the reverse breakdown voltage. These diodes are variable used in microwave applications. Varactor diodes are also used in resonant circuits where some level of voltage tuning or frequency control is required. This diode is also employed in Automatic Frequency Control (AFC) in FM radio and television receivers.

Tunnel diode

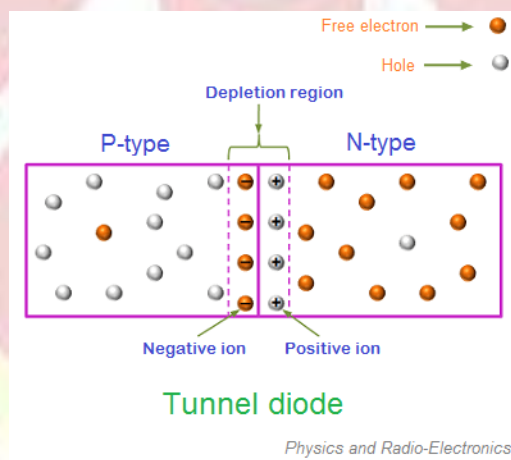
A Tunnel diode is a heavily doped p-n junction diode in which the electric current decreases as the voltage increases. In tunnel diode, electric current is caused by “Tunneling”. The tunnel diode is used as a very fast switching device in computers. It is also used in high-frequency oscillators and amplifiers.

Symbol of tunnel diode

The circuit symbol of tunnel diode is shown in the below figure. In tunnel diode, the p-type semiconductor act as an anode and the n-type semiconductor act as a cathode.



We know that a anode is a positively charged electrode which attracts electrons whereas cathode is a negatively charged electrode which emits electrons. In tunnel diode, n-type semiconductor emits or produces electrons so it is referred to as the cathode. On the other hand, p-type semiconductor attracts electrons emitted from the n-type semiconductor so p-type semiconductor is referred to as the anode. A tunnel diode is also known as Esaki diode which is named after Leo Esaki for his work on the tunneling effect. The operation of tunnel diode depends on the quantum mechanics principle known as “Tunneling”. In electronics, tunneling means a direct flow of electrons across the small depletion region from n-side conduction band into the p-side valence band.

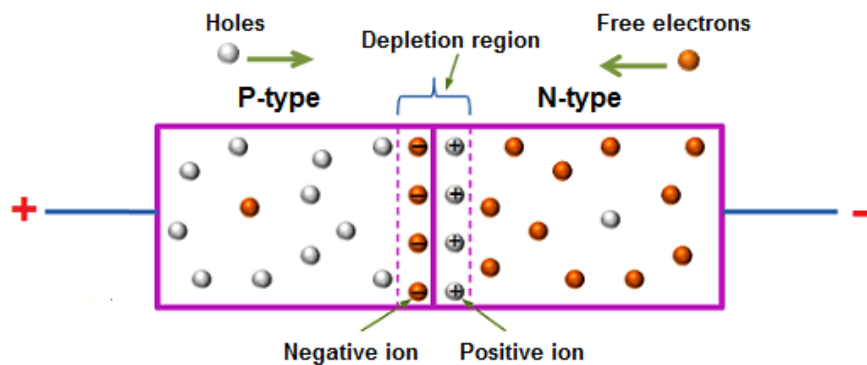


The germanium material is commonly used to make the tunnel diodes. They are also made from other types of materials such as gallium arsenide, gallium antimonide, and silicon.

Width of the depletion region in tunnel diode

The depletion region is a region in a p-n junction diode where mobile charge carriers (free electrons and holes) are absent. Depletion region acts like a barrier that opposes the flow of electrons from the n-type semiconductor and holes from the p-type semiconductor. The width of a depletion region depends on the number of impurities added. Impurities are the atoms introduced into the p-type and n-type semiconductor to increase electrical conductivity. If a small number of impurities are added to the p-n

junction diode (p-type and n-type semiconductor), a wide depletion region is formed. On the other hand, if large number of impurities are added to the p-n junction diode, a narrow depletion region is formed.



Forward bias tunnel diode

Physics and Radio-Electronics

In tunnel diode, the p-type and n-type semiconductor is heavily doped which means a large number of impurities are introduced into the p-type and n-type semiconductor. This heavy doping process produces an extremely narrow depletion region. The concentration of impurities in tunnel diode is 1000 times greater than the normal p-n junction diode. In normal p-n junction diode, the depletion width is large as compared to the tunnel diode. This wide depletion layer or depletion region in normal diode opposes the flow of current. Hence, depletion layer acts as a barrier. To overcome this barrier, we need to apply sufficient voltage. When sufficient voltage is applied, electric current starts flowing through the normal p-n junction diode.

Unlike the normal p-n junction diode, the width of a depletion layer in tunnel diode is extremely narrow. So applying a small voltage is enough to produce electric current in tunnel diode. Tunnel diodes are capable of remaining stable for a long duration of time than the ordinary p-n junction diodes. They are also capable of high-speed operations.

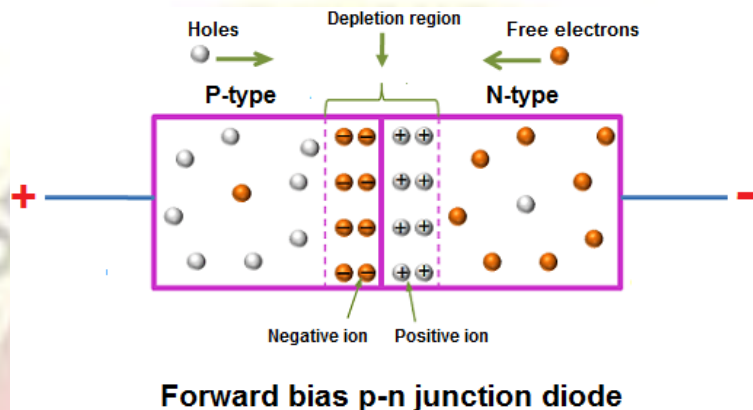
Concept of tunneling

The depletion region or depletion layer in a p-n junction diode is made up of positive ions and negative ions. Because of these positive and negative ions, there exists a built-in-potential or electric field in the depletion region. This electric field in the depletion region exerts electric force in a direction opposite to that of the external electric field (voltage).

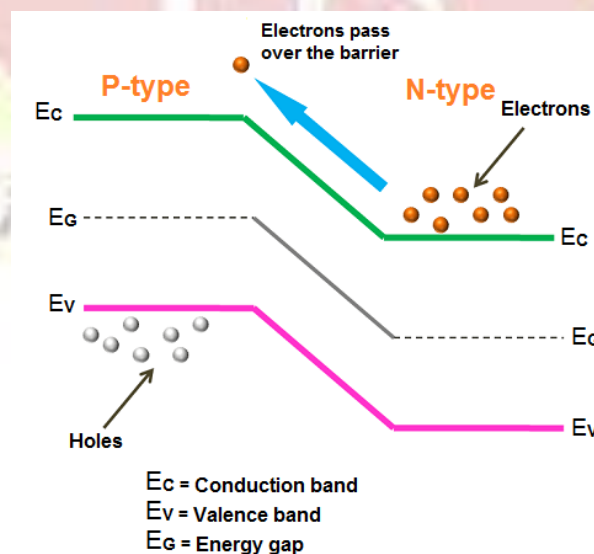
The valence band and conduction band energy levels in the n-type semiconductor are slightly lower than the valence band and conduction band energy levels in the p-type semiconductor. This difference in energy levels is due to the differences in the energy levels of the dopant atoms (donor or acceptor atoms) used to form the n-type and p-type semiconductor.

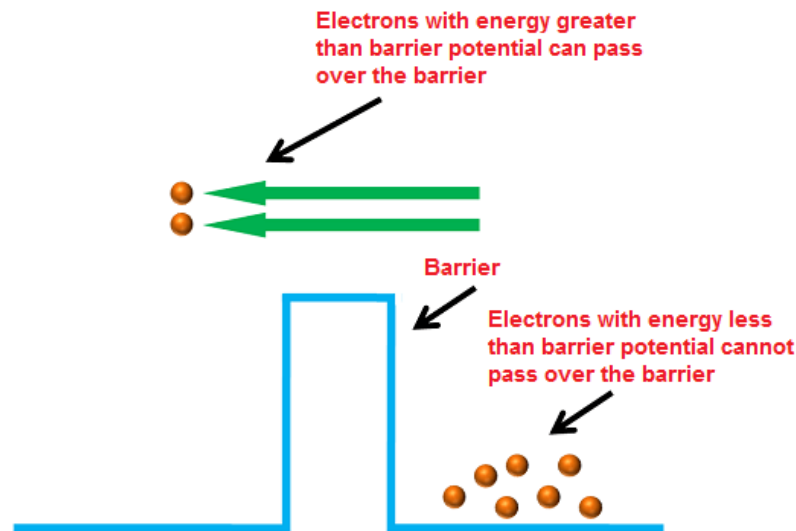
Electric current in ordinary p-n junction diode

When a forward bias voltage is applied to the ordinary p-n junction diode, the width of depletion region decreases and at the same time the barrier height also decreases. However, the electrons in the n-type semiconductor cannot penetrate through the depletion layer because the built-in voltage of depletion layer opposes the flow of electrons.



If the applied voltage is greater than the built-in voltage of depletion layer, the electrons from n-side overcome the opposing force from depletion layer and then enter into p-side. In simple words, the electrons can pass over the barrier (depletion layer) if the energy of the electrons is greater than the barrier height or barrier potential.



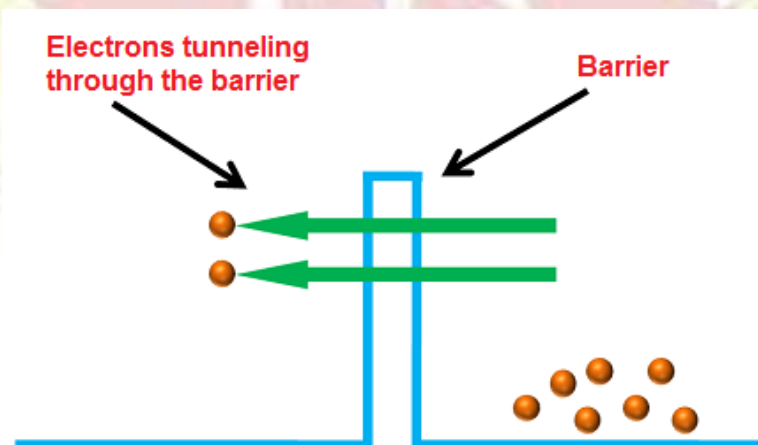


Ordinary p-n junction diode

Therefore, an ordinary p-n junction diode produces electric current only if the applied voltage is greater than the built-in voltage of the depletion region.

Electric current in tunnel diode

In tunnel diode, the valence band and conduction band energy levels in the n-type semiconductor are lower than the valence band and conduction band energy levels in the p-type semiconductor. Unlike the ordinary p-n junction diode, the difference in energy levels is very high in tunnel diode. Because of this high difference in energy levels, the conduction band of the n-type material overlaps with the valence band of the p-type material.



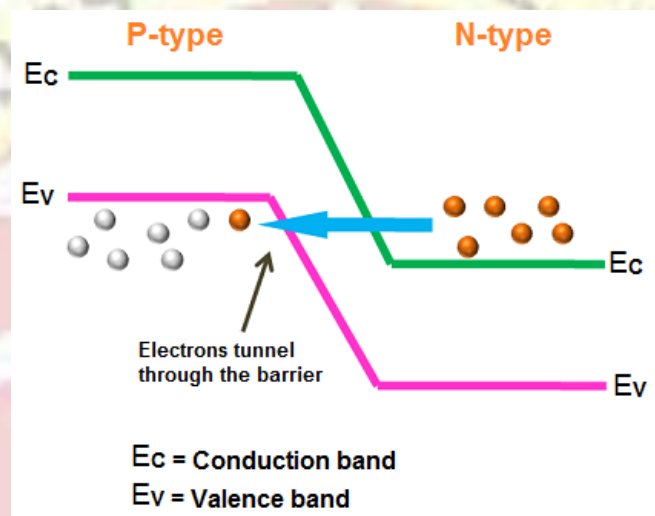
Tunnel diode

Quantum mechanics says that the electrons will directly penetrate through the depletion layer or barrier if the depletion width is very small. The depletion layer of tunnel diode is very small. It is in nanometers. So

the electrons can directly tunnel across the small depletion region from n-side conduction band into the p-side valence band.

In ordinary diodes, current is produced when the applied voltage is greater than the built-in voltage of the depletion region. But in tunnel diodes, a small voltage which is less than the built-in voltage of depletion region is enough to produce electric current.

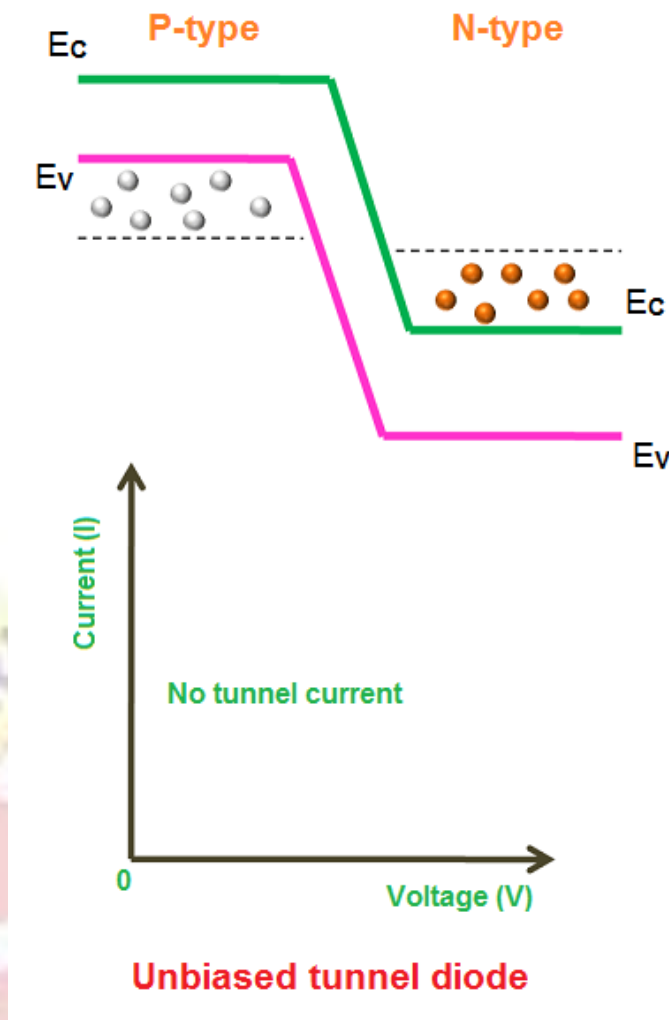
In tunnel diodes, the electrons need not overcome the opposing force from the depletion layer to produce electric current. The electrons can directly tunnel from the conduction band of n-region into the valence band of p-region. Thus, electric current is produced in tunnel diode.



Working :

Step 1: Unbiased tunnel diode

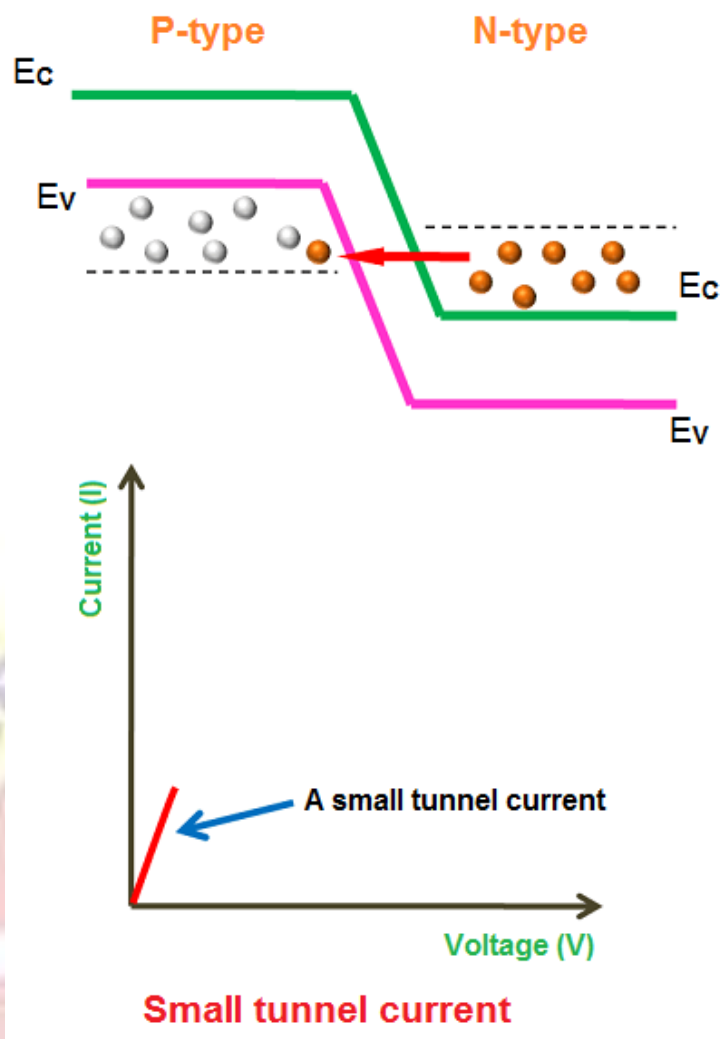
When no voltage is applied to the tunnel diode, it is said to be an unbiased tunnel diode. In tunnel diode, the conduction band of the n-type material overlaps with the valence band of the p-type material because of the heavy doping.



Because of this overlapping, the conduction band electrons at n-side and valence band holes at p-side are nearly at the same energy level. So when the temperature increases, some electrons tunnel from the conduction band of n-region to the valence band of p-region. In a similar way, holes tunnel from the valence band of p-region to the conduction band of n-region. However, the net current flow will be zero because an equal number of charge carriers (free electrons and holes) flow in opposite directions.

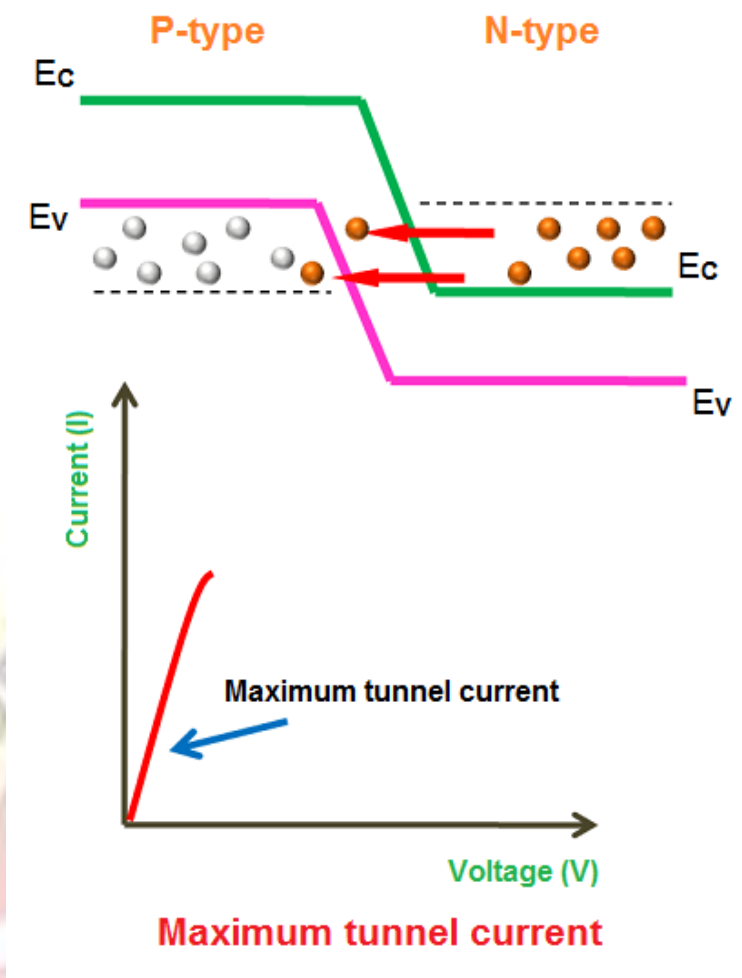
Step 2: Small voltage applied to the tunnel diode

When a small voltage is applied to the tunnel diode which is less than the built-in voltage of the depletion layer, no forward current flows through the junction. However, a small number of electrons in the conduction band of the n-region will tunnel to the empty states of the valence band in p-region. This will create a small forward bias tunnel current. Thus, tunnel current starts flowing with a small application of voltage.



Step 3: Applied voltage is slightly increased

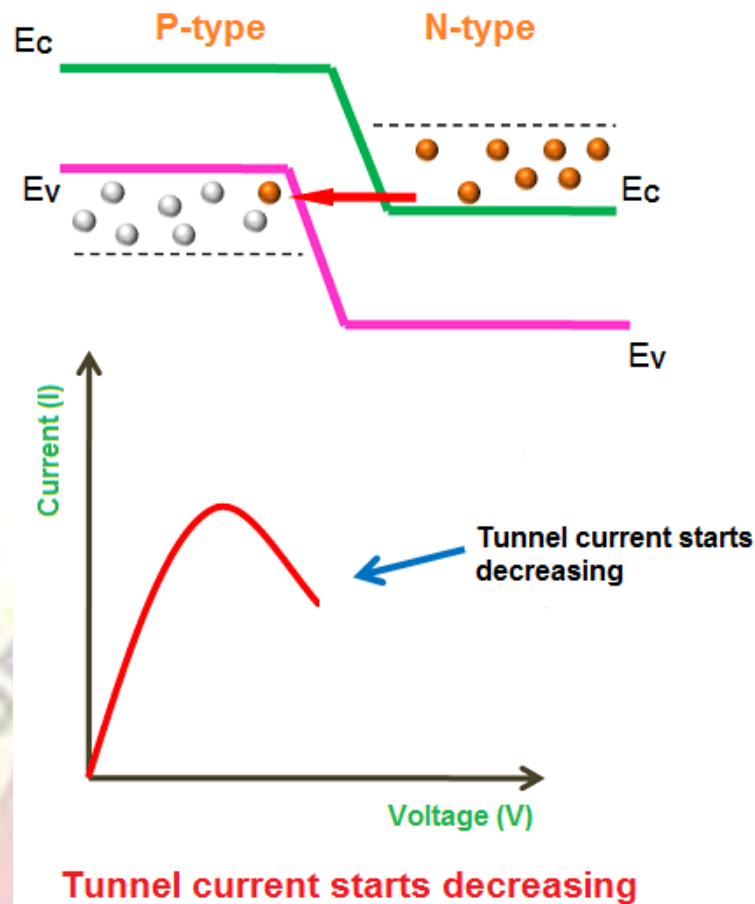
When the voltage applied to the tunnel diode is slightly increased, a large number of free electrons at n-side and holes at p-side are generated. Because of the increase in voltage, the overlapping of the conduction band and valence band is increased.



In simple words, the energy level of an n-side conduction band becomes exactly equal to the energy level of a p-side valence band. As a result, maximum tunnel current flows.

Step 4: Applied voltage is further increased

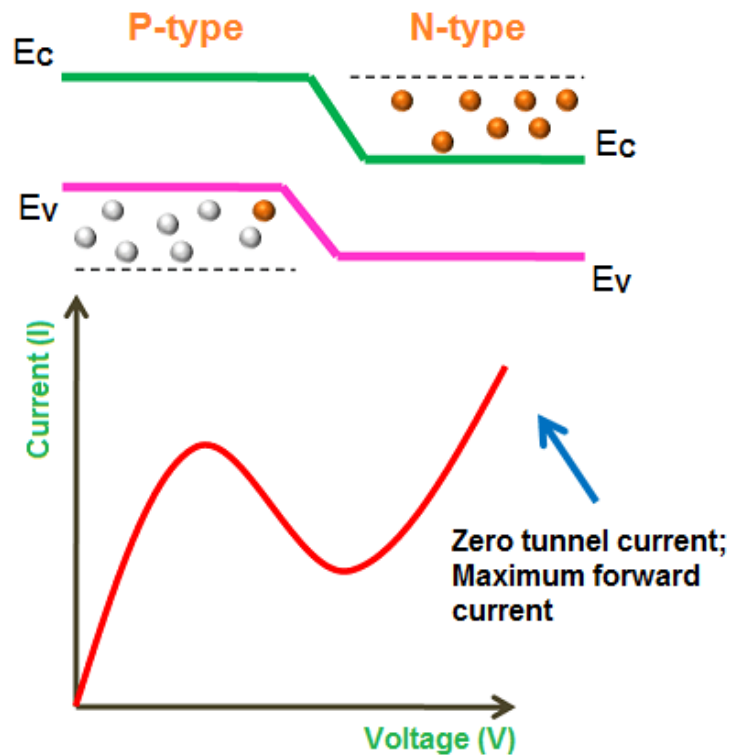
If the applied voltage is further increased, a slight misalign of the conduction band and valence band takes place.



Since the conduction band of the n-type material and the valence band of the p-type material still overlap, the electrons tunnel from the conduction band of the n-region to the valence band of the p-region and cause a small current flow. Thus, the tunneling current starts decreasing.

Step 5: Applied voltage is largely increased

If the applied voltage is largely increased, the tunneling current drops to zero. At this point, the conduction band and valence band no longer overlap and the tunnel diode operates in the same manner as a normal p-n junction diode.



Zero tunnel current; maximum forward current

If this applied voltage is greater than the built-in potential of the depletion layer, the regular forward current starts flowing through the tunnel diode. The portion of the curve in which current decreases as the voltage increases is the negative resistance region of the tunnel diode. The negative resistance region is the most important and most widely used characteristic of the tunnel diode. A tunnel diode operating in the negative resistance region can be used as an amplifier or an oscillator.

Advantages of tunnel diodes

- Long life
- High-speed operation
- Low noise
- Low power consumption

Disadvantages of tunnel diodes

- Tunnel diodes cannot be fabricated in large numbers
- Being a two terminal device, the input and output are not isolated from one another.

Applications of tunnel diodes

- Tunnel diodes are used as logic memory storage devices.

- Tunnel diodes are used in relaxation oscillator circuits.
- Tunnel diode is used as an ultra high-speed switch.
- Tunnel diodes are used in FM receivers.

UNIT – III

3.0 INTRODUCTION

W

hen a third doped element is added to a crystal diode in such a way that two $p-n$ junctions are formed, the result is a device known as a *transistor*. The transistor— an entirely new type of electronic device— is capable of achieving amplification of weak signals in a

fashion comparable and often superior to that realised by vacuum tubes. Transistors are far smaller than vacuum tubes, have no filament and hence need no heating power and may be operated in any position. They are mechanically strong, have practically unlimited life and can do some jobs better than vacuum tubes.



Invented in 1943 by J. Bardeen and W. H. Brattain of Bell Telephone Laboratories, U.S.A., transistor has now become the heart of most electronic applications. Though transistor is only slightly more than 53 years old, yet it is fast replacing vacuum tubes in almost all applications. In this chapter, we shall focus our attention on the various aspects of transistors and their increasing applications in the fast developing electronics industry.

3.1 Transistor

A **transistor** consists of two pn junctions formed by sandwiching either p -type or n -type semiconductor between a pair of opposite types. Accordingly, there are two types of transistors, namely;

(i) n - p - n transistor

(ii) p - n - p transistor

An n - p - n transistor is composed of two n -type semiconductors separated by a thin section of p -type as shown in Fig. 3.1(i). However, a p - n - p transistor is formed by two p -sections separated by a thin section of n -type as shown in Fig. 3.1(ii).



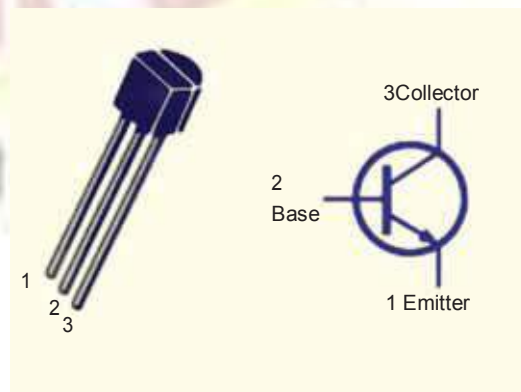
Fig.3.1

In each type of transistor, the following points may be noted:

- (i) These are two pn junctions. Therefore, a transistor may be regarded as a combination of two diodes connected back to back.
- (ii) There are three terminals, one taken from each type of semiconductor.
- (iii) The middle section is a very thin layer. This is the most important factor in the function of a transistor.

Origin of the name "Transistor". When new devices are invented, scientists often try to devise a name that will appropriately describe the device. A transistor has two pn junctions. As discussed later, one junction is forward biased and the other is reverse biased. The forward biased junction has a low resistance path whereas a reverse biased junction has a high resistance path. The weak signal is introduced in the low resistance circuit and output is taken from the high resistance circuit. Therefore, a transistor *transfers* a signal from a low resistance to a high resistance. The prefix 'trans' means the signal

transfer property of the device while 'istor' classifies it as a solid element in the same general family with resistors.



Naming the Transistor Terminals

A transistor (pnp or npn) has three sections of doped semiconductors. The section on one side is the **emitter** and the section on the opposite side is the **collector**. The middle section is called the **base** and forms two junctions between the emitter and collector.

(ii) Emitter. The section on one side that supplies charge carriers (electrons or holes) is called the **emitter**. *The emitter is always forward biased w.r.t. base* so that it can supply a large number of ***majority carriers**. In Fig. 3.2 (i), the emitter (p -type) of pnp transistor is forward biased and supplies hole charges to its junction with the base. Similarly, in Fig. 3.2 (ii), the emitter (n -type) of npn transistor has a forward bias and supplies free electrons to its junction with the base.

(iii) Collector. The section on the other side that collects the charges is called the **collector**. *The collector is always reverse biased*. Its function is to remove charges from its junction with the base. In Fig. 3.2 (i), the collector (p -type) of pnp transistor has a reverse bias and receives hole charges that flow in the output circuit. Similarly, in Fig. 3.2 (ii), the collector (n -type) of npn transistor has reverse bias and receives electrons.

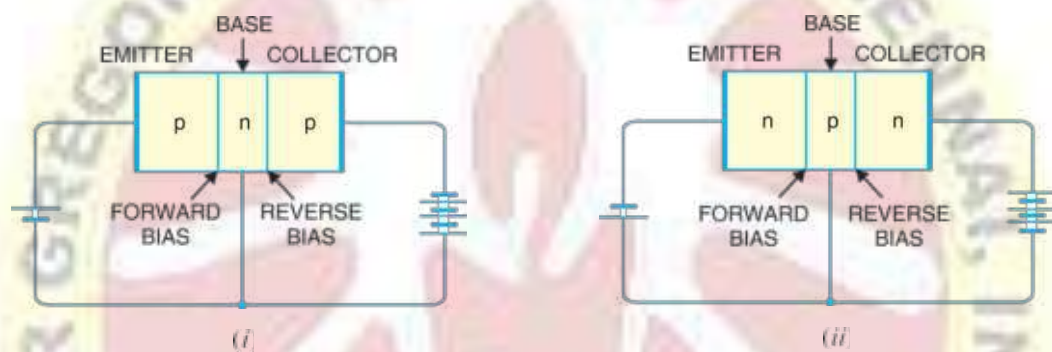


Fig.3.2

(iv) Base. The middle section which forms two p - n junctions between the emitter and collector is called the **base**. The base-emitter junction is forward biased, allowing low resistance for the emitter circuit. The base-collector junction is reverse biased and provides high resistance in the collector circuit.

Some Facts about the Transistor

Before discussing transistor action, it is important that the reader may keep in mind the following facts about the transistor:

(v) The transistor has three regions, namely ; **emitter**, **base** and **collector**. The base is much thinner than the emitter while ****collector** is wider than both as shown in Fig.3.3. However, for the sake of convenience, it is customary to show emitter and collector to be of equal size.

(vi) The emitter is heavily doped so that it can inject a large number of charge carriers (electrons or holes) into the base. The base is lightly doped and very thin; it passes most of the emitter injected charge carriers to the collector. The collector is moderately doped.



Fig.3.3

(vii) The transistor has two *pn* junctions *i.e.* it is like two diodes. The junction between emitter and base may be called *emitter-base diode* or simply the *emitter diode*. The junction between the base and collector may be called *collector-base diode* or simply *collector diode*.

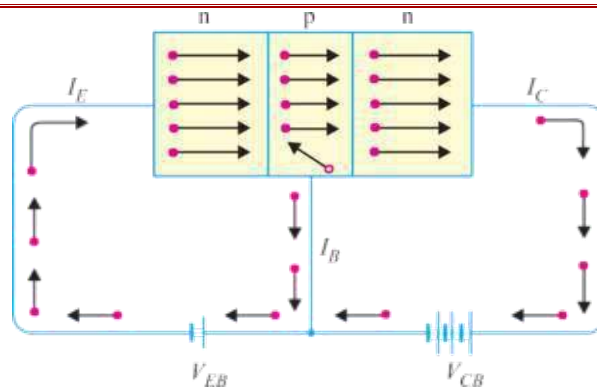
(viii) The emitter diode is always forward biased whereas collector diode is always reverse biased.

(ix) The resistance of emitter diode (forward biased) is very small as compared to collector diode (reverse biased). Therefore, forward bias applied to the emitter diode is generally very small whereas reverse bias on the collector diode is much higher.

3.2 Transistor Action

The emitter-base junction of a transistor is forward biased whereas collector-base junction is reverse biased. If for a moment, we ignore the presence of emitter-base junction, then *practically* no current would flow in the collector circuit because of the reverse bias. However, if the emitter-base junction is also present, then forward bias on it causes the emitter current to flow. It is seen that this emitter current almost entirely flows in the collector circuit. Therefore, the current in the collector circuit depends upon the emitter current. If the emitter current is zero, then collector current is nearly zero. However, if the emitter current is 1 mA, then collector current is also about 1 mA. This is precisely what happens in a transistor. We shall now discuss this transistor action for *npn* and *pnp* transistors.

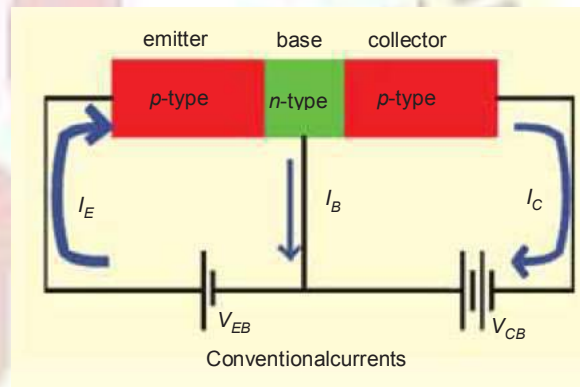
(x) **Working of npn transistor.** Fig. 3.4 shows the *npn* transistor with forward bias to emitter-base junction and reverse bias to collector-base junction. The forward bias causes the electrons in the *n*-type emitter to flow towards the base. This constitutes the emitter current I_E . As these electrons flow through the *p*-type base, they tend to combine with holes. As the base is lightly doped and very thin, therefore, only a few electrons (less than 5%) combine with holes to constitute base current I_B . The remainder (** more than 95%) cross over into the collector region to constitute collector current I_C . In this way, almost the entire emitter current flows in the collector circuit. It is clear that emitter current is the sum of collector and base currents
i.e. $I_E = I_B + I_C$



Basic connection of *nnp* transisto

Fig.3.4

(ii) Working of pnp transistor. Fig. 3.5 shows the basic connection of a *pnp* transistor. The forward bias causes the holes in the *p*-type emitter to flow towards the base. This constitutes the emitter current I_E . As these holes cross into *n*-type base, they tend to combine with the electrons. As the base is lightly doped and very thin, therefore, only a few holes (less than 5%) combine with the electrons. The remainder (more than 95%) cross into the collector region to constitute collector current I_C . In this way, almost the entire emitter current flows in the collector circuit. It may be noted that current conduction within *pnp* transistor is by holes. However, in the external connecting wires, the current is still by electrons.



Importance of transistor action. The input circuit (*i.e.* emitter-base junction) has low resistance because of forward bias whereas output circuit (*i.e.* collector-base junction) has high resistance due to reverse bias. As we have seen, the input emitter current almost entirely flows in the collector circuit. Therefore, a transistor transfers the input signal current from a low-resistance circuit to a high-resistance circuit. This is the key factor responsible for the amplifying capability of the transistor. We shall discuss the amplifying property of transistor later in this chapter.

note. There are two basic transistor types : the **bipolar junction transistor (BJT)** and **field-effect transistor (FET)**. As we shall see, these two transistor types differ in both their operating characteristics and their internal construction. **Note that when we use the term transistor, it means bipolar junction transistor (BJT).** The term comes from the fact that in a bipolar transistor, there are *two* types of charge carriers (*viz.* electrons and holes) that play part in conduction. Note that **bipolar** refers to polarities. The field-effect transistor is simply referred to as **FET**.

3.3 Transistor Symbols

In the earlier diagrams, the transistors have been shown in diagrammatic form. However, for the sake of convenience, the transistors are represented by schematic diagrams. The symbols used for $n p n$ and $p n p$ transistors are shown in Fig. 3.6.

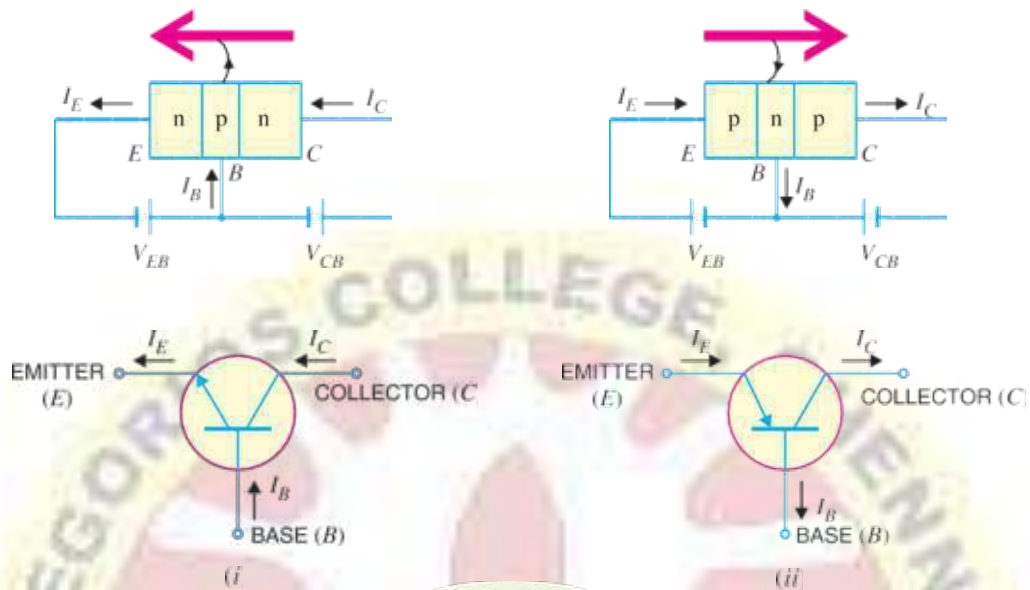


Fig.3.6

Note that emitter is shown by an arrow which indicates the direction of conventional current flow with forward bias. For $n p n$ connection, it is clear that conventional current flows out of the emitter as indicated by the outgoing arrow in Fig. 3.6(i). Similarly, for $p n p$ connection, the conventional current flows into the emitter as indicated by the inward arrow in Fig. 3.6(ii).

3.4 Transistor Circuit as an Amplifier

A transistor raises the strength of a weak signal and thus acts as an amplifier. Fig. 3.7 shows the basic circuit of a transistor amplifier. The weak signal is applied between emitter-base junction and output is taken across the load R_C connected in the collector circuit. In order to achieve faithful amplification, the input circuit should always remain forward biased. To do so, a d.c. voltage V_{EE} is applied in the input circuit in addition to the signals.

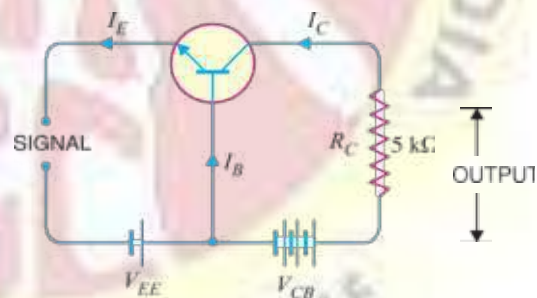


Fig.3.7

shown. This d.c. voltage is known as bias voltage and its magnitude is such that it always keeps the input circuit forward biased regardless of the polarity of the signal.

As the input circuit has low resistance, therefore, a small change in signal voltage causes an appreciable change in emitter current. This causes almost the same change in collector current due to transistor action. The collector current flowing through a high load resistance R_C produces a large voltage across it. Thus, a weak signal applied in the input circuit appears in the amplified form in the collector circuit. It is in this way that a transistor acts as an amplifier.

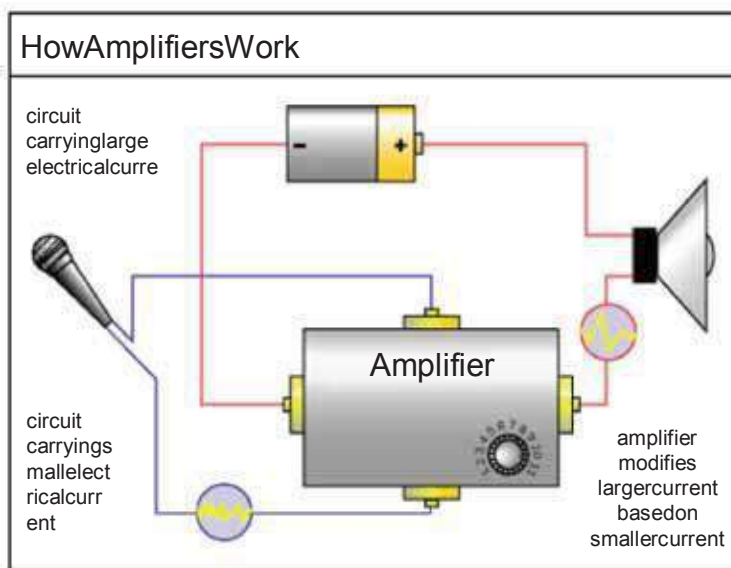


Illustration. The action of a transistor as an amplifier can be made more illustrative if we consider typical circuit values. Suppose collector load resistance $R_C = 5 \text{ k}\Omega$. Let us further assume that a change of 0.1 V in signal voltage produces a change of 1 mA in emitter current. Obviously, the change in collector current would also be approximately 1 mA . This collector current flowing through collector load R_C would produce a voltage $= 5 \text{ k}\Omega \times 1 \text{ mA} = 5 \text{ V}$. Thus, a change of 0.1 V in the signal has caused a change of 5 V

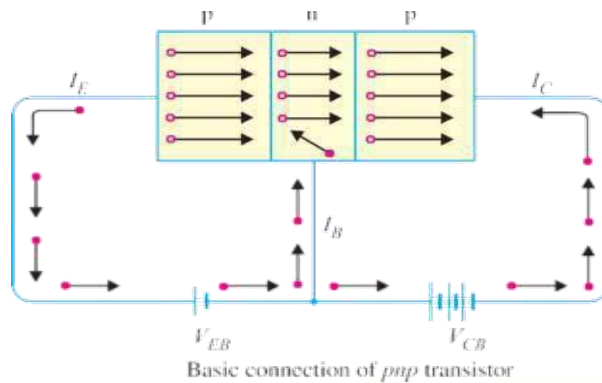
in the output circuit. In other words, the transistor has been able to raise the voltage level of the signal from 0.1 V to 5 V , i.e. voltage amplification is 50.

3.5 Transistor Connections

There are three leads in a transistor viz., emitter, base and collector terminals. However, when a transistor is to be connected in a circuit, we require four terminals; two for the input and two for the output. This difficulty is overcome by making one terminal of the transistor common to both input and output terminals. The input is fed between this common terminal and one of the other two terminals. The output is obtained between the common terminal and the remaining terminal. Accordingly, a transistor can be connected in a circuit in the following three ways:

- (xi) common base connection
- (ii) common emitter connection
- (iii) common collector connection

Each circuit connection has specific advantages and disadvantages. It may be noted here that regardless of circuit connection, the emitter is always biased in the forward direction, while the collector is always biased in the reverse direction.



3.6 commonBaseConnection

In this circuit arrangement, input is applied between emitter and base and output is taken from collector and base. Here, base of the transistor is common to both input and output circuits and hence the name common base connection. In Fig. 3.9(i), a common base *npn* transistor circuit is shown whereas Fig. 3.9(ii) shows the common base *pnp* transistor circuit.

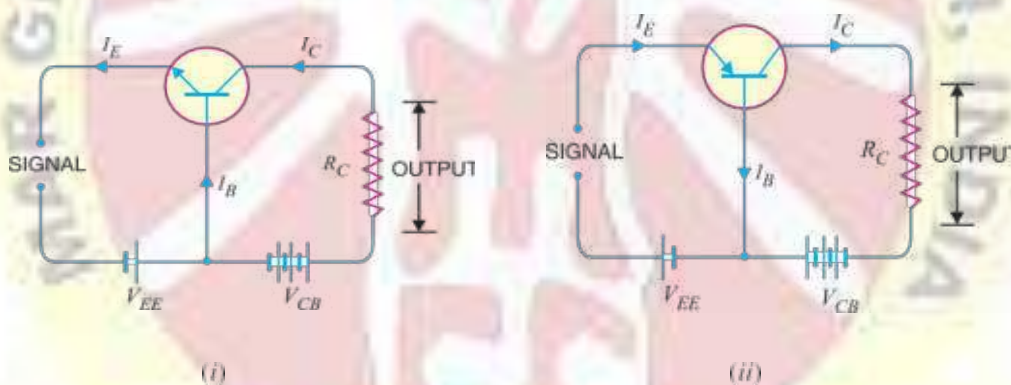


Fig.3.9

1. Current amplification factor (α). It is the ratio of output current to input current. In a common base connection, the input current is the emitter current I_E and output current is the collector current I_C .

The ratio of change in collector current to the change in emitter current at constant collector-base voltage V_{CB} is known as **current amplification factor** i.e.

$$\alpha = \frac{\Delta I_C}{\Delta I_E} \text{ at constant } V_{CB}$$

It is clear that current amplification factor is less than unity. This value can be increased (but not more than unity) by decreasing the base current. This is achieved by making the base thin and doping it lightly. Practical values of α in commercial transistors range from 0.9 to 0.99.

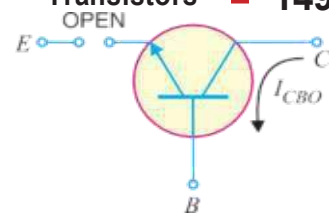


Fig. 3.10

2. Expression for collector current. The whole of emitter current does not reach the collector. It is because a small percentage of it, as a result of electron-hole combinations occurring in base area, gives rise to base current. Moreover, as the collector-base junction is reverse biased, therefore, some leakage current flows due to minority carriers. It follows, therefore, that total collector current consists of:

- (i) That part of emitter current which reaches the collector terminal i.e. αI_E .
- (ii) The leakage current $I_{leakage}$. This current is due to the movement of minority carriers across base-collector junction on account of it being reverse biased. This is generally much smaller than αI_E .

$$\therefore \text{Total collector current, } I_C = \alpha I_E + I_{leakage}$$

It is clear that if $I_E = 0$ (i.e., emitter circuit is open), a small leakage current still flows in the collector circuit. This $I_{leakage}$ is abbreviated as I_{CBO} , meaning collector-base current with emitter open. The I_{CBO} is indicated in Fig. 3.10.

$$\therefore I_C = \alpha I_E + I_{CBO} \quad \dots(i)$$

$$\text{Now } I_E = I_C + I_B$$

$$\therefore I_C = \alpha(I_C + I_B) + I_{CBO}$$

$$\text{or } I_C(1 - \alpha) = \alpha I_B + I_{CBO}$$

$$\text{or } I_C = \frac{\alpha I_B}{1 - \alpha} + \frac{I_{CBO}}{1 - \alpha} \quad \dots(ii)$$

Relation (i) or (ii) can be used to find I_C . It is further clear from these relations that the collector current of a transistor can be controlled by either the emitter or base current.

Fig. 3.11 shows the concept of I_{CBO} . In CB configuration, a small collector current flows even when the emitter current is zero. This is the leakage collector current (i.e. the collector current when emitter is open) and is denoted by I_{CBO} . When the emitter voltage V_{EE} is also applied, the various currents are as shown in Fig. 3.11(ii).

3.7 Characteristics of Common Base Connection

The complete electrical behavior of a transistor can be described by stating the interrelation of the various currents and voltages. These relationships can be conveniently displayed graphically and the curves thus obtained are known as the characteristics of transistor. The most important characteristics of common base connection are *input characteristics* and *output characteristics*.

1. Input characteristic. It is the curve between emitter current I_E and emitter-base voltage V_{EB} at constant collector-base voltage V_{CB} . The emitter current is generally taken along y-axis and emitter-base voltage along x-axis. Fig. 3.11

shows the input characteristics of a typical transistor in CB arrangement. The following points may be noted from these characteristics:

(i) The emitter current I_E increases rapidly with small increase in emitter-base voltage V_{EB} . It means that input resistance is very small.

(ii) The emitter current is almost independent of collector-base voltage V_{CB} . This leads to the conclusion that emitter current (and hence collector current) is almost independent of collector voltage.

Input resistance. It is the ratio of change

in emitter-base voltage (ΔV_{EB}) to the resulting change in emitter current (ΔI_E) at constant collector-base voltage (V_{CB}) i.e.

$$\text{Input resistance, } r_i = \frac{\Delta V_{BE}}{\Delta I_E} \text{ at constant } V_{CB}$$

In fact, input resistance is the opposition offered to the signal current. As a very small V_{EB} is sufficient to produce a large flow of emitter current I_E , therefore, input resistance is quite small, of the order of a few ohms.

2. Output characteristic. It is the curve between collector current I_C and collector-base voltage V_{CB} at constant emitter current I_E . Generally, collector current is taken along y-axis and collector-base voltage along x-axis. Fig. 3.12 shows the output characteristics of a typical transistor in CB arrangement.

The following points may be noted from the characteristics:

(i) The collector current I_C varies with V_{CB} only at very low voltages ($< 1\text{ V}$). The transistor is *never* operated in this region.

(ii) When the value of V_{CB} is raised above 1–2 V, the collector current becomes constant as indicated by straight horizontal curves. It means that now I_C is independent of V_{CB} and depends upon I_E only. This is consistent with the theory that the emitter current flows *almost* entirely to the collector terminal. The transistor is *always* operated in this region.

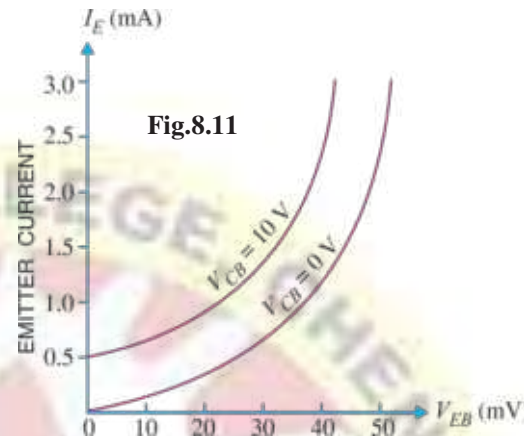


Fig. 3.11

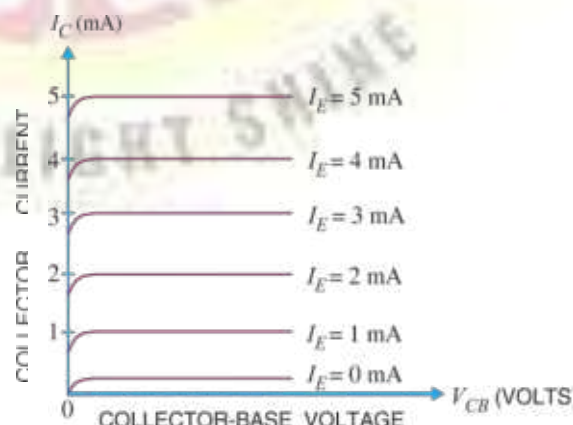


Fig. 3.12

(iii) A very large change in collector-base voltage produces only a tiny change in collector current. This means that output resistance is very high.

Output resistance. It is the ratio of change in collector-base voltage (ΔV_{CB}) to the resulting change in collector current (ΔI_C) at constant emitter current $i.e.$

Transistors ■ 151

$$\text{Output resistance, } r_o = \frac{\Delta V_{CB}}{\Delta I_C} \text{ at constant } I_E$$

The output resistance of CB circuit is very high, of the order of several tens of kilo-ohms. This is not surprising because the collector current changes very slightly with the change in V_{CB} .

3.8 Common Emitter Connection

In this circuit arrangement, input is applied between base and emitter and output is taken from the collector and emitter. Here, emitter of the transistor is common to both input and output circuits and hence the name common emitter connection. Fig. 3.13 (i) shows common emitter npn transistor circuit whereas Fig. 3.13 (ii) shows common emitter pnp transistor circuit.

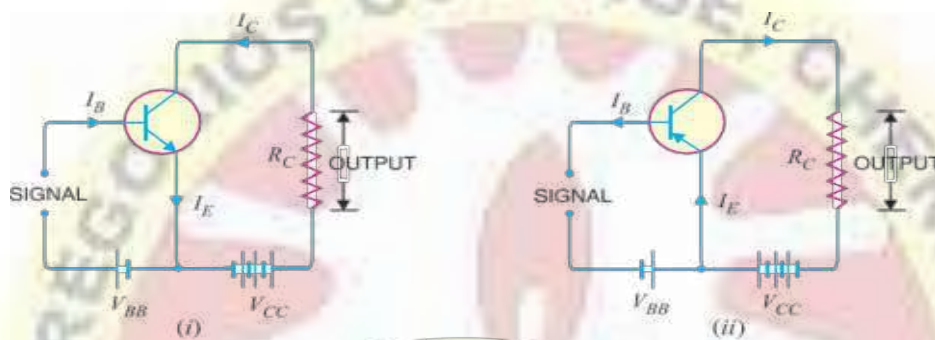


Fig.3.13

Base current amplification factor (β). In common emitter connection, input current is I_B

and output current is I_C .

The ratio of change in collector current (ΔI_C) to the change in base current (ΔI_B) is known as **base current amplification factor** i.e.

$$\beta^* = \frac{\Delta I_C}{\Delta I_B}$$

In almost any transistor, less than 5% of emitter current flows as the base current. Therefore, the value of β is generally greater than 20. Usually, its value ranges from 20 to 500. This type of connection is frequently used as it gives appreciable current gain as well as voltage gain.

Relation between β and α . A simple relation exists between β and α . This can be derived as follows:

$$\beta = \frac{\Delta I_C}{\Delta I_B} \quad \dots(i)$$

$$\alpha = \frac{\Delta I_C}{\Delta I_E} \quad \dots(ii)$$

Now

$$I_E = I_B + I_C$$

or

$$\Delta I_E = \Delta I_B + \Delta I_C$$

or

$$\Delta I_B = \Delta I_E - \Delta I_C$$

Substituting the value of ΔI_B in exp.(i), we get,

$$\beta = \frac{\Delta I_C}{\Delta I_E - \Delta I_C} \quad \dots(iii)$$

Dividing the numerator and denominator of R.H.S. of exp.(iii) by ΔI_E , we get,

$$\beta = \frac{\frac{\Delta I_C}{\Delta I_E}}{\frac{\Delta I_E}{\Delta I_E} - \frac{\Delta I_C}{\Delta I_E}} = \frac{\alpha}{1 - \alpha} \quad \left[\begin{array}{l} \alpha = \frac{\Delta I_C}{\Delta I_E} \\ \alpha = \frac{\Delta I_C}{\Delta I_E} \end{array} \right]$$

$\therefore$

$$\beta = \frac{\alpha}{1 - \alpha}$$

It is clear that as α approaches unity, β approaches infinity. In other words, the current gain in common emitter connection is very high. It is due to this reason that this circuit arrangement is used in about 90 to 95 percent of all transistor applications.

Expression for collector current. In common emitter circuit, I_B is the input current and I_C is the output current.

$$\text{We know } I_E = I_B + I_C \quad \dots(i)$$

$$\text{and } I_C = \alpha I_E + I_{CBO} \quad \dots(ii)$$

$$\text{From exp. (ii), we get, } I_C = \alpha I_E + I_{CBO} = \alpha(I_B + I_C) + I_{CBO}$$

$$\text{or } I_C(1 - \alpha) = \alpha I_B + I_{CBO}$$

$$\text{or } I_C = \frac{\alpha}{1 - \alpha} I_B + \frac{1}{1 - \alpha} I_{CBO} \quad \dots(iii)$$

From exp. (iii), it is apparent that if $I_B = 0$ (i.e. base circuit is open), the collector current will be the current of the emitter. This is abbreviated as I_{CEO} , meaning collector-emitter current with base open.

$$\therefore I_{CEO} = \frac{1}{1 - \alpha} I_{CBO}$$

Substituting the value of $\frac{1}{1 - \alpha} I_{CBO} = I_{CEO}$ in exp. (iii), we get,

$$I_C = \frac{\alpha}{1 - \alpha} I_B + I_{CEO}$$

$$\text{or } I_C = \beta I_B + I_{CEO} \quad \left(\because \beta = \frac{\alpha}{1 - \alpha} \right)$$

Concept of I_{CEO} . In CE configuration, a small collector current flows even when the base current is zero [See Fig. 3.14 (i)]. This is the collector cut off current (i.e. the collector current that flows when base is open) and is denoted by I_{CEO} . The value of I_{CEO} is much larger than I_{CBO} .

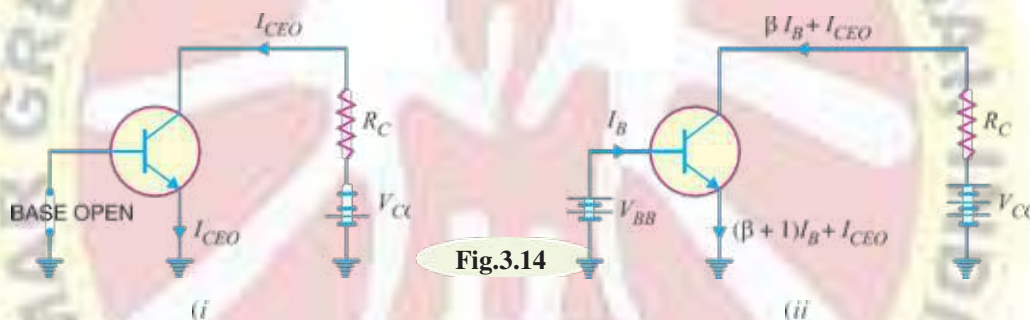


Fig.3.14

When the base voltage is applied as shown in Fig. 3.14 (ii), then the various currents are:

$$\begin{aligned} \text{Base current} &= I_B \\ \text{Collector current} &= \beta I_B + I_{CEO} \\ \text{Emitter current} &= \text{Collector current} + \text{Base current} \\ &= (\beta I_B + I_{CEO}) + I_B = (\beta + 1) I_B + I_{CEO} \end{aligned}$$

It may be noted here that:

$$I_{CEO} = \frac{1}{1 - \alpha} I_{CBO} = (\beta + 1) I_{CBO} \quad \left[\because \frac{1}{1 - \alpha} = \beta + 1 \right]$$

Measurement of Leakage Current

A very small leakage current flows in all transistor circuits. However, in most cases, it is quite small and can be neglected.

(i) **Circuit for I_{CEO} test.** Fig. 3.15 shows the circuit for measuring I_{CEO} . Since base is open

($I_B=0$), the transistor is in cutoff. Ideally, $I_C=0$ but actually there is a small current from collector to emitter due to minority carriers. It is called I_{CEO} (collector-to-emitter current with base open). This current is usually in the nA range for silicon. A faulty transistor will often have excessive leakage current.

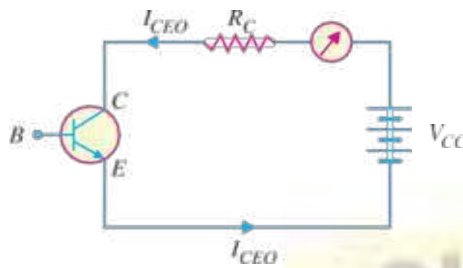


Fig.3.15

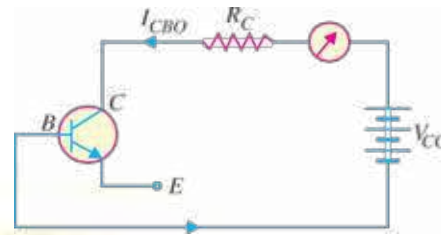


Fig.3.16

(ii) **Circuit for I_{CBO} test.** Fig. 3.16 shows the circuit for measuring I_{CBO} . Since the emitter is open ($I_E = 0$), there is a small current from collector to base. This is called I_{CBO} (collector-to-base current with emitter open). This current is due to the movement of minority carriers across base-collector junction. The value of I_{CBO} is also small. If in measurement, I_{CBO} is excessive, then there is a possibility that collector-base is shorted.

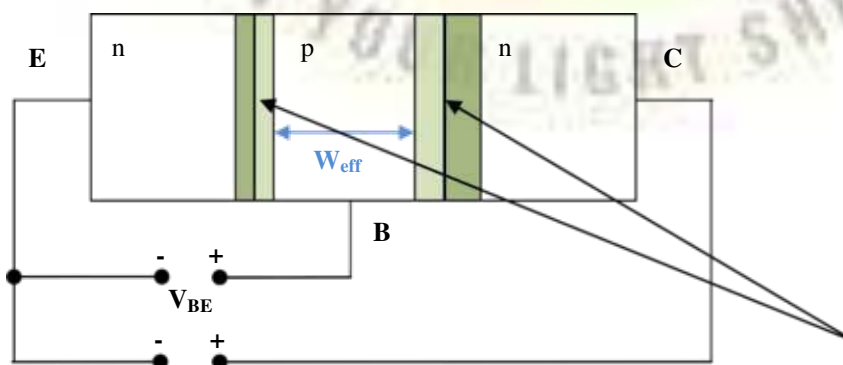
3.9 PUNCH THROUGH

As we increase the reverse bias voltage of base collector junction depletion width increases and effective base width decreases (since base is less doped compared to collector, depletion layer protrudes more into the base the collector hence effective base width gets affected), at some point effective base width approaches zero and transistor will breakdown. This phenomenon called reach through or punch through.

3.10 Early-effect

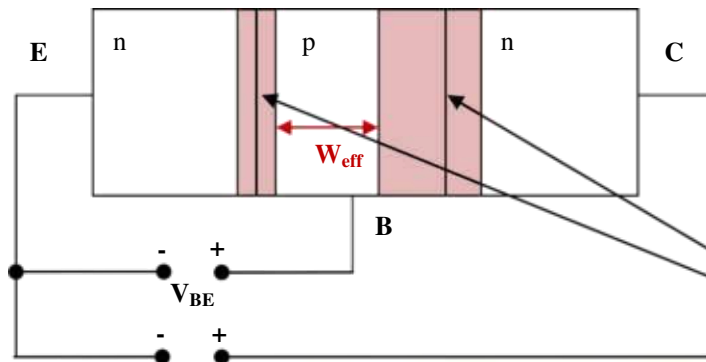
Early effect or base width modulation:

The early effect is the variation in the width of the base in a bipolar transistor due to a variation in the applied base-to-collector voltage. For example a greater reverse bias across the collector-base junction increases the collector-base depletion width. If V_{CE} increases V_{CB} increases too.



Depletion width

V_{CE1}



Depletion width

V_{CE2}

$$V_{BE} < V_{CE1} \ll V_{CE2}$$

The emitter-base junction is unchanged because the voltage V_{be} is the same. Base narrowing has two consequences that affect the current:

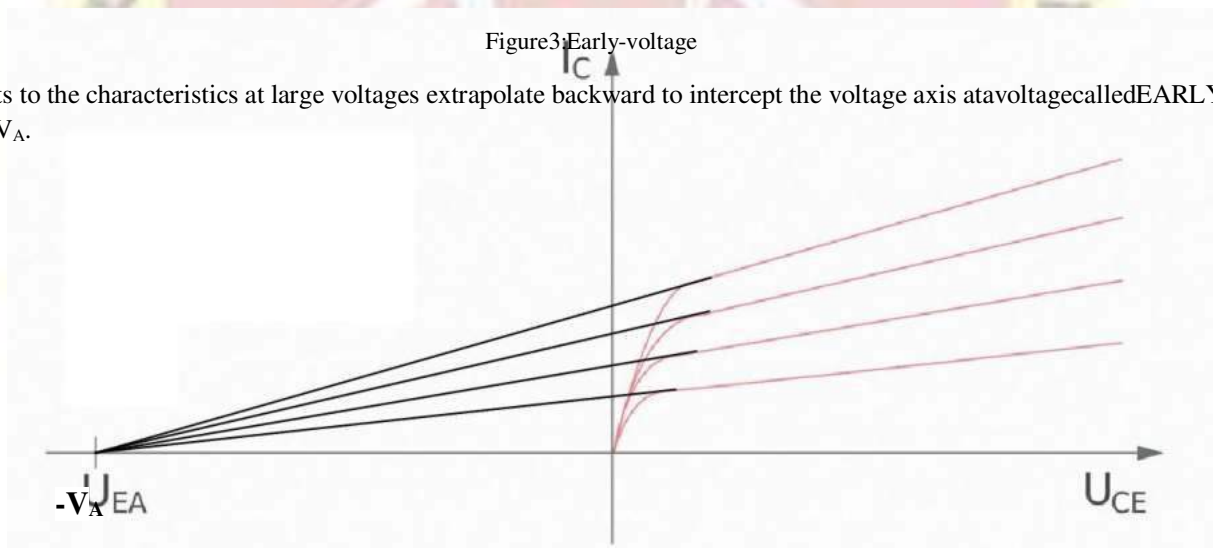
- There is a lesser chance for recombination within the "smaller" base region.
- The charge gradient is increased across the base, and consequently, the current of minority carriers injected across the emitter junction increases.

To counteract the Early-effect is a lightly doping of the collector region and a heavy doping of the emitter region.

Early-voltage:

Figure 3: Early-voltage

Tangents to the characteristics at large voltages extrapolate backward to intercept the voltage axis at a voltage called EARLY-voltage V_A .



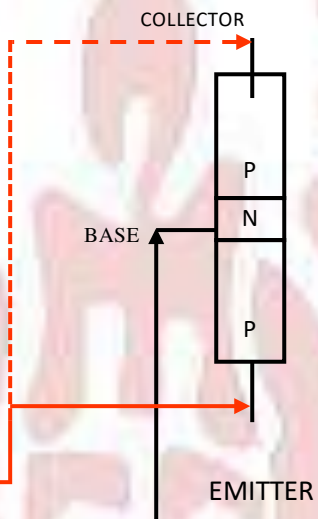
3.11 Transistor Test Procedure

An ohmmeter can be used to test the base-to-emitter PN junction and the base-to-collector PN junction of a bipolar junction transistor in the same way that a diode is tested. You can also identify the polarity (NPN or PNP) of an unknown device using this test. In order to do this you will need to be able to identify the emitter, base, and collector leads of the transistor. Refer to a semiconductor data reference manual if you are not sure of the lead identification. Note: While this test can be used to determine that the junctions are functional and that the transistor is not open or shorted, it will not convey any information about the common emitter current gain (amplification factor) of the device. A special transistor tester is required to measure this parameter known as the H_{fe} or Beta..



PNP Transistor

Simplified Diagram



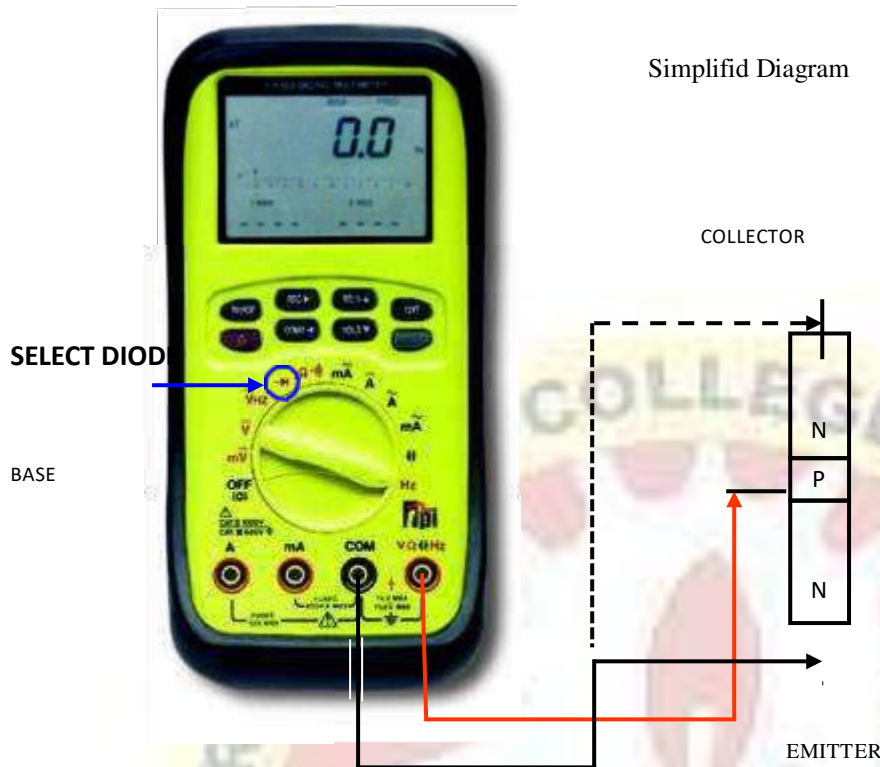
PNP Test Procedure

- Connect the meter leads with the polarity as shown and verify that the base-to-emitter and base-to-collector junctions read as a forward biased diode: 0.5 to 0.3 VDC.
- Reverse the meter connections to the transistor and verify that both PN junctions do not conduct. Meter should indicate an open circuit. (Display = OUCH or OL.)
- Finally read the resistance from emitter to collector and verify an open circuit reading in both directions. (Note: A short can exist from emitter to collector even if the individual PN junctions test properly.)

PNP Test Procedure

- Connect the meter leads with the polarity as shown and verify that the base-to-emitter and base-to-collector junctions read as a forward biased diode: 0.5 to 0.8 VDC.
- Reverse the meter connections to the transistor and verify that both PN junctions do not conduct. Meter should indicate an open circuit. (Display = OUCH or OL.)
- Finally read the resistance from emitter to collector and verify an open circuit reading in both directions. (Note: A short can exist from emitter to collector even if the individual PN junctions test properly.)

133DigitalMultimeterTransistor



NPN Test Procedure

- Connect the meter leads with the polarity as shown and verify that the base-to-emitter and base-to-collector junctions read as a forward biased diode: 0.5 to 0.3 VDC.
- Reverse the meter connections to the transistor and verify that both PN junctions do not conduct. Meter should indicate an open circuit. (Display = OUCH or OL.)
- Finally read the resistance from emitter to collector and verify an open circuit reading in both directions. (Note : A short can exist from emitter to collector even if the individual PN junctions test properly.)

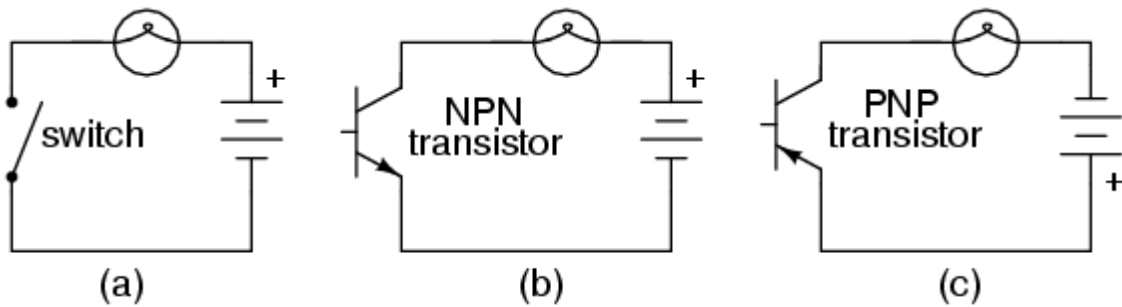
3.12 The transistor as a switch

Because a transistor's collector current is proportionally limited by its base current, it can be used as a sort of current-controlled switch. A relatively small flow of electrons sent through the base of the transistor has the ability to exert control over a much larger flow of electrons through the collector.

Suppose we had a lamp that we wanted to turn on and off with a switch. Such a circuit would be extremely simple as in Figure below(a).

For the sake of illustration, let's insert a transistor in place of the switch to show how it can control the flow of electrons through the lamp. Remember that the controlled current through a transistor must go between collector and emitter.

Since it is the current through the lamp that we want to control, we must position the collector and emitter of our transistor where the two contacts of the switch were. We must also make sure that the lamp's current will move *against* the direction of the emitter arrow symbol to ensure that the transistor's junction bias will be correct as in Figure below(b).

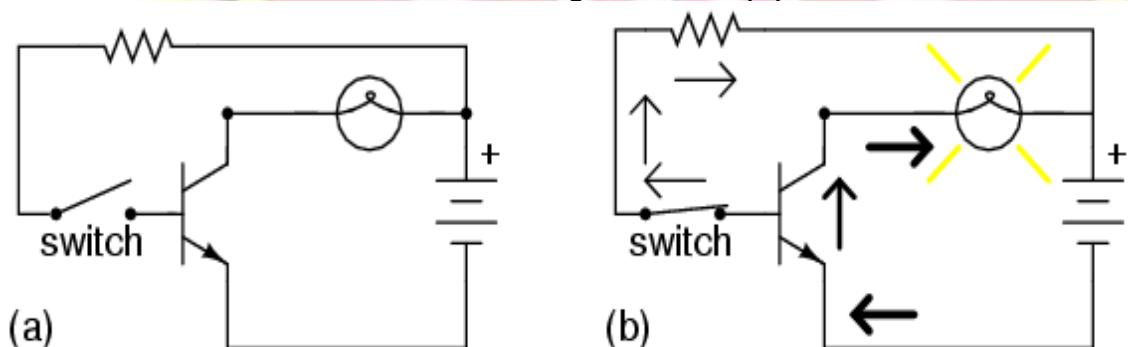


(a) mechanical switch, (b) NPN transistor switch, (c) PNP transistor switch.

A PNP transistor could also have been chosen for the job. Its application is shown in Figure above(c). The choice between NPN and PNP is really arbitrary. All that matters is that the proper current directions are maintained for the sake of correct junction biasing (electron flow going *against* the transistor symbol's arrow).

Going back to the NPN transistor in our example circuit, we are faced with the need to add something more so that we can have base current. Without a connection to the base wire of the transistor, base current will be zero, and the transistor cannot turn on, resulting in a lamp that is always off.

Remember that for an NPN transistor, base current must consist of electrons flowing from emitter to base (against the emitter arrow symbol, just like the lamp current). Perhaps the simplest thing to do would be to connect a switch between the base and collector wires of the transistor as in Figure below (a).



Transistor: (a) cutoff, lamp off; (b) saturated, lamp on.

If the switch is open as in (Figure above (a)), the base wire of the transistor will be left "floating" (not connected to anything) and there will be no current through it.

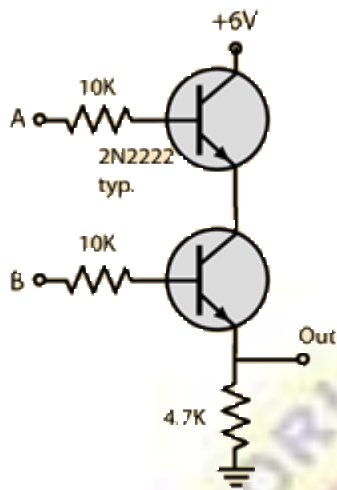
In this state, the transistor is said to be *cutoff*.

If the switch is closed as in (Figure above (b)), however, electrons will be able to flow from the emitter through to the base of the transistor, through the switch and up to the left side of the lamp, back to the positive side of the battery.

This base current will enable a much larger flow of electrons from the emitter through to the collector, thus lighting up the lamp. In this state of maximum circuit current, the transistor is said to be *saturated*. Of course, it may seem pointless to use a transistor in this capacity to control the lamp.

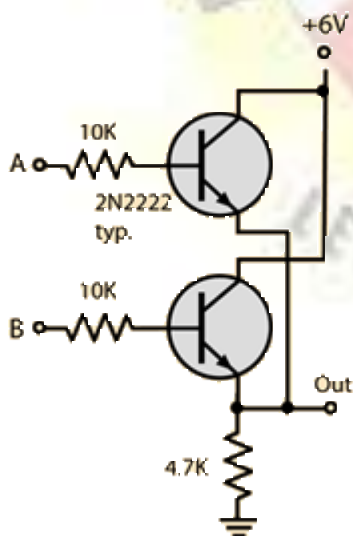
3.13 LOGIC GATES USING TRANSISTOR

TRANSISTOR AND GATE



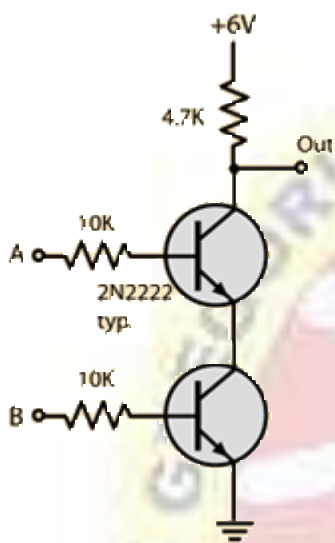
The use of [transistors](#) for the construction of logic [gates](#) depends upon their utility as fast [switches](#). When the base-emitter diode is turned on enough to be driven into [saturation](#), the collector voltage with respect to the emitter may be near zero and can be used to construct gates for the [TTL logic family](#). For the [AND](#) logic, the transistors are in series and both transistors must be in the conducting state to drive the output high

Transistor OR Gate



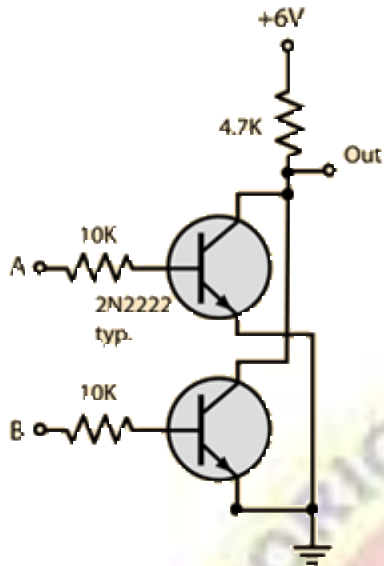
The use of [transistors](#) for the construction of logic [gates](#) depends upon their utility as fast [switches](#). When the base-emitter diode is turned on enough to be driven into [saturation](#), the collector voltage with respect to the emitter may be near zero and can be used to construct gates for the [TTL logic family](#). For the [OR](#) logic, the transistors are in parallel and the output is driven high if either of the transistors is conducting.

Transistor NAND Gate



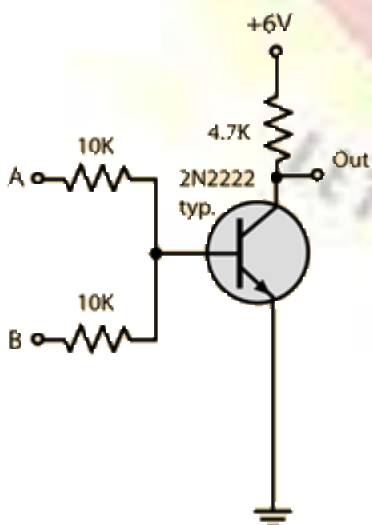
The use of [transistors](#) for the construction of logic [gates](#) depends upon their utility as fast [switches](#). When the base-emitter diode is turned on enough to be driven into [saturation](#), the collector voltage with respect to the emitter may be near zero and can be used to construct gates for the [TTL logic family](#). For the [NAND](#) logic, the transistors are in series, but the output is above them. The output is high unless both A and B inputs are high, in which case the output is taken down close to ground potential.

Transistor NOR GATE



The use of [transistors](#) for the construction of logic [gates](#) depends upon their utility as fast [switches](#). When the base-emitter diode is turned on enough to be driven into [saturation](#), the collector voltage with respect to the emitter may be near zero and can be used to construct gates for the [TTL logic family](#). For the [NOR](#) logic, the transistors are in parallel with the output above them so that if either or both of the inputs are high, the output is driven low.

Transistor NOR Gate



The use of [transistors](#) for the construction of logic [gates](#) depends upon their utility as fast [switches](#). When the base-emitter diode is turned on enough to be driven into [saturation](#),

the collector voltage with respect to the emitter may be near zero and can be used to construct gates for the [TTL logic family](#).

In this alternative way to achieve [NOR](#) logic, only one transistor is used with the two inputs tied to its base through resistors. If either or both of the inputs is high, the output is driven low

3.14 Unijunction Transistor (UJT)

A unijunction transistor (abbreviated as *UJT*) is a three-terminal semiconductor switching device. This device has a unique characteristic that when it is triggered, the emitter current increases regeneratively until it is limited by emitter power supply. Due to this characteristic, the unijunction transistor can be employed in a variety of applications *e.g.*, switching, pulse generator, saw-tooth generator etc.

Construction. Fig. 3.16(i) shows the basic structure of a unijunction

transistor. It consists of an *n*-type silicon bar with an electrical connection at each end. The leads to these connections are



Fig. 3.16

called base leads *base-one* B₁ and *base-two* B₂. Partway along the bar between the two bases, nearer to B₂ than B₁, a pn junction is formed between a p-type emitter and the bar. The lead to this junction is called the *emitter lead E*. Fig. 3.16 (ii) shows the symbol of unijunction transistor. Note that emitter is shown closer to B₂ than B₁. The following points are worth noting:

(i) Since the device has one pn junction and three leads, it is commonly called a unijunction transistor (*uni* means single).

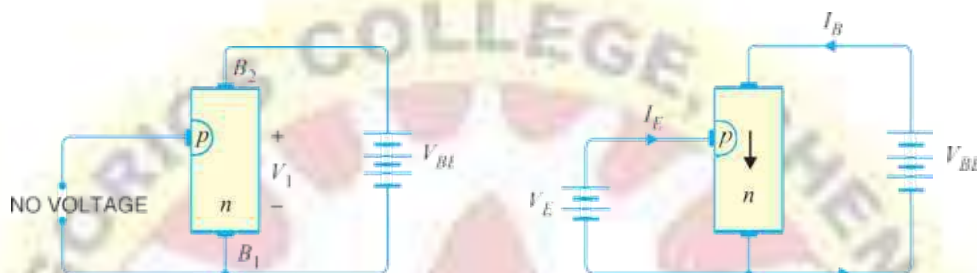
(ii) With only one pn junction, the device is really a form of diode. Because the two base terminals

nalsaretakenfromonesectionofthediode,thisdeviceisalsocalled *double-based diode*.

(iii) The emitter is heavily doped having many holes. The n region, however, is lightly doped. For this reason, the resistance between the base terminals is very high (5 to 10 k Ω) when emitter lead is open.

Operation. Fig. 3.17 shows the basic circuit operation of a unijunction transistor. The device has normally B_2 positive w.r.t. B_1 .

(i) If voltage V_{BB} is applied between B_2 and B_1 with emitter open [See Fig. 3.17(i)], a voltage gradient is established along the n -type bar. Since the emitter is located near to B_2 ,



(ii) more than half of V_{BB} appears between the emitter and B_1 . The voltage V_1 between emitter and B_1 establishes a

(i)

(ii)

reverse bias on the pn junction and the emitter current is cut off. Of course, a small leakage current flows from B_2 to emitter due to minority carriers.

(iii) If a positive voltage is applied at the emitter [See Fig. 21.25(ii)], the pn junction will remain reverse biased as long as the input voltage is less than V_1 . If the input voltage to the emitter exceeds V_1 , the pn junction becomes forward biased. Under these conditions, holes are injected from p -type material into the n -type bar. These holes are repelled by positive B_2 terminal and they are attracted towards B_1 terminal of the bar. This accumulation of holes in the emitter to B_1 region results in the decrease of resistance in this section of the bar. The result is that internal voltage drop from emitter to B_1 is decreased and hence the emitter current I_E increases. As more holes are injected, a condition of saturation will eventually be reached. At this point, the emitter current is limited by emitter power supply only. The device is now in the ON state.

(iv) If a negative pulse is applied to the emitter, the pn junction is reverse biased and the emitter current is cut off. The device is then said to be in the OFF state.

(v)

Equivalent Circuit of a UJT

Fig. 3.17

Fig. 3.18 shows the equivalent circuit of a UJT. The resistance of the silicon bar is called the inter-base resistance R_{BB} . The inter-base resistance is represented by two resistors in series viz.

(a) R_{B2} is the resistance of silicon bar between B_2 and the point at which the emitter junction lies.

(b) R_{B1} is the resistance of the bar between B_1 and emitter junction. This resistance is shown variable because its value depends upon the bias voltage across the pn junction.

The pn junction is represented in the emitter by a diode D .

The circuit action of a UJT can be explained more clearly from its equivalent circuit.

(i) With no voltage applied to the UJT, the inter-base resistance is given by;

$$R_{BB} = R_{B1} + R_{B2}$$

The value of R_{BB} generally lies between 4 k Ω and 10 k Ω .

(ii) If a

voltage V_{BB} is applied between the bases with emitter open, the voltage will divide across R_{B1} and R_{B2} .

$$\text{Voltage across } R_{B1}, V_{B1} = \frac{R_{B1}}{R_{B1} + R_{B2}} V_{BB}$$

*The main operational difference between the *JFET* and the *UJT* is that the *JFET* is normally operated with the gate junction reverse biased whereas the useful behaviour of the *UJT* occurs when the emitter is forward biased.

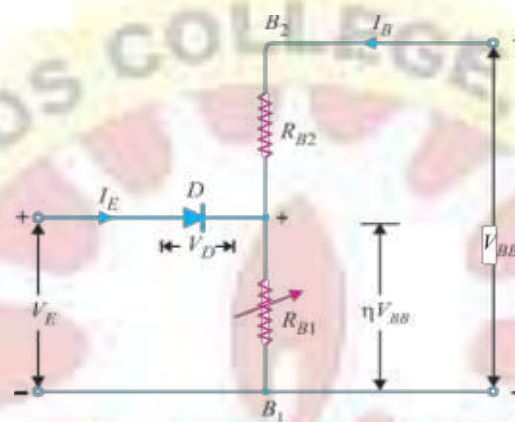


Fig.3.18

$$V_1/V_{BB} = \frac{R_{B1}R_B}{1+R_{B2}}$$

The ratio V_1/V_{BB} is called **intrinsic stand-off ratio** and is represented by Greek letter η .

$$\text{Obviously, } \eta = \frac{R_{B1}R_B}{1+R_{B2}}$$

The value of η usually lies between 0.51 and 0.32.

$\therefore$ Voltage across $R_{B1} = \eta V_{BB}$

The voltage ηV_{BB} appearing across R_{B1} reverse biases the diode. Therefore, the emitter current is zero.

(iii) If now a progressively rising positive voltage is applied to the emitter, the diode will become forward biased when input voltage exceeds ηV_{BB} by V_D , the forward voltage drop across the silicon diode i.e.

$$V_P = \eta V_{BB} + V_D$$

where

V_P = 'peak point voltage'

V_D = forward voltage drop across silicon diode (0.7V)

When the diode D starts conducting, holes are injected from p -type material to the n -type bar. These holes are swept down towards the terminal B_1 . This decreases the resistance between emitter and B_1 (indicated by variable resistance symbol for R_{B1}) and hence the internal drop from emitter to B_1 . The emitter current now increases regeneratively until it is limited by the emitter power supply.

Conclusion. The above discussion leads to the conclusion that when input positive voltage to the emitter is less than peak-point voltage V_P , the pn -junction remains reverse biased and the emitter current is practically zero. However, when the input voltage exceeds V_P , R_{B1} falls from several thousand ohms to a small value. The diode is now forward biased and the emitter current quickly reaches to a saturation value limited by R_{B1} (about 20A) and forward resistance of pn -junction (about 200A).

Characteristic of UJT

Fig. 3.19 shows the curve between emitter voltage (V_E) and emitter current (I_E) of a UJT at given voltage V_{BB} between the bases. This is known as the emitter characteristic of UJT. The following points may be noted from the characteristics:

(iv) Initially, in the cut-off region, as V_E increases from zero, slight leakage current flows from terminal B_2 to the emitter. This current is due to the minority carriers in the reverse biased diode.

(v) Above a certain value of V_E , forward I_E begins to flow, increasing until the peak voltage V_P and current I_P are reached at point P .

(vi) After the peak point P , an attempt to increase V_E is followed by a sudden increase in emitter current I_E with a corresponding decrease in V_E . This is a **negative resistance** portion of the curve because with increase in I_E , V_E decreases. The device, therefore, has a negative resistance region which is stable enough to be used with a great deal of reliability in many areas e.g., trigger circuits, saw-tooth generators, timing circuits.

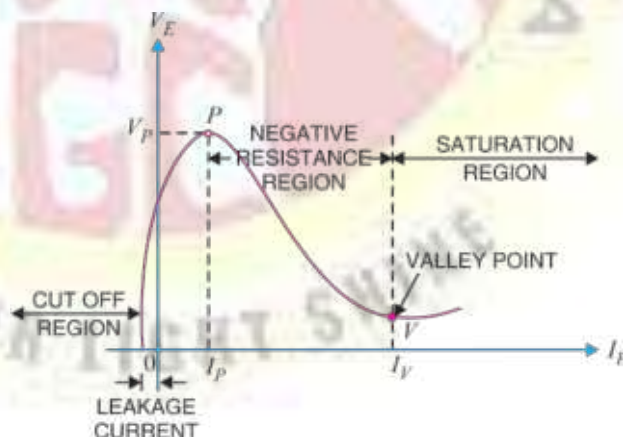


Fig.3.19

(vii) The negative portion of the curve lasts until the valley point V_V is reached with valley-point voltage V_V and valley-point current I_V . After the valley point, the device is driven to saturation.

Fig. 3.20 shows the typical family of V_E / I_E characteristics of a UJT at different voltages between the bases. It is clear that peak-point voltage ($= \eta V_{BB} + V_D$) falls steadily with reducing V_{BB} and so does the valley-point voltage V_V . The difference $V_P - V_V$ is a measure of the switching efficiency of UJT and can be seen to fall off as V_{BB} decreases. For a general purpose UJT, the peak-point current is of the order of $1 \mu A$ at $V_{BB} = 20V$ with a valley-point voltage

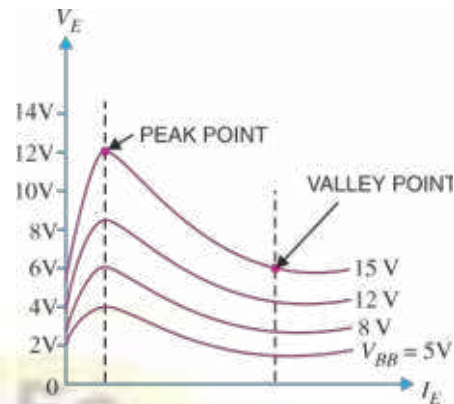


Fig.3.20

Advantages of UJT

The UJT was introduced in 1943 but did not become commercially available until 1952. Since then, the device has achieved great popularity due to the following reasons:

- (viii) It is a low cost device.
- (ix) It has excellent characteristics.
- (x) It is a low-power absorbing device under normal operating conditions.

Due to above reasons, this device is being used in a variety of applications. A few include oscillators, trigger circuits, saw-tooth generators, bistable network etc.

Applications of UJT

Unijunction transistors are used extensively in oscillator, pulse and voltage sensing circuits. Some of the important applications of UJT are discussed below:

(xi) **UJT Relaxation oscillator.** Fig. 3.21 shows UJT relaxation oscillator where the discharging of a capacitor through UJT can develop a saw-tooth output as shown.

When battery V_{BB} is turned on, the capacitor C charges through resistor R_1 . During the charging period, the voltage across the capacitor rises in an exponential manner until it reaches the peak-point voltage. At this instant of time, the UJT switches to its low resistance conducting mode and the capacitor is discharged between E and B_1 . As the capacitor voltage flies back to zero, the emitter ceases to conduct and the UJT is switched off. The next cycle then begins, allowing the capacitor C to charge again. The frequency of the output saw-tooth wave can be varied by changing the value of R_1 since this controls the time constant $R_1 C$ of the capacitor charging circuit.

The time period and hence the frequency of the saw-tooth wave can be calculated as follows. Assuming that the capacitor is initially uncharged, the voltage V_C across the capacitor prior to break-down is given by:

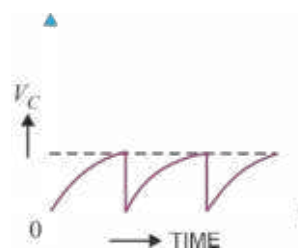
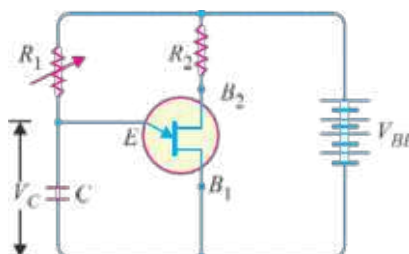


Fig.3.21

$$V_C = V_{BB} (1 - e^{-t/R_1 C})$$

where

$R_1 C$ = charging time constant of resistor-capacitor circuit

t = time from the commencement of waveform.

The discharge of the capacitor occurs when V_C is equal to the *peak-point voltage ηV_{BB} . i.e.

$$\eta V_{BB} = V_{BB} (1 - e^{-t/R_1 C})$$

or

$$\eta = 1 - e^{-t/R_1 C}$$

or

$$e^{-t/R_1 C} = 1 - \eta$$

or

$$t = R_1 C \log_{10} \frac{1}{1 - \eta}$$

∴

$$\text{Time period, } t = 2.3 R_1 C \log_{10} \frac{1}{1 - \eta}$$



$$\text{Frequency of saw-tooth wave, } f = \frac{1}{t \text{ in seconds}} \text{ Hz}$$

(xii) Overvoltage detector. Fig. 3.21 shows a simple d.c. over-voltage indicator. A warning pilot-lamp L is connected between the emitter and B_1 circuit. So long as the input voltage is less than the peak-point voltage (V_p) of the UJT , the device remains switched off. However, when the input voltage exceeds V_p , the UJT is switched on and the capacitor discharges through the low resistance path between terminals E and B_1 . The current flowing in the pilot lamp L lights it, thereby indicating the overvoltage in the circuit.

UNIT IV

FIELD EFFECT TRANSISTORS: FET-Construction-Working – statistics-Transfer characteristics – Parameters of FET – FET as an amplifier – MOSFET- Enhancement of MOSFET- Depletion MOSFET - Construction & Working – Drain characteristics of MOSFET - Comparison of JFET & MOSFET

4.0 TYPES of Field Effect Transistors

A bipolar junction transistor (BJT) is a current controlled device *i.e.*, output characteristics of the device are controlled by base current and not by base voltage. However, in a field effect transistor (FET), the output characteristics are controlled by input voltage (*i.e.*, electric field) and not by input current. This is probably the biggest difference between BJT and FET . There are two basic types of field effect transistors:

- (i) Junction field effect transistor ($JFET$)
- (ii) Metal oxide semiconductor field effect transistor ($MOSFET$)

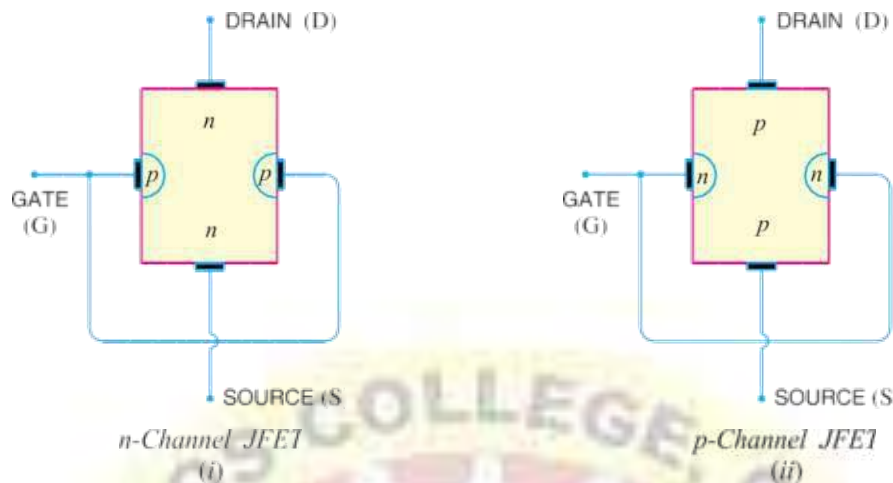
To begin with, we shall study about $JFET$ and then improved form of $JFET$, namely; $MOSFET$.

Junction Field Effect Transistor ($JFET$)

A **junction field effect transistor** is a three terminal semiconductor device in which current conduction is by one type of carrier *i.e.*, electrons or holes.

The $JFET$ was developed about the same time as the transistor but it came into general use only in the late 1960s. In a $JFET$, the current conduction is either by electrons or holes and is controlled by means of an electric field between the gate electrode and the conducting channel of the device. The $JFET$ has high input impedance and low noise level.

Constructional details. A $JFET$ consists of a p -type or n -type silicon bar containing two pn junctions at the sides as shown in Fig. 4.1. The bar forms the conducting channel for the charge carriers. If the bar is of n -type, it is called **n -channel $JFET$** as shown in Fig. 4.1 (i) and if the bar is of p -type, it is called a **p -channel $JFET$** as shown in Fig. 4.1 (ii). The two pn junctions forming diodes are connected internally and a common terminal called **gate** is taken out. Other terminals are **source** and **drain** taken out from the bar as shown. Thus a $JFET$ has essentially three terminals viz., **gate** (G), **source** (S) and **drain** (D).



JFET polarities. Fig. 4.2(i) shows *n*-channel *JFET* polarities whereas Fig. 4.2(ii) shows the *p*-channel *JFET* polarities. Note that in each case, the voltage between the gate and source is such that the gate is reverse biased. This is the normal way of *JFET* connection. The drain and source terminals are interchangeable i.e., either end can be used as source and the other end as drain.

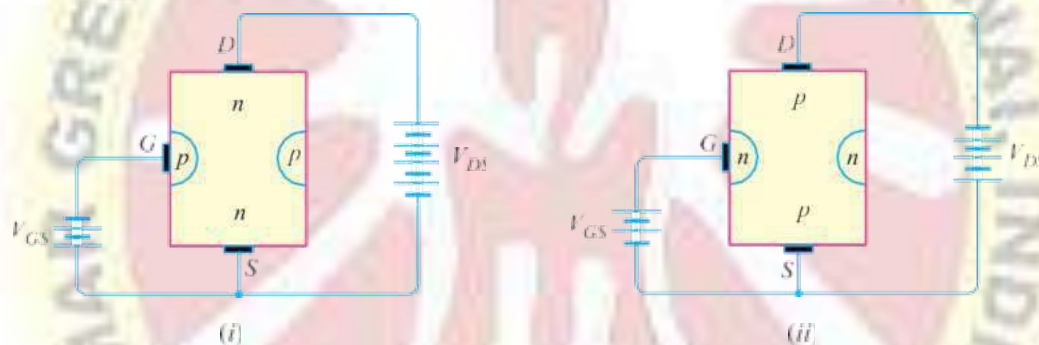


Fig. 4.2

The following points may be noted:

- (iii) The input circuit (i.e. gate to source) of a *JFET* is reverse biased. This means that the device has high input impedance.
- (iv) The drain is so biased w.r.t. source that drain current I_D flows from the source to drain.
- (v) In all *JFETs*, source current I_S is equal to the drain current i.e. $I_S = I_D$.

4.1 Principle and Working of JFET

Fig. 4.3 shows the circuit of *n*-channel *JFET* with normal polarities. Note that the gate is reverse biased.

Principle. The two *pn* junctions at the sides form two wide depletion layers. The current conduction by charge carriers (i.e. free electrons in this case) is through the channel between the two wide depletion layers and out of the drain. The width and hence resistance of this channel can be controlled by changing the input voltage V_{GS} . The greater the reverse voltage V_{GS} , the wider will be the depletion layers and narrower will be the conducting channel. The narrower channel means greater resistance and hence source to drain current decreases. Reverse will happen should V_{GS} decrease. **Thus, JFET operates on the principle that width and hence resistance of the conducting channel can be varied by changing the reverse voltage V_{GS} .** In other words, the magnitude of drain current (I_D) can be changed by altering V_{GS} .

Working. The working of *JFET* is as under:

- (vi) When a voltage V_{DS} is applied between drain and source terminals and voltage on the gate is zero [See Fig. 4.3 (i)], the two *pn* junctions at the sides of the bar establish depletion layers. The electrons will flow from source to drain through a channel between the depletion layers. The size of these

layers determine the width of the channel and hence the current conduction through the bar.

(vii) When a reverse voltage V_{GS} is applied between the gate and source [See Fig. 4.3(ii)], the width of the depletion layers is increased. This reduces the width of conducting channel, thereby increasing the resistance of n -type bar. Consequently, the current from source to drain is decreased. On the other hand, if the reverse voltage on the gate is decreased, the width of the depletion layers also decreases. This increases the width of the conducting channel and hence source to drain current.

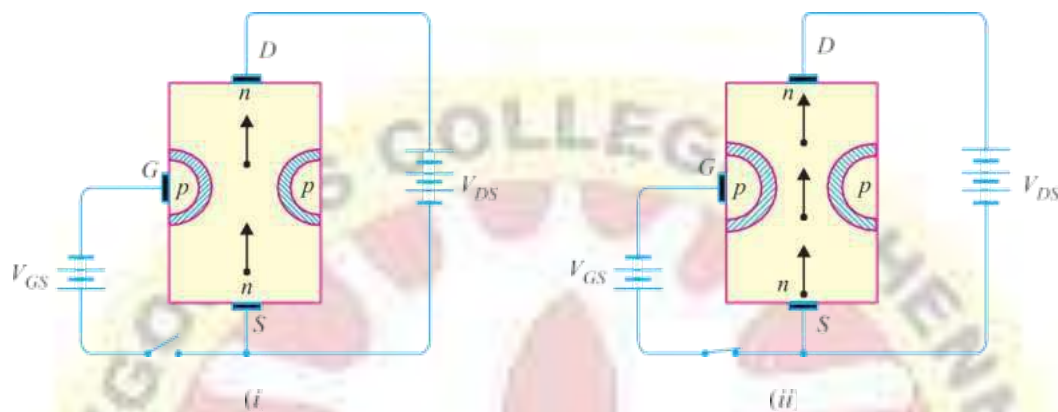
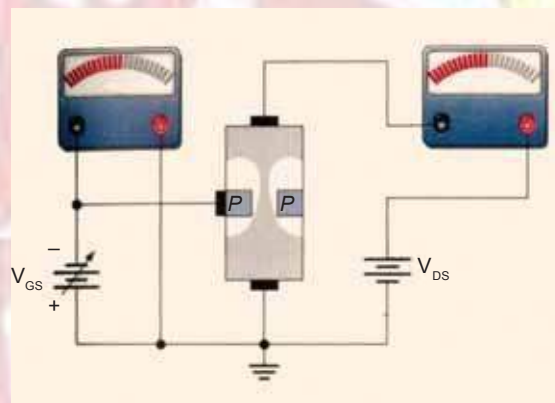


Fig.4.3

It is clear from the above discussion that current from source to drain can be controlled by the application of potential (i.e. electric field) on the gate. For this reason, the device is called **field effect transistor**. It may be noted that a p -channel $JFET$ operates in the same manner as an n -channel $JFET$ except that channel current carriers will be the holes instead of electrons and the polarities of V_{GS} and V_{DS} are reversed.

Note. If the reverse voltage V_{GS} on the gate is continuously increased, a state is reached when the two depletion layers touch each other and the channel is cut off. Under such conditions, the channel becomes an on-conductor.



Schematic Symbol of JFET

Fig.4.4 shows the schematic symbol of JFET. The vertical line in the symbol may be thought

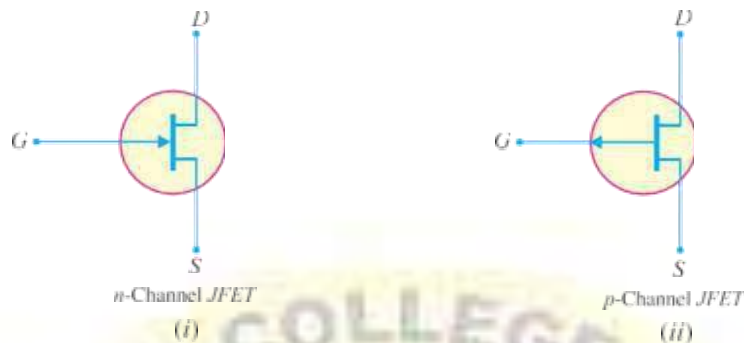


Fig.4.4

as channel and source (S) and drain (D) connected to this line. If the channel is *n*-type, the arrow on the gate points towards the channel as shown in Fig.4.4(i). However, for *p*-type channel, the arrow on the gate points from channel to gate [See Fig. 4.4(ii)].

Importance of JFET

A JFET acts like a voltage controlled device *i.e.* input voltage (V_{GS}) controls the output current. This is different from ordinary transistor (or bipolar transistor) where input current controls the output current. Thus JFET is a semiconductor device acting **like* a vacuum tube. The need for JFET arose because as modern electronic equipment became increasingly transistorised, it became apparent that there were many functions in which bipolar transistors were unable to replace vacuum tubes. Owing to their extremely high input impedance, JFET devices are more like vacuum tubes than are the bipolar transistors and hence are able to take over many vacuum-tube functions. Thus, because of JFET, electronic equipment is close to today to being completely solid state.

The JFET devices have not only taken over the functions of vacuum tubes but they now also threaten to displace the bipolar transistors as the most widely used semiconductor devices. As an amplifier, the JFET has higher input impedance than that of a conventional transistor, generates less noise and has greater resistance to nuclear radiations.

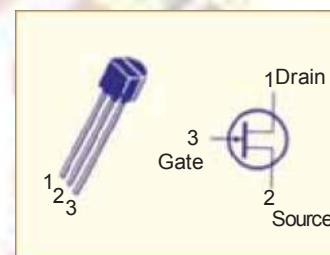
Difference Between JFET and Bipolar Transistor

The JFET differs from an ordinary or bipolar transistor in the following ways:

(viii) In a JFET, there is only one type of carrier, holes in *p*-type channel and electrons in *n*-type channel. For this reason, it is also called a *unipolar transistor*. However, in an ordinary transistor, both holes and electrons play part in conduction. Therefore, an ordinary transistor is sometimes called a *bipolar transistor*.

(ix) As the input circuit (*i.e.*, gate to source) of a JFET is reverse biased, therefore, the device has high input impedance. However, the input circuit of an ordinary transistor is forward biased and hence has low input impedance.

(x) The primary functional difference between the JFET and the BJT is that no current (actually, a very, very small current) enters the gate of JFET (*i.e.* $I_G = 0$ A). However, typical BJT base current might be a few μ A while JFET gate current is a thousand times smaller [See Fig. 4.5].



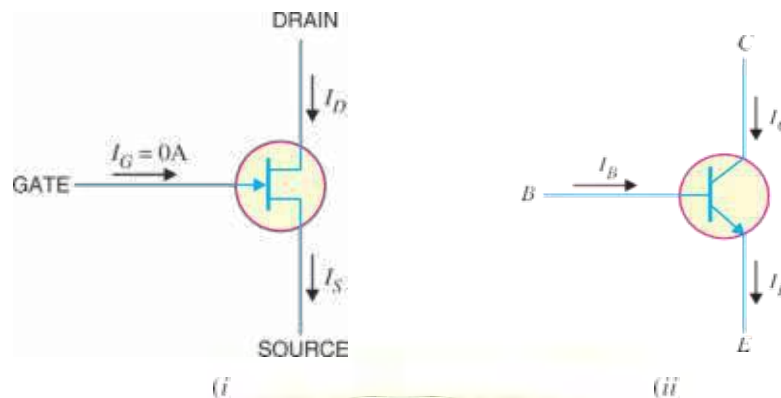


Fig.4.5

* The gate, source and drain of a *JFET* correspond to grid, cathode and anode of a vacuum tube.

(xi) A bipolar transistor uses a current into its base to control a large current between collector and emitter whereas a *JFET* uses voltage on the 'gate' (= base) terminal to control the current between drain (= collector) and source (= emitter). Thus a bipolar transistor gain is characterised by current gain whereas the *JFET* gain is characterised as a transconductance i.e., the ratio of change in output current (drain current) to the input (gate) voltage.

(xii) In *JFET*, there are no junctions as in an ordinary transistor. The conduction is through an *n*-type or *p*-type semi-conductor material. For this reason, noise level in *JFET* is very small.

4.2 *JFET* as an Amplifier

Fig. 4.6 shows *JFET* amplifier circuit. The weak signal is applied between gate and source and amplified output is obtained in the drain-source circuit. For the proper operation of *JFET*, the gate must be negative w.r.t. source i.e., input circuit should always be reverse biased. This is achieved either by inserting a battery V_{GG} in the gate circuit or by a circuit known as biasing circuit. In the present case, we are providing biasing by the battery V_{GG} .

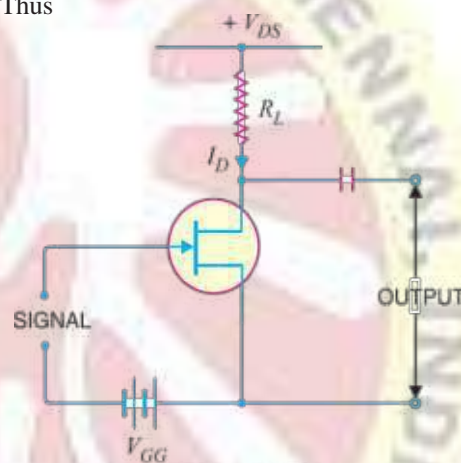
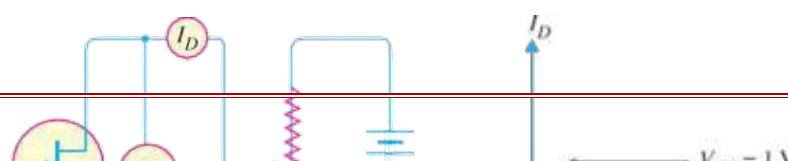


Fig.4.6

A small change in the reverse bias on the gate produces a large change in drain current. This fact makes *JFET* capable of raising the strength of a weak signal. During the positive half of signal, the reverse bias on the gate decreases. This increases the channel width and hence the drain current. During the negative half-cycle of the signal, the reverse voltage on the gate increases. Consequently, the drain current decreases. The result is that a small change in voltage at the gate produces a large change in drain current. These large variations in drain current produce large output across the load R_L . In this way, *JFET* acts as an amplifier.

4.3 Output Characteristics of *JFET*

The curve between drain current (I_D) and drain-source voltage (V_{DS}) of a *JFET* at constant gate-source voltage (V_{GS}) is called the output characteristics of *JFET*. Fig. 4.7 shows the circuit for determining the output characteristics of *JFET*. Keeping V_{GS} fixed at some value, say 1V, the drain-source voltage is changed in steps. Corresponding to each value of V_{DS} , the drain current I_D is noted. A plot of these values gives the output characteristic of *JFET* at $V_{GS}=1V$. Repeating similar procedure, output characteristics at other gate-source voltages can be drawn. Fig. 4.8 shows a family of output characteristics.



The following points may be noted from the characteristics:

(xiii) At first, the drain current I_D rises rapidly with drain-source voltage V_{DS} but then becomes constant. The drain-source voltage above which drain current becomes constant is known as **pinch-off voltage**. Thus in Fig. 4.8, $O A$ is the **pinch-off voltage** V_p .

(xiv) After pinch-off voltage, the channel width becomes so narrow that depletion layers almost touch each other. The drain current passes through the small passage between these layers. Therefore, increase in drain current is very small with V_{DS} above pinch-off voltage. Consequently, drain current remains constant.

(xv) The characteristics resemble that of a pentode valve.

Salient Features of JFET

The following are some salient features of JFET:

(xvi) A JFET is a three-terminal **voltage-controlled** semiconductor device i.e. input voltage controls the output characteristics of JFET.

(xvii) The JFET is **always** operated with gate-source *pn* junction *reverse biased.

(xviii) In a JFET, the gate current is zero i.e. $I_G = 0$ A.

(xix) Since there is no gate current, $I_D = I_S$.

(xx) The JFET must be operated between V_{GS} and $V_{GS(off)}$. For this range of gate-to-source voltages, I_D will vary from a maximum of I_{DSS} to a minimum of almost zero.

(xxi) Because the two gates are at the same potential, both depletion layers widen or narrow down by an equal amount.

(xxii) The JFET is not subjected to thermal runaway when the temperature of the device increases.

(xxiii) The drain current I_D is controlled by changing the channel width.

(xxiv) Since JFET has no gate current, there is no β rating of the device. We can find drain current I_D by using the eq. mentioned in Art. 4.11.

Important Terms

In the analysis of a JFET circuit, the following important terms are often used:

1. Shorted-gate drain current (I_{DSS})
2. Pinch-off voltage (V_p)
3. Gate-source cutoff voltage [$V_{GS(off)}$]

1. Shorted-gate drain current (I_{DSS}). It is the drain current with source short-circuited to gate (i.e. $V_{GS} = 0$) and drain voltage (V_{DS}) equal to pinch-off voltage. It is sometimes called zero-bias current.

Fig 4.9 shows the JFET circuit with $V_{GS} = 0$ i.e., source shorted-circuited to gate. This is normally called shorted-gate condition. Fig. 4.10 shows the graph between I_D and V_{DS} for the shorted gate condition. The drain current rises rapidly at first and then levels off at pinch-off voltage V_p . The drain current has now reached the maximum value I_{DSS} . When V_{DS} is increased beyond V_p , the depletion layers expand at the top of the channel. The channel now acts as a current limiter and

** holds drain current constant at I_{DSS} .

* Forward biasing gate-source *pn* junction may destroy the device.

** When drain voltage equals V_p , the channel becomes narrow and the depletion layers almost touch each other. The channel now acts as a current limiter and holds drain current at a constant value of I_{DSS} .

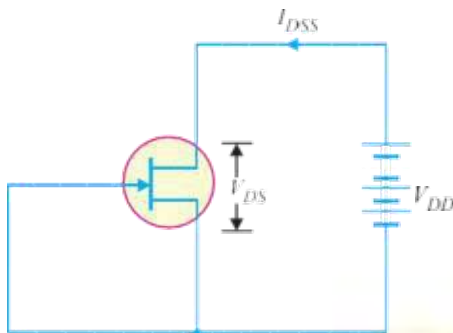


Fig. 4.9

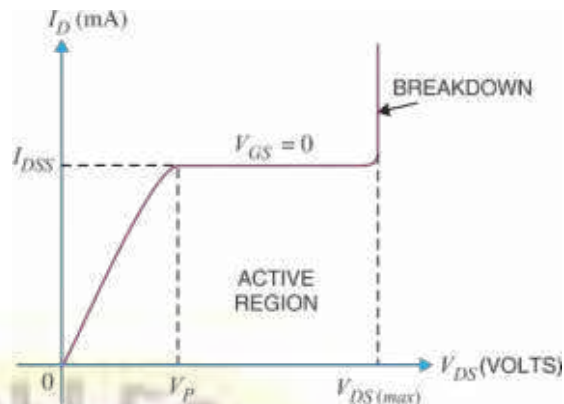


Fig. 4.10

The following points may be noted carefully:

(xxv) Since I_{DSS} is measured under shorted gate conditions, it is the maximum drain current that you can get with normal operation of *JFET*.

(xxvi) There is a maximum drain voltage [$V_{DS(max)}$] that can be applied to a *JFET*. If the drain voltage exceeds $V_{DS(max)}$, *JFET* would breakdown as shown in Fig. 4.10.

(xxvii) The region between V_P and $V_{DS(max)}$ (breakdown voltage) is called **constant-current region** or **active region**. As long as V_{DS} is kept within this range, I_D will remain constant for a constant value of V_{GS} . In other words, in the active region, *JFET* behaves as a constant-current device. For proper working of *JFET*, it must be operated in the active region.

(xxviii)

2. Pinch off Voltage (V_P). It is the minimum drain-source voltage at which the drain current essentially becomes constant.

Figure 4.11 shows the drain curves of a *JFET*. Note that pinch off voltage is V_P . The highest curve is for $V_{GS} = 0V$, the shorted-gate condition. For values of V_{DS} greater than V_P , the drain current is almost constant. It is because when V_{DS} equals V_P , the channel is effectively closed and does not allow further increase in drain current. It may be noted that for proper function of *JFET*, it is always operated for $V_{DS} > V_P$. However, V_{DS} should not exceed $V_{DS(max)}$ otherwise *JFET* may breakdown.

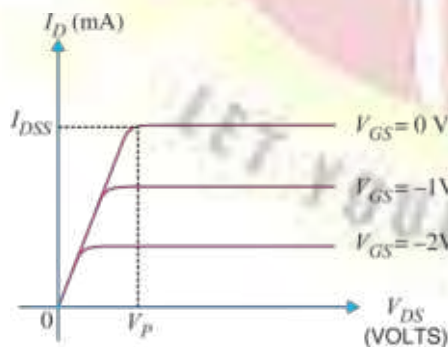


Fig. 4.11

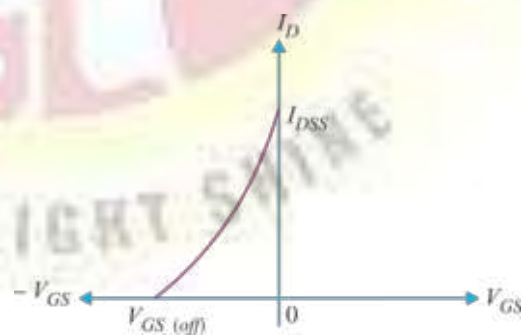


Fig. 4.12

3. Gate-source cut off voltage $V_{GS(off)}$. It is the gate-source voltage where the channel is completely cut off and the drain current becomes zero.

The idea of gate-source cut off voltage can be easily understood if we refer to the transfer characteristic of a JFET shown in Fig. 4.12. As the reverse gate-source voltage is increased, the cross-sectional area of the channel decreases. This in turn decreases the drain current. At some reverse gate-source voltage, the depletion layers extend completely across the channel. In this condition, the channel is cut off and the drain current reduces to zero. The gate voltage at which the channel is cut off (i.e. channel becomes non-conducting) is called gate-source cut off voltage $V_{GS(off)}$.

Notes. (i) It is interesting to note that $V_{GS(off)}$ will always have the same magnitude value as V_P . For example if $V_P = 6\text{ V}$, then $V_{GS(off)} = -6\text{ V}$. Since these two values are always equal and opposite, only one is listed on the specifications sheet for a given JFET.

(iii) There is a distinct difference between V_P and $V_{GS(off)}$. Note that V_P is the value of V_{DS} that causes the JFET to become a constant current device. It is measured at $V_{GS} = 0\text{ V}$ and will have a constant drain current $= I_{DSS}$. However, $V_{GS(off)}$ is the value of V_{GS} that causes I_D to drop to nearly zero.

Expression for Drain Current (I_D)

The relation between I_{DSS} and V_P is shown in Fig. 4.13. We note that gate-source cut off voltage [i.e. $V_{GS(off)}$] on the transfer characteristic is equal to pinch off voltage V_P on the drain characteristic i.e.

$$V_P = |V_{GS(off)}|$$

For example, if a JFET has $V_{GS(off)} = -4\text{ V}$, then $V_P = 4\text{ V}$.

The transfer characteristic of JFET shown in Fig. 4.13 is part of a parabola. A rather complex mathematical analysis yields the following expression for drain current:

$$I_D = I_{DSS} \left[1 - \frac{V_{GS}}{V_{GS(off)}} \right]^2$$

where

I_D = drain current at given
 V_{GS} , I_{DSS} = shorted – gate drain
 current V_{GS} = gate–source voltage
 $V_{GS(off)}$ = gate–source cutoff voltage

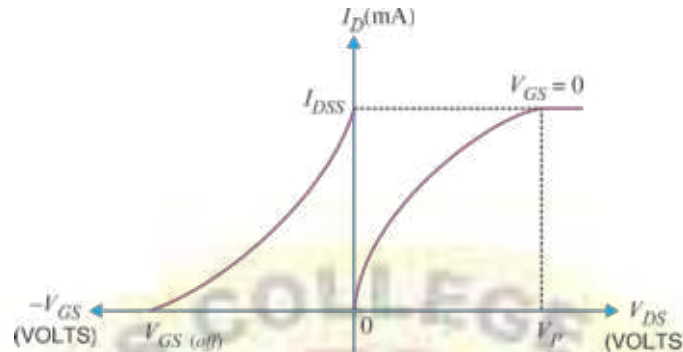


Fig.4.13

Advantages of JFET

A *JFET* is a voltage controlled, constant current device (similar to a vacuum pentode) in which variations in input voltage control the output current. It combines the many advantages of both bipolar transistor and vacuum pentode. Some of the advantages of a *JFET* are:

(xxix) It has a very high input impedance (of the order of 100 M Ω). This permits high degree of isolation between the input and output circuits.

(xxx) The operation of a *JFET* depends upon the bulk material current carriers that do not cross junctions. Therefore, the inherent noise of tubes (due to high-temperature operation) and those of transistors (due to junction transitions) are not present in a *JFET*.

(xxxi) A *JFET* has a negative temperature coefficient of resistance. This avoids the risk of thermal runaway.

(xxxii) A *JFET* has a very high power gain. This eliminates the necessity of using driver stages.

(xxxiii) A *JFET* has a small size, longer life and high efficiency.

(xxxiv)

Parameters of JFET

Like vacuum tubes, a *JFET* has certain parameters which determine its performance in a circuit. The main parameters of a *JFET* are (i) a.c. drain resistance (ii) transconductance (iii) amplification factor.

(xxxv) **a.c. drain resistance (r_d)**. Corresponding to the a.c. plate resistance, we have a.c. drain resistance in a *JFET*. It may be defined as follows:

It is the ratio of change in drain-source voltage (ΔV_{DS}) to the change in drain current (ΔI_D) at constant gate-source voltage i.e.

$$\text{a.c. drain resistance, } r_d = \frac{\Delta V_{DS}}{\Delta I_D} \text{ at constant } V_{GS}$$

For instance, if a change in drain voltage of 2V produces a change in drain current of 0.02mA, then,

$$\text{a.c. drain resistance, } r_d = \frac{2\text{V}}{0.02\text{mA}} = 100\text{k}\Omega$$

Referring to the output characteristics of a JFET in Fig. 4.8, it is clear that above the pinch-off voltage, the change in I_D is small for a change in V_{DS} because the curve is almost flat. Therefore, drain resistance of a JFET has a large value, ranging from 10k Ω to 1M Ω .

(xxxvi) Transconductance (g_{fs}). The control that the gate voltage has over the drain current is measured by transconductance g_{fs} and is similar to the transconductance g_m of the tube. It may be defined as follows:

It is the ratio of change in drain current (ΔI_D) to the change in gate-source voltage (ΔV_{GS}) at constant drain-source voltage i.e.

$$\text{Transconductance, } g_{fs} = \frac{\Delta I_D}{\Delta V_{GS}} \text{ at constant } V_{DS}$$

The transconductance of a JFET is usually expressed either in mA/volt or micromho. As an example, if a change in gate voltage of 0.1V causes a change in drain current of 0.3mA, then,

$$\begin{aligned} \text{Transconductance, } g_{fs} &= \frac{0.3\text{mA}}{0.1\text{V}} = 3\text{mA/V} = 3 \times 10^{-3} \text{A/V or mho or S (siemens)} \\ &= 3 \times 10^{-3} \times 10^6 \mu\text{mho} = 3000 \mu\text{mho (or } \mu\text{S)} \end{aligned}$$

(xxxvii) Amplification factor (μ). It is the ratio of change in drain-source voltage (ΔV_{DS}) to the change in gate-source voltage (ΔV_{GS}) at constant drain current i.e.

$$\text{Amplification factor, } \mu = \frac{\Delta V_{DS}}{\Delta V_{GS}} \text{ at constant } I_D$$

Amplification factor of a JFET indicates how much more control the gate voltage has over drain current than has the drain voltage. For instance, if the amplification factor of a JFET is 50, it means that gate voltage is 50 times as effective as the drain voltage in controlling the drain current.

Relation Among JFET Parameters

The relationship among JFET parameters can be established as under:

$$\text{We know } \mu = \frac{\Delta V_{DS}}{\Delta V_{GS}}$$

Multiplying the numerator and denominator on R.H.S. by ΔI_D , we get,

$$\mu = \frac{\Delta V_{DS} \times \Delta I_D}{\Delta V_{GS} \times \Delta I_D} = \frac{\Delta V_{DS}}{\Delta V_{GS}} \times \frac{\Delta I_D}{\Delta I_D}$$

$\therefore$

$$\mu = r_d \times g_{fs}$$

amplification factor = a.c. drain resistance $\times$ transconductance

Variation of Transconductance (g_m or g_{fs}) of JFET

We have seen that transconductance g_m of a JFET is the ratio of change in drain current (ΔI_D) to change in gate-source voltage (ΔV_{GS}) at constant V_{DS} i.e.

$$g_m = \frac{\Delta I_D}{\Delta V_{GS}}$$

The transconductance g_m of a JFET is an important parameter because it is a major factor in determining the voltage gain of JFET amplifiers. However, the transfer characteristic curve for a JFET is nonlinear so that the value of g_m depends upon the location on the curve. Thus the value of

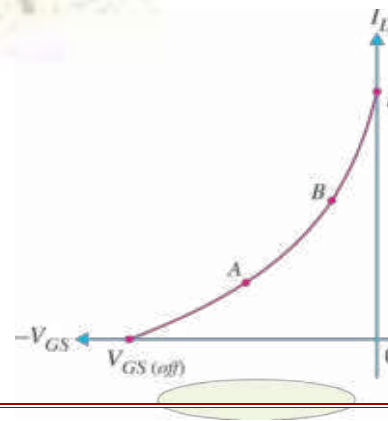


Fig. 4.17

$$g_m = g_{m0} \left(1 - \frac{V_{GS}}{V_{GS(off)}} \right)$$

where g_m = value of transconductance at any point on the transfer characteristic curve
 g_{m0} = value of transconductance (maximum) at $V_{GS} = 0$

Normally, the data sheet provides the value of g_{m0} . When the value of g_{m0} is not available, you can approximately calculate g_{m0} using the following relation:

$$g_{m0} = \frac{2I_{DSS}}{|V_{GS(off)}|}$$

JFET Biasing

For the proper operation of n -channel JFET, gate must be negative w.r.t. source. This can be achieved either by inserting a battery in the gate circuit or by a circuit known as biasing circuit. The latter method is preferred because batteries are costly and require frequent replacement.

1. Bias battery. In this method, JFET is biased by a bias battery V_{GG} . This battery ensures that gate is always negative w.r.t. source during all parts of the signal.

2. Biasing circuit. The biasing circuit uses supply voltage V_{DD} to provide the necessary bias. Two most commonly used methods are (i) self-bias (ii) potential divider method. We shall discuss each method in turn.

Self-Bias for JFET

In the self-bias method for n -channel JFET, the resistor R_S is the bias resistor. The d.c. component of drain current flowing through R_S produces the desired bias voltage.

$$\text{Voltage across } R_S, V_S = I_D R_S$$

Since gate current is negligibly small, the gate terminal is at d.c. ground i.e., $V_G = 0$.

$$\therefore V_{GS} = V_G - V_S = 0 - I_D R_S$$

$$\text{or } V_{GS} = -I_D R_S$$

Thus bias voltage V_{GS} keeps gate negative w.r.t. source.

Operating point. The operating point (i.e., zero signal I_D and V_{DS}) can be easily determined. Since the parameters of the *JFET* are usually known, zero signal I_D can be calculated from the following relation:

$$I_D = I_{DSS} \left(1 - \frac{V_{GS}}{V_{GS(off)}} \right)^2$$

Also
$$V_{DS} = V_{DD} - I_D(R_D + R_S)$$

Thus d.c. conditions of *JFET* amplifier are fully specified i.e. operating point for the circuit is V_{DS}, I_D .

Also,
$$R_S = \frac{|V_{GS}|}{|I_D|}$$

Note that gate resistor R_G does not affect bias because voltage across it is zero.

Midpoint Bias. It is often desirable to bias a *JFET* near the midpoint of its transfer characteristic curve where $I_D = I_{DSS}/2$. When signal is applied, the midpoint bias allows a maximum amount of drain current swing between I_{DSS} and 0. It can be proved that when $V_{GS} = V_{GS(off)}/3.4$, midpoint bias conditions are obtained for I_D .

$$I_D = I_{DSS} \left(1 - \frac{V_{GS}}{V_{GS(off)}} \right)^2 = I_{DSS} \left(1 - \frac{V_{GS(off)}/3.4}{V_{GS(off)}} \right)^2 = 0.5 I_{DSS}$$

To set the drain voltage at midpoint ($V_D = V_{DD}/2$), select a value of R_D to produce the desired voltage drop.

JFET with Voltage-Divider Bias

Fig. 4.15 shows potential divider method of biasing a *JFET*. This circuit is identical to that used for a transistor. The resistors R_1 and R_2 form a voltage divider across drain supply V_{DD} . The voltage $V_2 (= V_G)$ across R_2 provides the necessary bias.

$$V_2 = V_G = \frac{V_{DD} R_2}{R_1 + R_2}$$

Now $V_2 = V_{GS} + I_D R_S$

or $V_{GS} = V_2 - I_D R_S$

The circuit is so designed that $I_D R_S$ is larger than V_2 so that V_{GS} is negative. This provides correct bias voltage. We can find the operating point as under:

$$I_D = \frac{V_2 - V_{GS} R_S}{R_S}$$

and
$$V_{DS} = V_{DD} - I_D(R_D + R_S)$$

Although the circuit of voltage-divider bias is a bit complex, yet the advantage of this method of biasing is that it provides good stability of the operating point. The input impedance Z_i of this circuit is given by;

$$Z_i = R_1 \parallel R_2$$

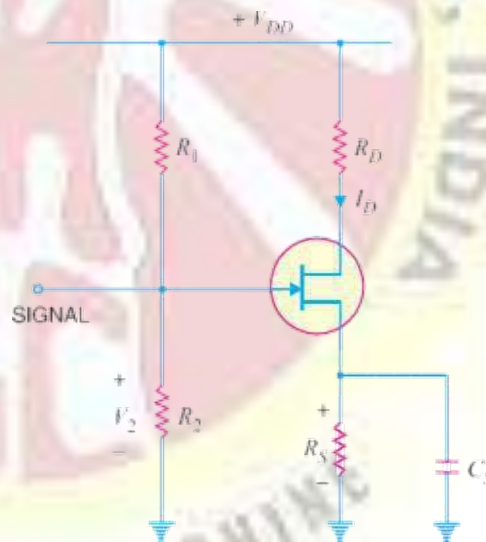


Fig.4.15

JFET Connections

There are three leads in a *JFET* viz., source, gate and drain terminals. However, when *JFET* is to be connected in a circuit, we require four terminals; two for the input and two for the output. This difficulty is overcome by making one terminal of the *JFET* common to both input and output terminals. Accordingly, a *JFET* can be connected in a circuit in the following three ways:

- (i) Common source connection (ii) Common gate connection
- (iii) Common drain connection

The common source connection is the most widely used arrangement. It is because this connection provides high input impedance, good voltage gain and a moderate output impedance. However, the circuit produces a phase reversal *i.e.*, output signal is 180° out of phase with the input signal. Fig. 4.29 shows a common source *n*-channel *JFET* amplifier. Note that source terminal is common to both input and output.

Note. A common source *JFET* amplifier is the *JFET* equivalent of common emitter amplifier. Both amplifiers have a 180° phase shift from input to output. Although the two amplifiers serve the same basic purpose, the means by which they operate are quite different.

4.5 Metal Oxide Semiconductor FET (MOSFET)

The main drawback of *JFET* is that its gate *must* be reverse biased for proper operation of the device *i.e.* it can only have negative gate operation for *n*-channel and positive gate operation for *p*-channel. This means that we can *only* decrease the width of the channel (*i.e.* decrease the conductivity of the channel) from its zero-bias size. This type of operation is referred to as *depletion-mode* operation. Therefore, a *JFET* can only be operated in the depletion-mode. However, there is a field effect transistor (*FET*) that can be operated to enhance (or increase) the width of the channel (with consequent increase in conductivity of the channel) *i.e.* it can have *enhancement-mode* operation. Such a *FET* is called *MOSFET*.

A field effect transistor (*FET*) that can be operated in the enhancement-mode is called a **MOSFET**.

A *MOSFET* is an important semiconductor device and can be used in any of the circuits covered for *JFET*. However, a *MOSFET* has several advantages over *JFET* including high input impedance and low cost of production.

Types of MOSFETs

There are two basic types of *MOSFETs* viz.

1. **Depletion-type MOSFET or D-MOSFET.** The *D-MOSFET* can be operated in both the depletion-mode and the enhancement-mode. For this reason, a *D-MOSFET* is sometimes called *depletion/enhancement MOSFET*.
2. **Enhancement-type MOSFET or E-MOSFET.** The *E-MOSFET* can be operated *only* in enhancement-mode.

The manner in which a *MOSFET* is constructed determines whether it is *D-MOSFET* or *E-MOSFET*.

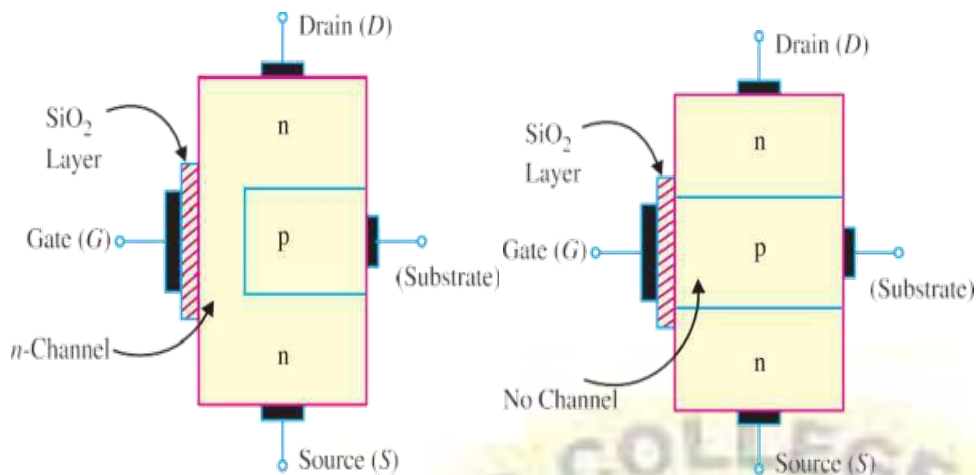
1. D-MOSFET. Fig. 4.43 shows the constructional details of *n*-channel *D-MOSFET*. It is similar to *n*-channel *JFET* except with the following modifications/remarks:

(i) The *n*-channel *D-MOSFET* is a piece of *n*-type material with a *p*-type region (called *substrate*) on the right and an *insulated gate* on the left as shown in Fig. 4.43. The free electrons (Q) in *n*-channel flowing from source to drain must pass through the narrow channel between the gate and the *p*-type region (*i.e.* substrate).

(ii) Note carefully the gate construction of *D-MOSFET*. A thin layer of metal oxide (usually silicon dioxide, SiO_2) is deposited over a small portion of the channel. A metallic gate is deposited over the oxide layer. As SiO_2 is an insulator, therefore, gate is insulated from the channel. Note that the arrangement forms a capacitor. One plate of this capacitor is the gate and the other plate is the channel with SiO_2 as the dielectric. Recall that we have a gated diode in a *JFET*.

(iii) It is a usual practice to connect the substrate to the source (*S*) internally so that a *MOSFET* has three terminals viz. *source* (*S*), *gate* (*G*) and *drain* (*D*).

(iv) Since the gate is insulated from the channel, we can apply either negative or positive voltage to the gate. Therefore, *D-MOSFET* can be operated in both depletion-mode and enhancement-mode. However, *JFET* can be operated only in depletion-mode.



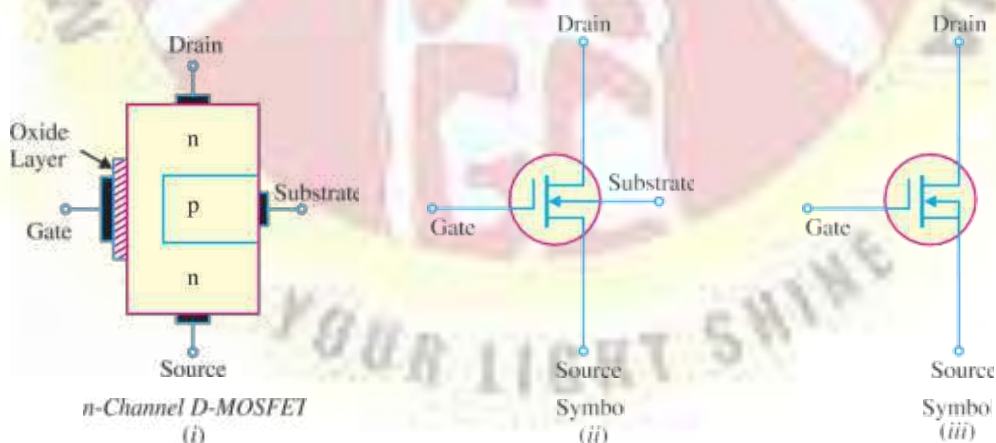
E-MOSFET. Fig. 4.16 shows the constructional details of *n*-channel *E-MOSFET*. Its gate construction is similar to that of *D-MOSFET*. The *E-MOSFET* has no channel between source and drain unlike the *D-MOSFET*. Note that the substrate extends completely to the SiO_2 layers so that no channel exists. The *E-MOSFET* requires a proper gate voltage to *form* a channel (called induced channel). It is reminded that *E-MOSFET* can be operated *only* in enhancement mode. In short, the construction of *E-MOSFET* is quite similar to that of the *D-MOSFET* except for the absence of a channel between the drain and source terminals.

The SiO_2 layer is an insulator. The gate terminal is made of a metal conductor. Thus, going from gate to substrate, you have a *metal oxide semiconductor* and hence the name *MOSFET*. Since the gate is insulated from the channel, the *MOSFET* is sometimes called *insulated-gate FET (IGFET)*. However, this term is rarely used in place of the term *MOSFET*.

4.6 Symbols for D-MOSFET

There are two types of *D-MOSFET* viz (i) *n*-channel *D-MOSFET* and (ii) *p*-channel *D-MOSFET*.

(i) ***n*-channel D-MOSFET.** Fig. 4.17 (i) shows the various parts of *n*-channel *D-MOSFET*. The *p*-type substrate constricts the channel between the source and drains so that only a small passage



remains at the left side. Electrons flowing from source (when drain is positive w.r.t. source) must pass through this narrow channel. The symbol for *n*-channel *D-MOSFET* is shown in Fig. 4.17 (ii). The gate appears like a capacitor plate. Just to the right of the gate is a thick vertical line representing the channel. The drain lead comes out of the top of the channel and the source lead connects to the bottom. The arrow is on the substrate and points to the *n*-material, therefore we have *n*-channel *D-MOSFET*. It is a usual practice to connect the substrate to source internally as shown in Fig. 4.17 (iii). This gives rise to a three-terminal device.

(ii) p -channel D-MOSFET. Fig. 4.18 (i) shows the various parts of p -channel D-MOSFET. Then- type substrate constricts the channel between the source and drain so that only a small passage remains at the left side. The conduction takes place by the flow of holes from source to drain through this narrow channel. The symbol for p -channel D-MOSFET is shown in Fig. 4.18 (ii). It is a usual practice to connect the substrate to source internally. This results in a three-terminal device

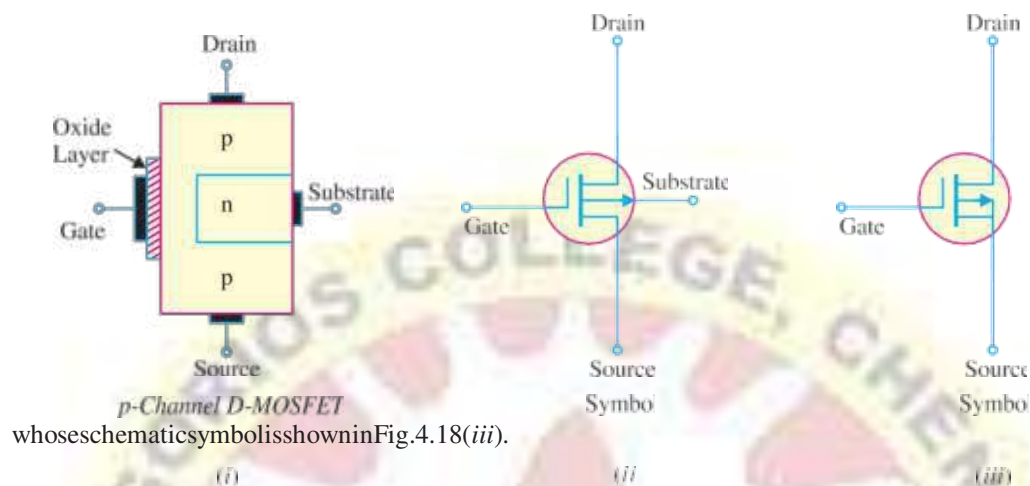


Fig.4.18

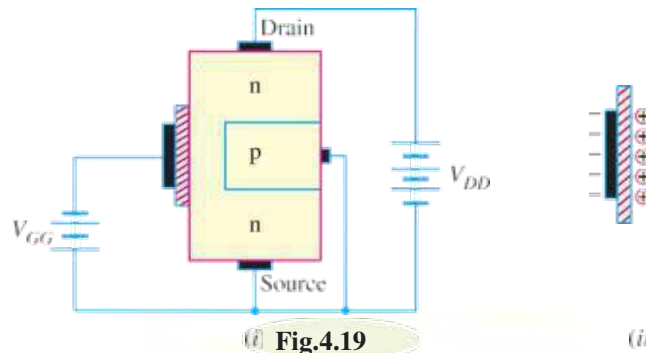
Circuit Operation of D-MOSFET

Fig.4.19(i) shows the circuit of n -channel D-MOSFET. The gate forms a small capacitor. One plate of this capacitor is the gate and the other plate is the channel with metal oxide layer as the dielectric. When gate voltage is changed, the electric field of the capacitor changes which in turn changes the resistance of the n -channel. Since the gate is insulated from the channel, we can apply either negative or positive voltage to the gate. The negative-gate operation is called **depletion mode** whereas positive-gate operation is known as **enhancement mode**.

(iii) Depletion mode. Fig.4.19(i) shows depletion-mode operation of n -channel D-MOSFET. Since gate is negative, it means electrons are on the gate as shown in Fig.4.19(ii). These electrons

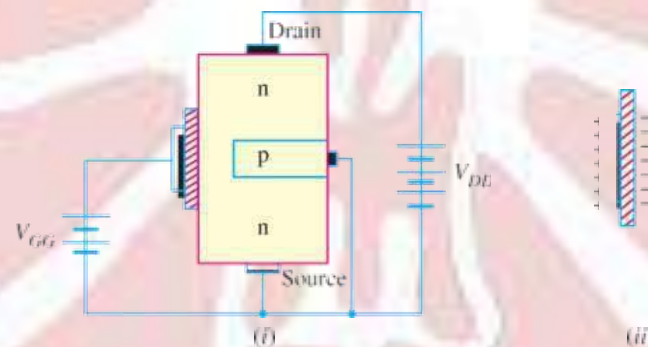
*repel the free electrons in the n -channel, leaving a layer of positive ions in a part of the channel as shown in Fig.4.19(ii). In other words, we have depleted (*i.e.* emptied) the n -channel of some of its free electrons. Therefore, lesser number of free electrons are made available for current conduction through the n -channel. This is the same thing as if the resistance of the channel is increased. The greater the negative voltage on the gate, the lesser is the current from source to drain.

Thus by changing the negative voltage on the gate, we can vary the resistance of the n -channel and hence the current from source to drain. Note that with negative voltage to the gate, the action of D-MOSFET is similar to JFET. Because the action with negative gate depends upon depleting (*i.e.* emptying) the channel of free electrons, the negative-gate operation is called **depletion mode**.



(iv) Enhancement mode. Fig. 4.20 (i) shows enhancement-mode operation of *n*-channel *D-MOSFET*. Again, the gate acts like a capacitor. Since the gate is positive, it induces negative charges in the *n*-channel as shown in Fig. 4.20 (ii). These negative charges are the free electrons drawn into the channel. Because these free electrons are added to those already in the channel, the total number of free electrons in the channel is increased. Thus a positive gate voltage *enhances* or *increases* the conductivity of the channel. The greater the positive voltage on the gate, the greater the conduction from source to drain.

Thus by changing the positive voltage on the gate, we can change the conductivity of the channel. The main difference between *D-MOSFET* and *JFET* is that we can apply positive gate voltage to *D-MOSFET* and still have essentially *zero* current. Because the action with a positive gate depends upon *enhancing* the conductivity of the channel, the positive gate operation is called *enhancement mode*.



The following points may be noted about *D-MOSFET* operation:

- (i)** In a *D-MOSFET*, the source to drain current is controlled by the electric field of capacitor formed at the gate.
- (ii)** The gate of *JFET* behaves as a reverse-biased diode whereas the gate of a *D-MOSFET* acts like a capacitor. For this reason, it is possible to operate *D-MOSFET* with positive or negative gate voltage.
- (iii)** As the gate of *D-MOSFET* forms a capacitor, therefore, negligible gate current flows whether

positive or negative voltage is applied to the gate. For this reason, the input impedance of *D-MOSFET* is very high, ranging from 10,000 M Ω to 10,000,000 M Ω .

(iv) The extremely small dimensions of the oxide layer under the gate terminal result in a very low capacitance and the *D-MOSFET* has, therefore, a very low input capacitance. This characteristic makes the *D-MOSFET* useful in high-frequency applications.

4.7 *D-MOSFET* Transfer Characteristic

Fig. 4.21 shows the transfer characteristic curve (or transconductance curve) for *n-channel D-MOSFET*. The behaviour of this device can be beautifully explained with the help of this curve as under:

- (v) The point on the curve where $V_{GS}=0, I_D=I_{DSS}$. It is expected because I_{DSS} is the value of I_D when gate and source terminals are shorted i.e. $V_{GS}=0$.
- (vi) As V_{GS} goes **negative**, I_D decreases below the value of I_{DSS} till I_D reaches zero when $V_{GS}=V_{GS(off)}$ just as with *JFET*.
- (vii) When V_{GS} is **positive**, I_D increases above the value of I_{DSS} . The maximum allowable value of I_D is given on the data sheet of *D-MOSFET*.

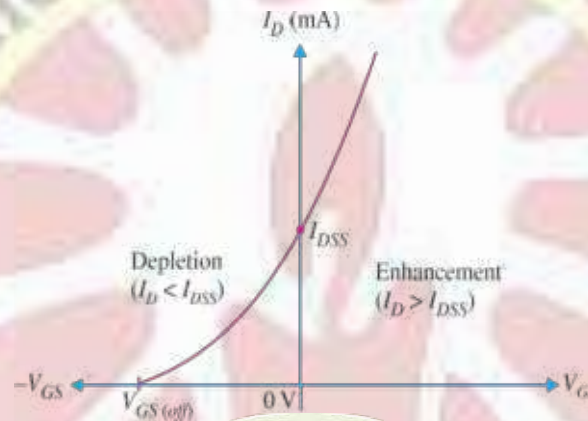


Fig. 4.21

Note that the transconductance curve for the *D-MOSFET* is very similar to the curve for a *JFET*. Because of this similarity, the *JFET* and the *D-MOSFET* have the same transconductance equation viz.

Transconductance and Input Impedance of D-MOSFET

These are important parameters of a D-MOSFET and a brief discussion on them is desirable.

(viii) **D-MOSFET Transconductance (g_m)**. The value of g_m is found for a D-MOSFET in the same way that it is for the JFET, i.e.,

$$g_m = g_{mo} \left(1 - \frac{V_{GS}}{V_{GS(off)}} \right)$$

(ix) **D-MOSFET Input Impedance**. The gate impedance of a D-MOSFET is extremely high. For example, a typical D-MOSFET may have a maximum gate current of 10 pA when $V_{GS} = 35$ V.

$$\therefore \text{Input impedance} = \frac{35 \text{ V}}{10 \text{ pA}} = \frac{35 \text{ V}}{10 \times 10^{-12} \text{ A}} = 3.5 \times 10^{12} \Omega$$

With an input impedance in this range, D-MOSFET would present virtually no load to a source circuit.

4.7 D-MOSFET Biasing

The following methods may be used for D-MOSFET biasing:

4.7.1

Gate bias (ii) Self-bias

(iii) Voltage-divider bias

(iv) Zero bias

The first three methods are exactly the same as those used for JFETs and are not discussed here. However, the last method of zero-bias is widely used in D-MOSFET circuits.

Zero bias. Since a D-MOSFET can be operated with either positive or negative values of V_{GS} , we can set its Q-point at $V_{GS} = 0$ V as shown in Fig. 4.22. Then an input a.c. signal to the gate can produce variations above and below the Q-point.

*

We can only change V_{GS} because the values of I_{DSS} and $V_{GS(off)}$ are constant for a given D-MOSFET.

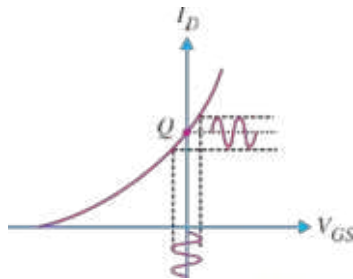


Fig.4.22

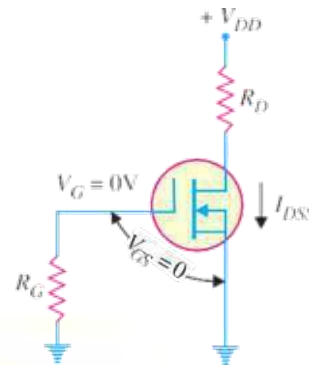


Fig.4.23

We can use the simple circuit of Fig. 4.23 to provide zero bias. This circuit has $V_{GS}=0V$ and $I_D=I_{DSS}$. We can find V_{DS} as under:

$$V_{DS}=V_{DD}-I_{DSS}R_D$$

Note that for the *D-MOSFET* zero bias circuit, the source resistor (R_S) is not necessary. With no source resistor, the value of V_S is 0 V. This gives us a value of $V_{GS}=0V$. This biases the circuit at $I_D=I_{DSS}$ and $V_{GS}=0V$. For mid-point biasing, the value of R_D is so selected that $V_{DS}=V_{DD}/2$.

Common-Source *D-MOSFET* Amplifier

Fig. 4.24 shows a common-source amplifier using *n-channel D-MOSFET*. Since the source terminal is common to the input and output terminals, the circuit is called *common-source amplifier. The circuit is zero biased with an a.c. source coupled to the gate through the coupling capacitor C_1 . The gate is at approximately 0V d.c. and the source terminal is grounded, thus making $V_{GS}=0V$.

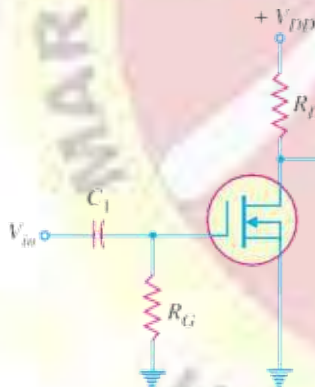


Fig.4.24

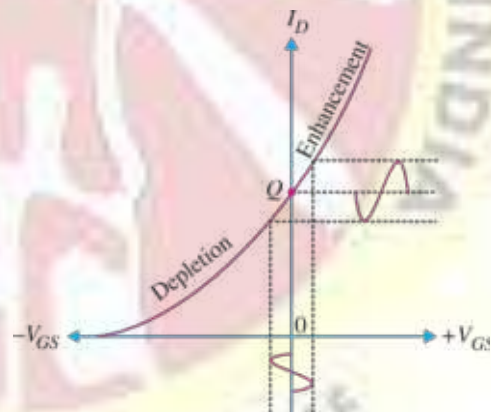


Fig.4.25

* It is comparable to common-emitter transistor amplifier

Operation. The input signal (V_{in}) is capacitively coupled to the gate terminal. In the absence of the signal, d.c. value of $V_{GS}=0V$. When signal (V_{in}) is applied, V_{gs} swings above and below its zero value (d.c. value of $V_{GS}=0V$), producing a swing in drain current I_d .

4.7.2 A small change in gate voltage produces a large change in drain current as in a *JFET*. This fact makes *MOSFET* capable of raising the strength of a weak signal; thus acting as an amplifier.

4.7.3 During the positive half-cycle of the signal, the positive voltage on the gate increases and produces the enhancement-mode. This increases the channel conductivity and hence the drain current.

4.7.4 During the negative half-cycle of the signal, the positive voltage on the gate decreases and produces depletion-mode. This decreases the conductivity and hence the drain current.

The result of above action is that a small change in gate voltage produces a large change in the drain current. This large variation in drain current produces a large a.c. output voltage across drain resistance R_D . In this way, *D-MOSFET* acts as an amplifier. Fig. 4.25 shows the amplifying action of *D-MOSFET* on transconductance curve.

Voltage gain. The a.c. analysis of *D-MOSFET* is similar to that of the *JFET*. Therefore, voltage gain expressions derived for *JFET* are also applicable to *D-MOSFET*.

Voltage gain, $A_v = g_m R_D$... for unloaded *D-MOSFET* amplifier
 $= g_m R_{AC}$... for loaded *D-MOSFET* amplifier Not the total a.c. drain resistance $R_{AC} = R_D || R_L$

D-MOSFETs Versus JFETs

Table below summarises many of the characteristics of *JFETs* and *D-MOSFETs*.

Devices:	JFETs	D-MOSFETs
Schematic symbol:		
Transconductance curve:		
Modes of operation:	Depletion only	Depletion and enhancement
Commonly used bias circuits:	Gate bias Self bias Voltage-divider bias	Gate bias Self bias Voltage-divider bias Zero bias
Advantages:	Extremely high input impedance.	Higher input impedance than comparable <i>JFET</i> . Can operate in both modes (depletion and enhancement).
Disadvantages:	Bias instability.	Bias instability.

	Can operate only in the depletion mode.	More sensitive to changes in temperature than the JFE T.
--	---	--

4.9 E-MOSFET

Two things are worth noting about *E-MOSFET*. First, *E-MOSFET* operates *only* in the enhancement mode and has no depletion mode. Secondly, the *E-MOSFET* has no physical channel from source to drain because the substrate extends completely to the SiO_2 layer [See Fig. 4.26(i)]. It is only by the application of V_{GS} (gate-to-source voltage) of proper magnitude and polarity that the device starts conducting. The minimum value of V_{GS} of proper polarity that turns on the *E-MOSFET* is called **Threshold voltage** [$V_{GS(th)}$]. Then *n*-channel device requires positive V_{GS} ($\geq V_{GS(th)}$) and the *p*-channel device requires negative V_{GS} ($\leq -V_{GS(th)}$).

Operation. Fig. 4.26 (i) shows the circuit of *n*-channel *E-MOSFET*. The circuit action is as under:

4.7.5 When $V_{GS} = 0\text{V}$ [See Fig. 4.26(i)], there is no channel connecting the source and drain. The substrate has only a few thermally produced free electrons (minority carriers) so that drain current is essentially zero. For this reason, *E-MOSFET* is normally *OFF* when $V_{GS} = 0\text{V}$. Note that this behaviour of *E-MOSFET* is quite different from *JFET* or *D-MOSFET*.

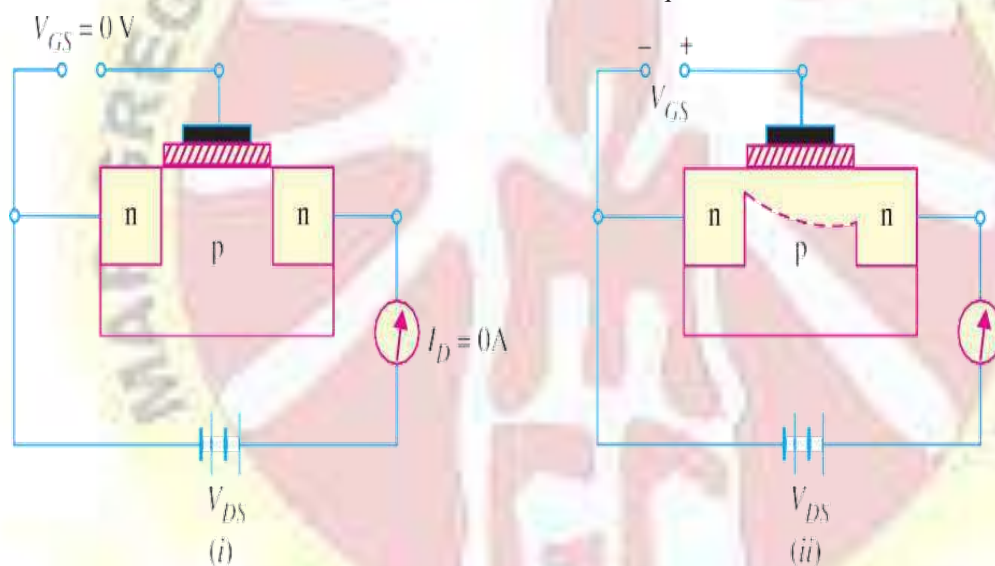


Fig.4.26

4.7.6 When gate is made positive (i.e. V_{GS} is positive) as shown in Fig. 4.26 (ii), it attracts free electrons into the p region. The free electrons combine with the holes next to the SiO_2 layer. If V_{GS} is positive enough, all the holes touching the SiO_2 layer are filled and free electrons begin to flow from the source to drain. The effect is the same as creating a thin layer of n -type material (i.e. inducing a thin n -channel) adjacent to the SiO_2 layer. Thus the E -MOSFET is turned ON and drain current I_D starts flowing from the source to the drain.

The minimum value of V_{GS} that turns the E -MOSFET ON is called **threshold voltage** [$V_{GS(th)}$].

4.7.7 When V_{GS} is less than $V_{GS(th)}$, there is no induced channel and the drain current I_D is zero. When V_{GS} is equal to $V_{GS(th)}$, the E -MOSFET is turned ON and the induced channel conducts drain current from the source to the drain. Beyond $V_{GS(th)}$, if the value of V_{GS} is increased, then the newly formed channel becomes wider, causing I_D to increase. If the value of V_{GS} decreases [not less than $V_{GS(th)}$], the channel becomes narrower and I_D will decrease. This fact is revealed by the transconductance curve of n -channel E -MOSFET shown in Fig. 4.57. As you can see, $I_D = 0$ when $V_{GS} = 0$. Therefore, the value of I_{DSS} for the E -MOSFET is zero. Note also that there is no drain current until V_{GS} reaches $V_{GS(th)}$.

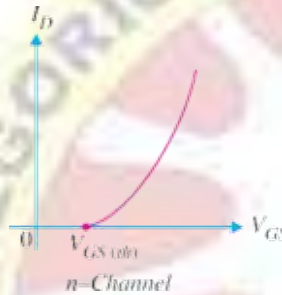


Fig.4.27

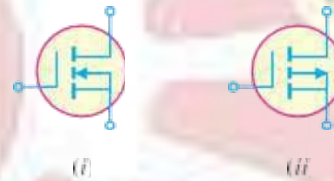


Fig.4.28

Schematic Symbols. Fig. 4.28 (i) shows the schematic symbols for n -channel E -MOSFET whereas Fig. 4.28 (ii) shows the schematic symbol for p -channel E -MOSFET. When $V_{GS} = 0$, the E -MOSFET is OFF because there is no conducting channel between source and drain. The broken channel line in the symbols indicates the normally OFF condition.

Equation for Transconductance Curve. Fig. 4.29 shows the transconductance curve for n -channel E -MOSFET. Note that this curve is different from the transconductance curve for n -channel JFET or n -channel D -MOSFET. It is because it starts at $V_{GS(th)}$ rather than $V_{GS(off)}$ on the horizontal axis and never intersects the vertical axis. The equation for the E -MOSFET transconductance curve (for $V_{GS} > V_{GS(th)}$) is

$$I_D = K(V_{GS} - V_{GS(th)})^2$$

The constant K depends on the particular E -MOSFET and its value is determined from the following equation:

$$K = \frac{I_{D(on)}}{(V_{GS(on)} - V_{GS(th)})^2}$$

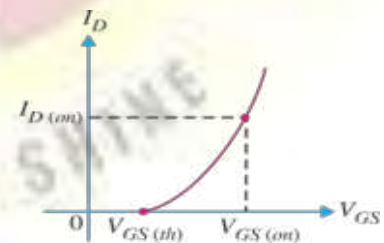
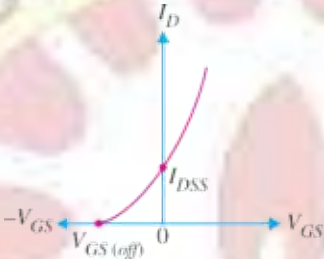
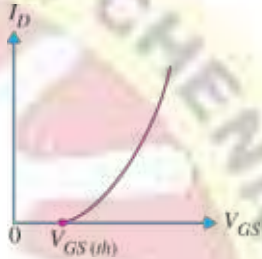


Fig4.29

Any datasheet for an *E-MOSFET* will include the current $I_{D(on)}$ and the voltage $V_{GS(on)}$ for one point well above the threshold voltage as shown in Fig. 4.29.

D-MOSFETs Versus E-MOSFETs

Table below summarizes many of the characteristics of *D-MOSFETs* and *E-MOSFETs*

Devices:	D-MOSFETs	E-MOSFETs
Schematic symbol:		
Transconductance curve:		
Modes of operation:	Depletion and enhancement.	Enhancement only.
Commonly used bias circuits:	Gate bias Self bias Voltage-divider bias Zero bias	Gate bias Voltage-divider bias Drain-feedback bias

UNIT V

POWER DEVICE: Power transistor – SCR – Triac – Diac – IGBT –Characteristics and working

5.0 SiliconControlledRectifier(SCR)

A silicon*controlled rectifier is a semiconductor** device that acts as a true electronics switch. It can

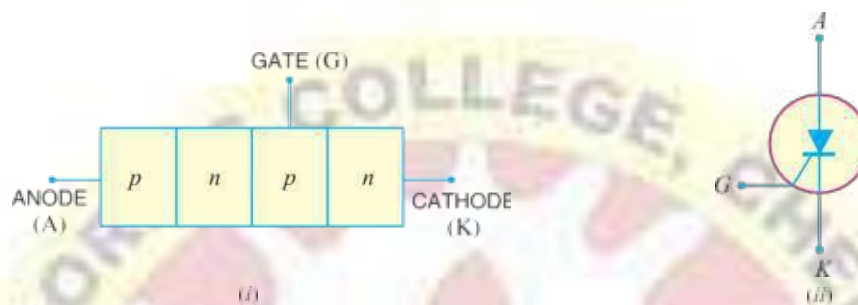


Fig.5.1

change alternating current into direct current and at the same time can control the amount of power fed to the load. Thus, SCR combines the features of a rectifier and a transistor.

Constructional details. When a pn junction is added to a junction transistor, the resulting three pn junction device is called a silicon controlled rectifier. Fig. 5.1 (i) shows its construction. It is clear that it is essentially an ordinary rectifier (pn) and a junction transistor (nnp) combined in one unit to form a pn device. Three terminals are taken; one from the outer p -type material called anode A , second from the outer n -type material called cathode K and the third from the base of transistor section and is called gate G . In the normal operating conditions of SCR, anode is held at high positive potential w.r.t. cathode and gate at small positive potential w.r.t. cathode. Fig. 5.1 (ii) shows the symbol of SCR.

The silicon controlled rectifier is a solid state equivalent of thyatron. The gate, anode and cathode of SCR correspond to the grid, plate and cathode of thyatron. For this reason, SCR is sometimes called *thyristor*.



TypicalSCRPackages

Working of SCR

In a silicon controlled rectifier, load is connected in series with anode. The anode is always kept at positive potential w.r.t. cathode. The working of SCR can be studied under the following two heads:

(i) **When gate is open.** Fig. 5.2 shows the SCR circuit with gate open i.e. no voltage applied to the gate. Under this condition, junction J_2 is reverse biased while junctions J_1 and J_3 are forward biased. Hence, the situation in the junctions J_1 and J_3 is just as in a $n-p-n$ transistor with base open. Consequently, no current flows through the load R_L and the SCR is *cut off*. However, if the applied voltage is gradually increased, a stage is reached when *reverse biased junction J_2 breaks down*. The SCR now conducts *heavily* and is said to be in the *ON* state. The applied voltage at which SCR conducts heavily without gate voltage is called *Breakover voltage*.

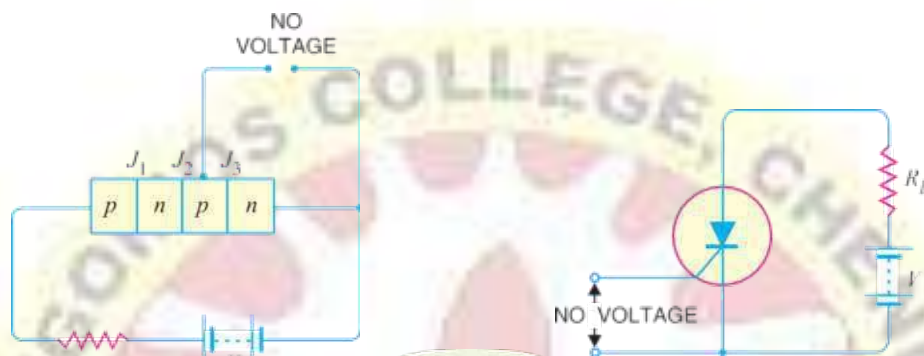


Fig.5.2

(ii) **When gate is positive w.r.t. cathode.** The SCR can be made to conduct heavily at smaller applied voltage by applying a small positive potential to the gate as shown in Fig. 5.3. Now junction J_3 is forward biased and junction J_2 is reverse biased. The electrons from n -type material start moving across junction J_3 towards left whereas holes from p -type towards the right. Consequently, the electrons from junction J_3 are attracted across junction J_2 and gate current starts flowing. As soon as the gate current flows, anode current increases. The increased anode current in turn makes more electrons available at junction J_2 . This process continues and in an extremely small time, junction J_2 breaks down and the SCR starts conducting heavily. *Once SCR starts conducting, the gate (the reason for this name is obvious) loses all control*. Even if gate voltage is removed, the anode current does not decrease at all. The only way to stop conduction (i.e. bring SCR in off condition) is to reduce the applied voltage to zero.

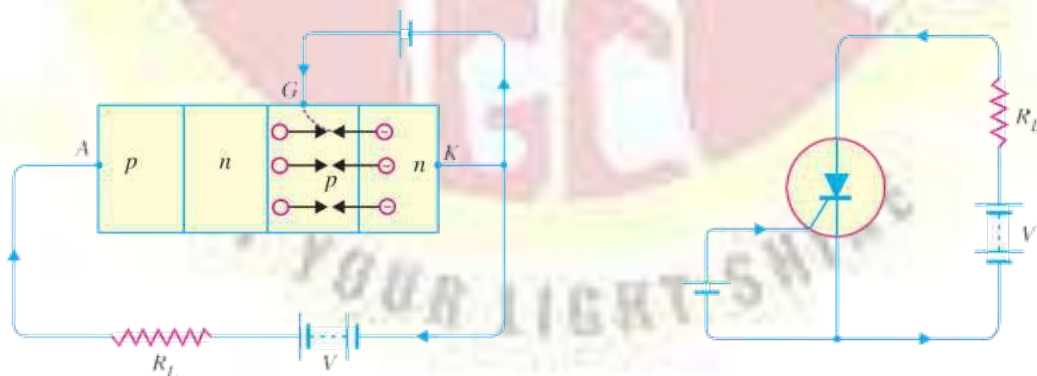


Fig.5.3

- * The whole applied voltage V appears as reverse bias across junction J_2 as junctions J_1 and J_3 are forward biased.
- ** Because J_1 and J_3 are forward biased and J_2 has broken down.

Conclusion. The following conclusions are drawn from the working of SCR:

- (i) An SCR has two states, i.e. either it does not conduct or it conducts heavily. There is no state in between. Therefore, SCR behaves like a switch.
- (ii) There are two ways to turn on the SCR. The first method is to keep the gate open and make the supply voltage equal to the breakover voltage. The second method is to operate SCR with supply voltage less than breakover voltage and then turn it on by means of a small voltage (typically 1.5V, 30 mA) applied to the gate.
- (iii) Applying a small positive voltage to the gate is the normal way to close an SCR because the breakover voltage is usually much greater than supply voltage.
- (iv) To open the SCR (i.e. to make it non-conducting), reduce the supply voltage to zero.

Equivalent Circuit of SCR

The SCR shown in Fig. 5.4(i) can be visualised as separated into two transistors as shown in

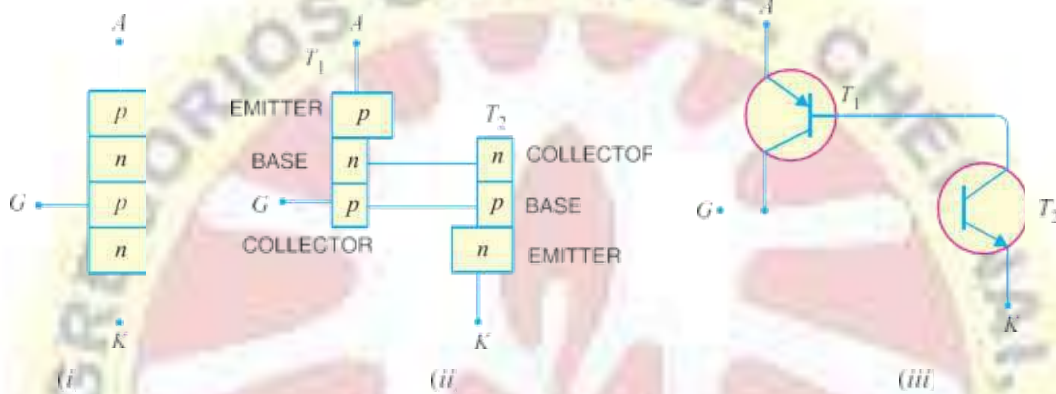


Fig.5.4

Fig. 5.4 (ii). Thus, the equivalent circuit of SCR is composed of pnp transistor and npn transistor connected as shown in Fig. 5.4. (iii). It is clear that collector of each transistor is coupled to the base of the other, thereby making a positive feedback loop.

The working of SCR can be easily explained from its equivalent circuit. Fig. 5.5 shows the equivalent circuit of SCR with supply voltage V and load resistance R_L . Assume the supply voltage V is less than breakover voltage as is usually the case. With gate open (i.e. switch S open), there is no base current in transistor T_2 . Therefore, no current flows in the collector of T_2 and hence that of T_1 . Under such conditions, the SCR is open. However, if switch S is closed, a small gate current will flow through the base of T_2 which means its collector current will increase. The collector current of T_2 is the base current of T_1 . Therefore, collector current of T_1 increases. But collector current of T_1 is the base current of T_2 . This action is accumulative since an increase of current in

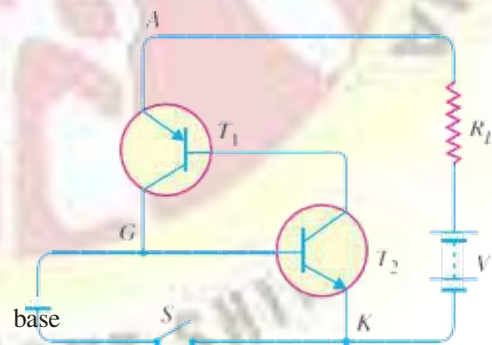


Fig.5.5

one transistor causes an increase of current in the other transistor. As a result of this action, both

transistors are driven to saturation, and heavy current flows through the load R_L . Under such conditions, the SCR closes.

Important Terms

The following terms are much used in the study of SCR:

- (iii) Breakover voltage
- (ii) Peak reverse voltage
- (iii) Holding current
- (iv) Forward current rating
- (v) Circuit fusing rating
- (i) **Breakover voltage.** It is the minimum forward voltage, gate being open, at which SCR starts conducting heavily i.e. turned on.

Thus, if the breakover voltage of an SCR is 200 V, it means that it can block a forward voltage (i.e. SCR remains open) as long as the supply voltage is less than 200 V. If the supply voltage is more than this value, then SCR will be turned on. In practice, the SCR is operated with supply voltage less than breakover voltage and it is then turned on by means of a small voltage applied to the gate. Commercially available SCRs have breakover voltages from about 50 V to 500 V.

- (ii) **Peak reverse voltage (PRV).** It is the maximum reverse voltage (cathode positive w.r.t. anode) that can be applied to an SCR without conducting in the reverse direction.

Peak reverse voltage (PRV) is an important consideration while connecting an SCR in an a.c. circuit. During the negative half of a.c. supply, reverse voltage is applied across SCR. If PRV is exceeded, there may be avalanche breakdown and the SCR will be damaged if the external circuit does not limit the current. Commercially available SCRs have PRV ratings up to 2.5 kV.

- (iii) **Holding current.** It is the maximum anode current, gate being open, at which SCR is turned off from ON conditions.

As discussed earlier, when SCR is in the conducting state, it cannot be turned OFF even if gate voltage is removed. The only way to turn off or open the SCR is to reduce the supply voltage to almost zero at which point the internal transistor comes out of saturation and opens the SCR. The anode current under this condition is very small (a few mA) and is called *holding current*. Thus, if an SCR has a holding current of 5 mA, it means that if anode current is made less than 5 mA, then SCR will be turned off.

- (iv) **Forward current rating.** It is the maximum anode current that an SCR is capable of passing without destruction.

Every SCR has a safe value of forward current which it can conduct. If the value of current exceeds this value, the SCR may be destroyed due to intensive heating at the junctions. For example, if an SCR has a forward current rating of 40 A, it means that the SCR can safely carry only 40 A. Any attempt to exceed this value will result in the destruction of the SCR. Commercially available SCRs have forward current ratings from about 30 A to 100 A.

- (v) **Circuit fusing (I^2t) rating.** It is the product of square of forward surge current and the time of duration of the surge i.e.,

$$\text{Circuit fusing rating} = I^2t$$

The circuit fusing rating indicates the maximum forward surge current capability of SCR. For example, consider an SCR having circuit fusing rating of $90 \text{ A}^2\text{s}$. If this rating is exceeded in the SCR circuit, the device will be destroyed by excessive power dissipation.

Example 5.1. An SCR has a breakover voltage of 400 V, a trigger current of 10 mA and holding current of 10 mA. What do you infer from it? What will happen if gate current is made 15 mA?

Solution. (i) **Breakover voltage of 400 V.** It means that if gate is open and the supply voltage is

400 V, then SCR will start conducting heavily. However, as long as the supply voltage is less than 400V, the SCR stays open i.e. it does not conduct.

(ii) **Trigger current of 10mA.** It means that if the supply voltage is less than breakover voltage (i.e. 400 V) and a minimum gate current of 10mA is passed, the SCR will close i.e. starts conducting heavily. The SCR will not conduct if the gate current is less than 10mA. It may be emphasised that triggering is the normal way to close an SCR as the supply voltage is normally much less than the breakover voltage.

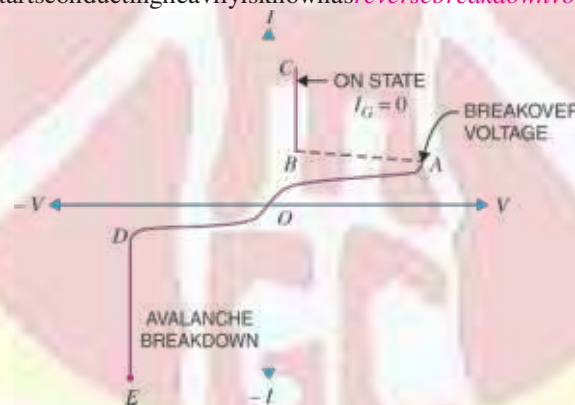
(iii) **Holding current of 10mA.** When the SCR is conducting, it will not open (i.e. stop conducting) even if triggering current is removed. However, if supply voltage is reduced, the anode current also decreases. When the anode current drops to 10mA, the holding current, the SCR is turned off. If gate current is increased to 15mA, the SCR will be turned on at a lower supply voltage.

5.1 V-I Characteristics of SCR

It is the curve between anode-cathode voltage (V) and anode current (I) of an SCR at constant gate current. Fig. 5.7 shows the V-I characteristics of a typical SCR.

(i) **Forward characteristics.** When anode is positive w.r.t. cathode, the curve between V and I is called the forward characteristic. In Fig. 5.7, $OABC$ is the forward characteristic of SCR at $I_G = 0$. If the supply voltage is increased from zero, a point is reached (point A) when the SCR starts conducting. Under this condition, the voltage across SCR suddenly drops as shown by dotted curve AB and most of supply voltage appears across the load resistance R_L . If proper gate current is made to flow, SCR can close at much smaller supply voltage.

(ii) **Reverse characteristics.** When anode is negative w.r.t. cathode, the curve between V and I is known as **reverse characteristic**. The reverse voltage does come across SCR when it is operated with a.c. supply. If the reverse voltage is gradually increased, at first the anode current remains small (i.e. leakage current) and at some reverse voltage, avalanche breakdown occurs and the SCR starts conducting heavily in the reverse direction as shown by the curve DE . This maximum reverse voltage at which SCR starts conducting heavily is known as **reverse breakdown voltage**.



SCR in Normal Operation

In order to operate the SCR in normal operation, the following points are kept in view:

- (iv) The supply voltage is generally much less than breakover voltage.
- (v) The SCR is turned on by passing an appropriate amount of gate current (a few mA) and not by breakover voltage.
- (vi) When SCR is operated from a.c. supply, the peak reverse voltage which comes during negative half-cycles should not exceed the reverse breakdown voltage.
- (vii) When SCR is to be turned OFF from the ON state, anode current should be reduced to holding current.
- (viii) If gate current is increased above the required value, the SCR will close at much reduced supply voltage.

5.2 The Triac

The major drawback of an SCR is that it can conduct current in one direction only. Therefore, an SCR can only

control d.c. power or forward biased half-cycles of a.c. in a load. However, in an a.c. system, it is often desirable and necessary to exercise control over both positive and negative half-cycles. For this purpose, a semiconductor device called *triac* is used.

A triac is a three-terminal semiconductor switching device which can control alternating current in a load.

Triac is an abbreviation for *triode a.c. switch*. 'Tri' – indicates that the device has three terminals and 'ac' means that the device controls alternating current or can conduct current in either direction.

The key function of a triac may be understood by referring to the simplified Fig. 5.8. The control circuit of triac can be adjusted to pass the desired portion of positive and negative half-cycle of a.c. supply through the load R_L . Thus referring to Fig. 5.8(ii), the triac passes the positive



Fig.5.8

half-cycle of the supply from θ_1 to 180° i.e. the shaded portion of positive half-cycle. Similarly, the shaded portion of negative half-cycle will pass through the load. In this way, the alternating current and hence a.c. power flowing through the load can be controlled.

Since a triac can control conduction of both positive and negative half-cycles of a.c. supply, it is sometimes called a bidirectional semiconductor triode switch. The above action of a triac is as follows –

- * For example, the speed of a ceiling fan can be changed in four to five steps by electrical method.
- ** Although it appears that 'triac' has two terminals, there is also a third terminal connected to the control circuit.

tainly not a rectifying action (as in an *SCR*) so that the triac makes no mention of rectification in its name.

Triac Construction

A **triac** is a three-terminal, five-layer semiconductor device whose forward and reverse characteristics are identical to the forward characteristics of the *SCR*. The three terminals are designated as main terminal *MT1*, main terminal *MT2* and gate *G*.

Fig. 5.9 (i) shows the basic structure of a triac. As we shall see, a triac is equivalent to two separate *SCRs* connected in inverse parallel (i.e. anode of each connected to the cathode of the other) with gates commoned as shown in Fig. 5.9(ii). Therefore, a triac acts like a bidirectional switch i.e. it can conduct current in either direction. This is unlike an *SCR* which can conduct current only in one direction. Fig. 5.9 (iii) shows the schematic symbol of a triac. The symbol consists of two parallel diodes connected in opposite directions with a single gate lead. It can be seen that even the symbol of a triac indicates that it can conduct current for either polarity of the main terminals (*MT1* and *MT2*) i.e. it can act as a bidirectional switch. The gate provides control over conduction in either direction.

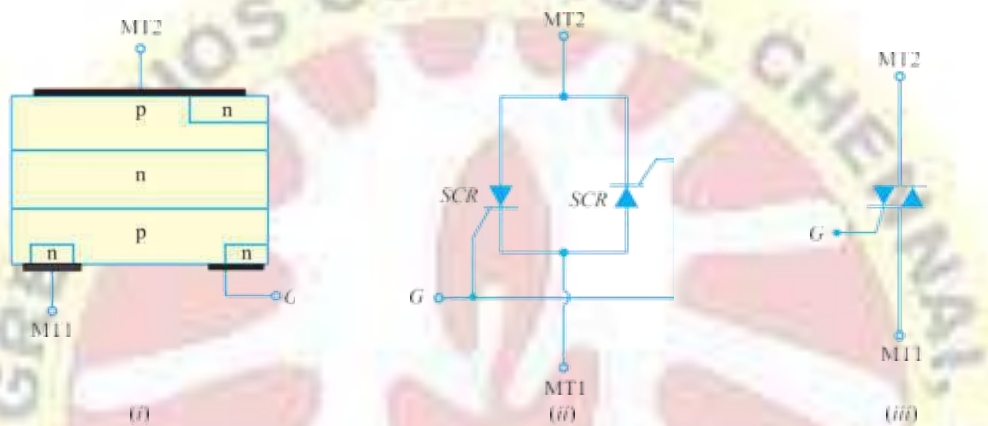


Fig.5.9

The following points may be noted about the triac:

(xiii) The triac can conduct current (of course with proper gate current) regardless of the polarities of the main terminals *MT1* and *MT2*. Since there is no longer a specific anode or cathode, the main leads are referred to as *MT1* and *MT2*.

(xiv) A triac can be turned on either with a positive or negative voltage at the gate of the device.

(xv) Like the *SCR*, once the triac is fired into conduction, the gate loses all control. The triac can be turned off by reducing the circuit current to the value of holding current.

(xvi) The main disadvantage of triacs over *SCRs* is that triacs have considerably lower current-handling capabilities. Most triacs are available in ratings of less than 40 A at voltages up to 600 V.

SCR Equivalent Circuit of Triac

We shall now see that a triac is equivalent to two separate *SCRs* connected in inverse parallel (i.e. anode of each connected to the cathode of the other) with gates commoned. Fig. 5.10 (i) shows the basic structure of a triac. If we split the basic structure of a triac into two halves as shown in Fig. 5.10(ii), it is easy to see that we have two *SCRs* connected in inverse parallel. The left half in Fig. 5.10(ii) consists of a *pnpn* device ($p_1n_2p_2n_4$) having three *pn* junctions and constitutes *SCR1*. Similarly, the

*

SCR is a controlled rectifier. It is a unidirectional switch and can conduct only in one direction. Therefore, it can

right half in Fig. 5.3(ii) consists of pnp device ($p_2n_3p_1n_1$) having three pn junctions and constitutes SCR_2 . The SCR equivalent circuit of the triac is shown in Fig. 5.4.

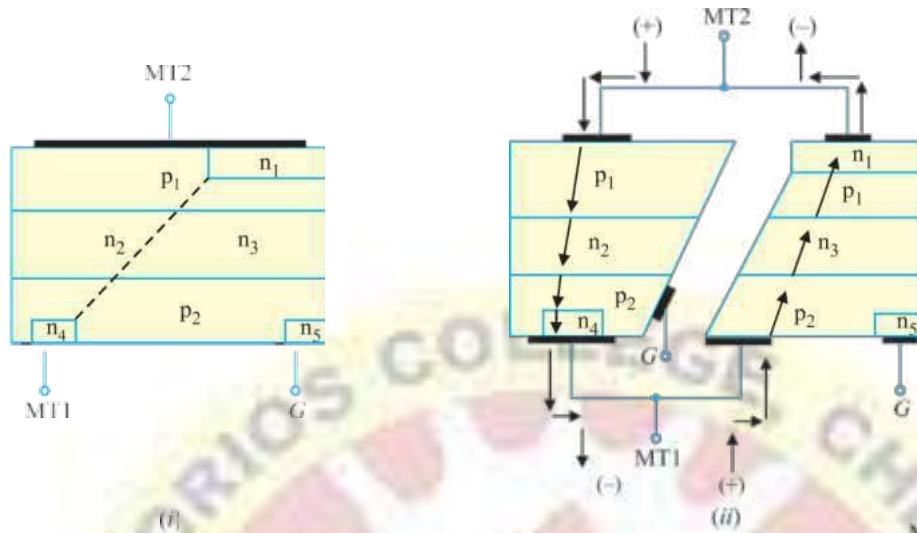


Fig.5.10

Suppose the main terminal MT_2 is positive and main terminal MT_1 is negative. If the triac is now fired into conduction by proper gate current, the triac will conduct current following the path (left half) shown in Fig. 5.11 (ii). In relation to Fig. 5.11, the SCR_1 is *ON* and the SCR_2 is *OFF*. Now suppose that MT_2 is negative and MT_1 is positive. With proper gate current, the triac will be fired into conduction. The current through the devices follows the path (right half) as shown in Fig. 5.10 (ii). In relation to Fig. 5.11, the SCR_2 is *ON* and the SCR_1 is *OFF*. Note that the triac will conduct current in the appropriate direction as long as the current through the device is greater than its holding current.

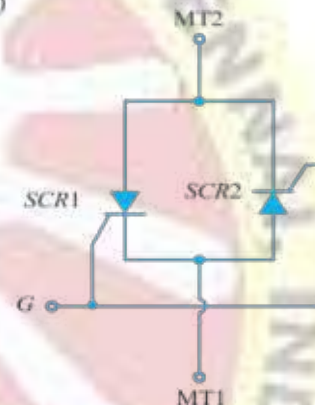


Fig.5.11

Triac Operation

Fig. 5.12 shows the simple triac circuit. The a.c. supply to be controlled is connected across the main terminal of triac through a load resistance R_L . The gate circuit consists of battery, a current limiting resistor R and a switch S . The circuit action is as follows:

(xvii) With switch S open, there will be no gate current and the triac is cut off. Even with no gate current, the triac can be turned on provided the supply voltage becomes equal to the breakover voltage of triac. However, the normal way to turn on a triac is by introducing a proper gate current.

(xviii) When switch S is closed, the gate current starts flowing in the gate circuit. In a similar manner to SCR , the breakover voltage of the triac can be varied by making proper gate current to flow. With a few milliamperes introduced at the gate, the triac will

start conducting whether terminal MT_2 is positive or negative w.r.t. MT_1 .

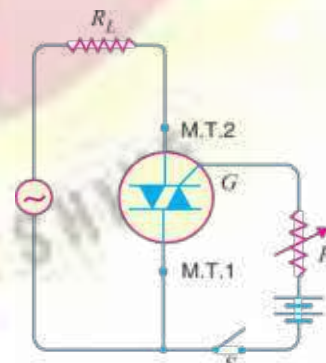


Fig.5.12

(xix) If terminal $MT2$ is positive w.r.t. $MT1$, the triac turns on and the conventional current will flow from $MT2$ to $MT1$. If the terminal $MT2$ is negative w.r.t. $MT1$, the triac is again turned on but this time the conventional current flows from $MT1$ to $MT2$.

The above action of triac reveals that it can act as a *a.c.* contactor to switch on or off alternating current to a load. The additional advantage of triac is that by adjusting the gate current to a proper value, any portion of both positive and negative half-cycles of *a.c.* supply can be made to flow through the load. This permits to adjust the transfer of *a.c.* power from

Example 5.2. Draw the transistor equivalent circuit of a triac and explain its operation from this equivalent circuit.

the source to the load.

Solution. We have seen that transistor equivalent circuit of an SCR is composed of *pnp* transistor and *npn* transistor with collector of each transistor coupled to the base of the other. Since a triac is equivalent to two SCRs connected in inverse parallel (Refer back to Fig. 5.11), the transistor equivalent circuit of triac will be composed of four transistors arranged as shown in Fig. 5.13 (i). The transistors Q_1 and Q_2 constitute the equivalent circuit of SCR1 while the transistors Q_3 and Q_4 constitute the equivalent circuit of SCR2. We can explain the action of triac from its transistor equivalent circuit as under:

(i) When $MT2$ is positive w.r.t. $MT1$ and appropriate gate current is allowed in the gate circuit, SCR1 is turned ON while SCR2 remains OFF. In terms of transistor equivalent circuit, Q_1 and Q_2 are forward biased while Q_3 and Q_4 are reverse biased. Therefore, transistors Q_1 and Q_2 conduct current as shown in Fig. 5.13 (i). Since Q_1 and Q_2 form a positive feedback, both transistors are quickly driven to saturation and a large current flows through the load R_L . This is as if switch between $MT2$ and $MT1$ were closed.

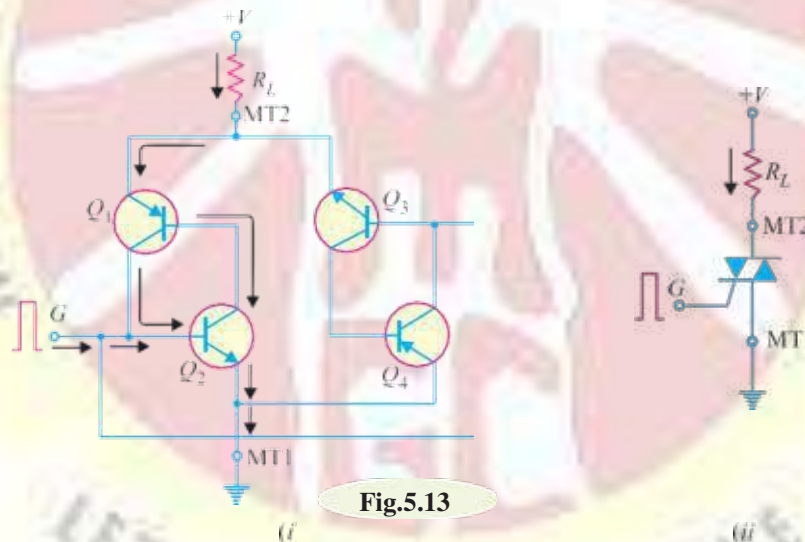


Fig. 5.13 (ii) shows the action of triac ($MT2$ positive w.r.t. $MT1$) by replacing the triac with its symbol.

(ii) When $MT2$ is negative w.r.t. $MT1$ and appropriate gate current is allowed in the gate circuit, SCR2 is turned ON and SCR1 is OFF. In terms of transistor equivalent circuit, Q_3 and Q_4 are forward biased while Q_1 and Q_2 are reverse biased. Therefore, transistors Q_3 and Q_4 will conduct as shown in Fig. 5.13 (i). As explained above, the current in load R_L will quickly attain a large value. The circuit will behave as if a switch is closed between $MT2$ and $MT1$.

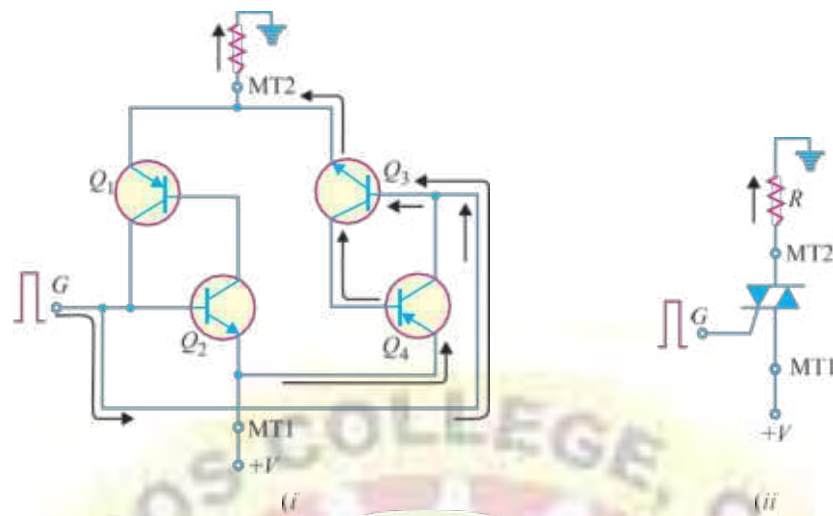


Fig.5.14

Fig.5.14(ii) shows the action of triac (MT2 negative w.r.t. MT1) by replacing the triac with its symbol.

Traic Characteristics

Fig. 5.15 shows the $V-I$ characteristics of a triac. Because the triac essentially consists of two SCRs of opposite orientation fabricated in the same crystal, its operating characteristics in the first and third quadrants are the same except for the direction of applied voltage and current flow. The following points may be noted from the triac characteristics:

(xx) The $V-I$ characteristics for triac in the 1st and 3rd quadrants are essentially identical to those of an SCR in the 1st quadrant.

(xxi) The triac can be operated with either positive or negative gate control voltage but in normal operation usually the gate voltage is positive in quadrant I and negative in quadrant III.

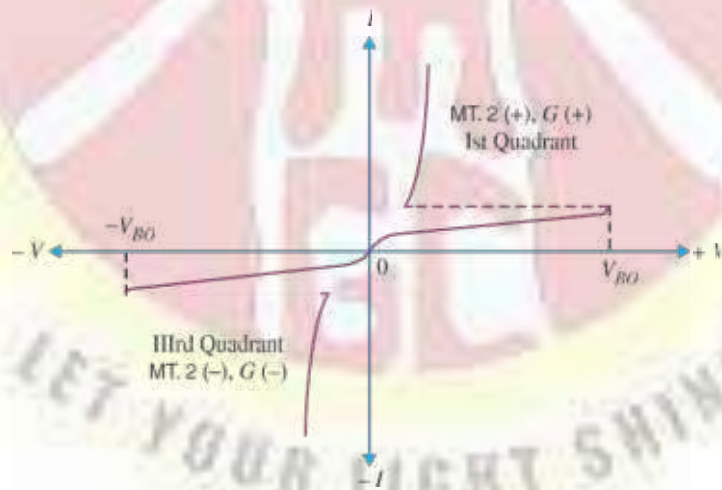
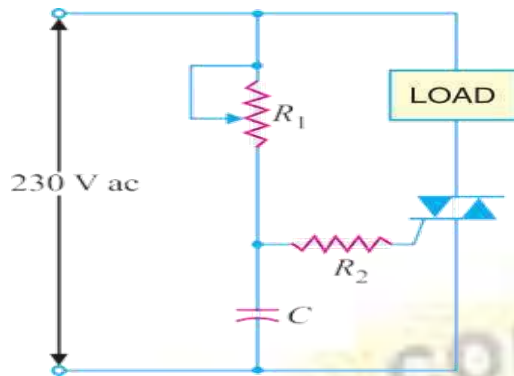


Fig.5.15

* With this arrangement, less charge is required to turn on the triac.

The supply voltage at which the triac is turned ON depends upon the gate current. The greater the gate current, the smaller the supply voltage at which the triac is turned on. This permits to use a triac to control a.c. power in a load from zero to full power in a smooth and continuous manner with no loss in the controlling device.



The Diac

A **diac** is a two-terminal, three-layer bidirectional device which can be switched from its OFF state to ON state for either polarity of applied voltage.



Fig.5.17

The diac can be constructed in either *npn* or *pnp* form. Fig.5.17(i) shows the basic structure of a diac in *pn* form. The two leads are connected to *p*-regions of silicon separated by an *n*-region. The structure of diac is very much similar to that of a transistor. However, there are several important differences:

- (i) There is no terminal attached to the base layer.
- (ii) The three regions are nearly identical in size.
- (iii) The doping concentrations are identical (unlike a bipolar transistor) to give the device symmetrical properties.

Fig.5.17(ii) shows the symbol of a diac.

Operation. When a positive or negative voltage is applied across the terminals of a diac, only a small leakage current I_{BO} will flow through the device. As the applied voltage is increased, the leakage current will continue to flow until the voltage reaches the breakover voltage V_{BO} . At this point, avalanche breakdown of the reverse-biased junction occurs and the device exhibits negative resistance, i.e. current through the device increases with the decreasing values of applied voltage. The voltage across the device then drops to 'breakback' voltage V_W .

Fig. 5.18 shows the V - I characteristics of a diac. For applied positive voltage less than $+V_{BO}$ and negative voltage less than $-V_{BO}$, a small leakage current ($\pm I_{BO}$) flows through the device. Under such conditions, the diac blocks the flow of current and effectively behaves as an open circuit. The voltages $+V_{BO}$ and $-V_{BO}$ are the breakdown voltages and usually have a range of 30 to 50 volts.

When the positive or negative applied voltage is equal to or greater than the breakdown voltage, the diac begins to conduct and the voltage drop across it becomes a few volts. Conduction then continues until the device current drops below its holding current. Note that the breakover voltage and holding current values are identical for the forward and reverse regions of operation.

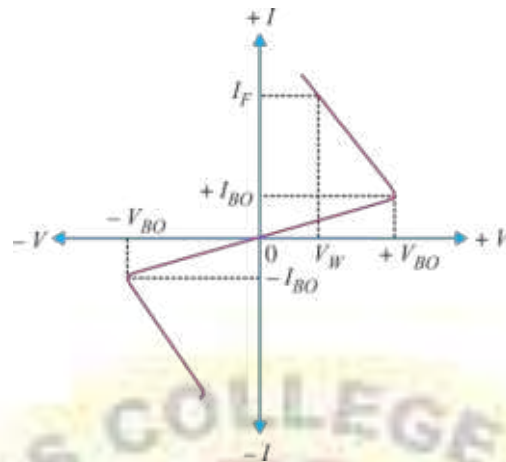


Fig.5.18

Diacs are used primarily for triggering of triacs in adjustable phase control of a.c. mains power. Some of the circuit applications of diac are (i) light dimming (ii) heat control and (iii) universal motor speed control.

Applications of Diac

Although a triac may be fired into the conducting state by a simple resistive triggering circuit, more reliable and faster turn-on may be achieved if a switching device is used in series with the gate. One of the switching devices that can trigger a triac is the diac. This is illustrated in the following applications.

(iv) Lamp dimmer. Fig. 5.19 shows a typical circuit that may be used for smooth control of a.c. power fed to a lamp. This permits control of the light output from the lamp. The basic control is by an RC variable gate voltage arrangement. The series R_4 – C_1 circuit across the triac is designed to limit the rate of voltage rise across the device during switch off.

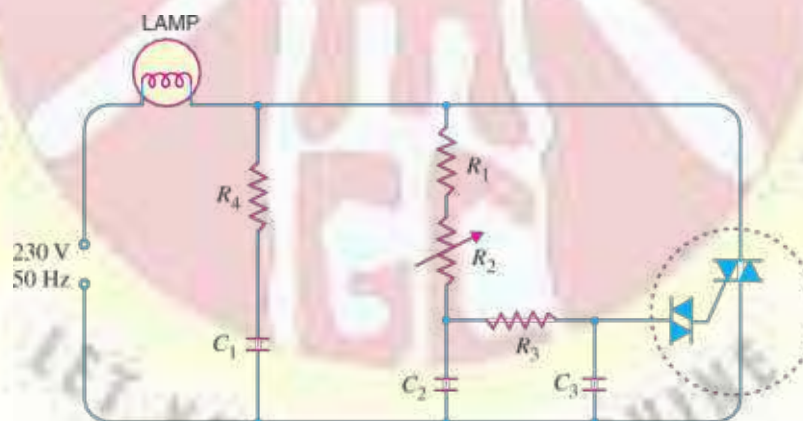


Fig.5.19

The circuit action is as follows: As the input voltage increases positively or negatively, C_1 and C_2 charge at a rate determined primarily by R_2 . When the voltage across C_3 exceeds the breakover voltage of the diac, the diac is fired into the conducting state. The capacitor C_3 discharges through the conducting diac into the gate of the triac. Hence, the triac is turned on to pass the a.c. power to the

lamp. By adjusting the value of R_2 , the rate of charge of capacitors and hence the point at which triac will trigger on the positive or negative half-cycle of input voltage can be controlled. Fig. 5.20 shows the waveforms of supply voltage and load voltage in a diac-triac control circuit.



Fig. 5.20

The firing of triac can be controlled up to a maximum of 180° . In this way, we can provide a continuous control of load voltage from practically zero to full r.m.s. value.

(v) Heat control. Fig. 5.21 shows a typical diac-triac circuit that may be used for the smooth control of a.c. power in a heater. This is similar to the circuit shown in Fig. 5.4. The capacitor C_1 in series with choke L across the triac helps to slow-up the voltage rise across the device during switch-off. The resistor R_4 in parallel with the diac ensures smooth control at all positions of variable resistance R_2 .

The circuit action is as follows: As the input voltage increases positively or negatively, C_1 and C_2 charge at a rate determined primarily by R_2 . When the voltage across C_3 exceeds the breakover voltage of the diac, the diac conducts. The capacitor C_3 discharges through the conducting diac into the gate of the triac. This turns on the triac and hence a.c. power to the heater. By adjusting the value of R_2 , any portion of positive and negative half-cycles of the supply voltage can be passed through the heater. This permits a smooth control of the heat output from the heater.

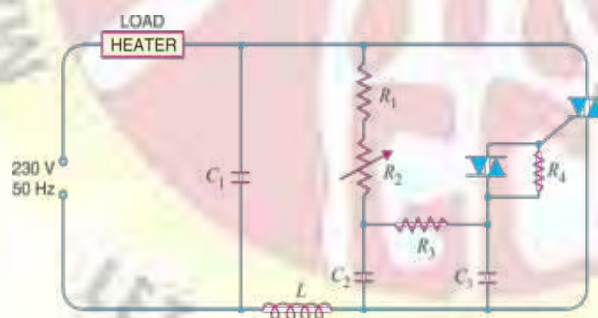


Fig. 5.21

Example 53. We require a small gate triggering voltage V_{GT} (say $\pm 2V$) to fire a triac into conduction. How will you raise the trigger level of the triac?

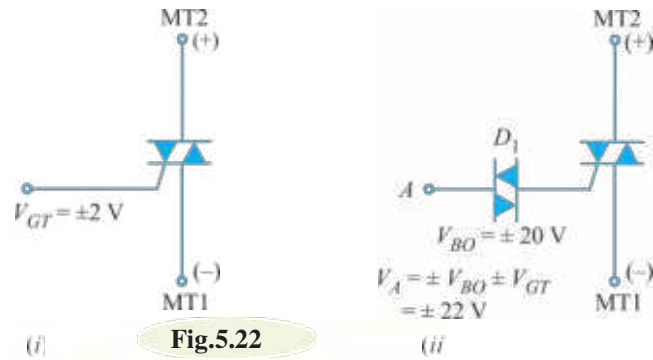


Fig. 5.22

$$V_A = \pm V_{BO} \pm V_{GT}$$

$$= \pm 20V \pm 2V = \pm 22V$$

Therefore, in order to turn on the triac, the gate trigger signal V_A must be $\pm 22V$. Note that triac triggering control is the primary function of a diac.



Fig. 5.23
