

Course File

Name of the subject : Electronic devices and circuits
Name of the Course Coordinator & Faculty Member : ch.krishnaveni
Academic Year : 2026-27
Year & Sem : II & I Sem



Department of Electronics & Communication Engineering

TKR College of Engineering & Technology
Meerpet, Balapur, Saroornagar, Hyderabad - 500097, Telangana.
Accredited by NBA & NAAC A+

Course Objectives

- To introduce components such as diodes, BJTs and FETs.
- To know the applications of components.
- To know the switching characteristics of components
- To give understanding of various types of amplifier circuits

Course Outcomes

Upon completion of the Course, the students will be able to:

- Know the characteristics of various components.
- Understand the utilization of components.
- Understand the biasing techniques
- Design and analyze small signal amplifier circuits.

JNTU Syllabus

Unit – I	Diode and Applications: Diode - Static and Dynamic resistances, Equivalent circuit, Load line analysis, Diffusion and Transition Capacitances, Diode Applications: Switch-Switching times. Rectifier - Half Wave Rectifier, Full Wave Rectifier, Bridge Rectifier, Rectifiers with Capacitive and Inductive Filters, Clippers-Clipping at two independent levels, Clamper-Clamping Circuit Theorem, Clamping Operation, Types of Clampers.
Unit – II	Bipolar Junction Transistor (BJT): Principle of Operation, Common Emitter, Common Base and Common Collector Configurations, Transistor as a switch, switching times, Transistor Biasing and Stabilization - Operating point, DC & AC load lines, Biasing - Fixed Bias, Self-Bias, Bias Stability, Bias Compensation using Diodes.
Unit – III	Junction Field Effect Transistor (FET): Construction, Principle of Operation, Pinch-Off Voltage, Volt- Ampere Characteristic, Comparison of BJT and FET, Biasing of FET, FET as Voltage Variable Resistor. Special Purpose Devices: Zener Diode - Characteristics, Voltage Regulator. Principle of Operation -SCR, Tunnel diode, UJT, Varactor Diode.
Unit – IV	Analysis and Design of Small Signal Low Frequency BJT Amplifiers: Transistor Hybrid model, Determination of h-parameters from transistor characteristics, Typical values of h-parameters in CE, CB and CC configurations, Transistor amplifying action, Analysis of CE, CC, CB Amplifiers and CE Amplifier with emitter resistance, low frequency response of BJT Amplifiers, effect of coupling and bypass capacitors on CE Amplifier.
Unit – V	FET Amplifiers: Small Signal Model, Analysis of JFET Amplifiers, Analysis of CS, CD, CG JFET Amplifiers. MOSFET Characteristics in Enhancement and Depletion mode, Basic Concepts of MOS Amplifiers.

Books / Material

Text Books	
a	Electronic Devices and Circuits- Jacob Millman, McGraw Hill Education
b	Electronic Devices and Circuits theory– Robert L. Boylestead, Louis Nashelsky, 11th Edition, 2009, Pearson.

Suggested / Reference Books	
a	The Art of Electronics, Horowitz, 3rd Edition Cambridge University Press
b	Electronic Devices and Circuits, David A. Bell – 5th Edition, Oxford.
c	Pulse, Digital and Switching Waveforms –J. Millman, H. Taub and Mothiki S. Prakash Rao, 2Ed., 2008, Mc Graw Hill.

TEACHING NOTES

UNIT I PN DIODE CHARACTERISTICS

Introduction:

Electronics Engineering is a branch of engineering which deals with the flow of electrons in vacuum tubes, gas and semiconductor.

Applications of Electronics:

Home Appliances, Medical Applications, Robotics, Mobile Communication, Computer Communication etc.

Atomic Structure:

- Atom is the basic building block of all the elements. It consists of the central nucleus of positive charge around which small negatively charged particles called electrons revolve in different paths or orbits.
- An Electrostatic force of attraction between electrons and the nucleus holds up electrons in different orbits.

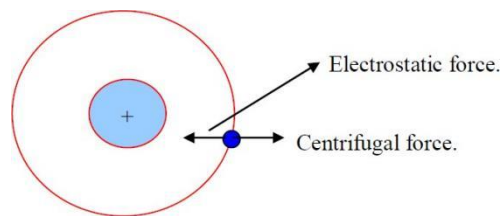


Figure 1.1: Atomic structure

- Nucleus is the central part of an atom and contains protons and neutrons. A proton is positively charged particle, while the neutron has the same mass as the proton, but has no charge. Therefore, nucleus of an atom is positively charged.
- **atomic weight = no. of protons + no. of neutrons**
- An electron is a negatively charged particle having negligible mass. The charge on an electron is equal but opposite to that on a proton. Also the number of electrons is equal to the number of protons in an atom under ordinary conditions. Therefore an atom is neutral as a whole.
- **atomic number = no. of protons or electrons in an atom**
- The number of electrons in any orbit is given by $2n^2$ where n is the number of the orbit.

For example, I orbit contains $2 \times 1^2 = 2$ electrons

II orbit contains $2 \times 2^2 = 8$ electrons

III orbit contains $2 \times 3^2 = 18$ electrons and so on

- The last orbit cannot have more than 8 electrons.
- The last but one orbit cannot have more than 18 electrons.

Positive and negative ions:

- Protons and electrons are equal in number hence if an atom loses an electron it has lost negative charge therefore it becomes positively charged and is referred as positive ion.
- If an atom gains an electron it becomes negatively charged and is referred to as negative ion.

Valence electrons:

The electrons in the outermost orbit of an atom are known as valence electrons.

- The outermost orbit can have a maximum of 8 electrons.
- The valence electrons determine the physical and chemical properties of a material.
- When the number of valence electrons of an atom is less than 4, the material is usually a metal and a conductor. Examples are sodium, magnesium and aluminium, which have 1, 2 and 3 valence electrons respectively.
- When the number of valence electrons of an atom is more than 4, the material is usually a non-metal and an insulator. Examples are nitrogen, sulphur and neon, which have 5, 6 and 8 valence electrons respectively.
- When the number of valence electrons of an atom is 4 the material has both metal and non-metal properties and is usually a semi-conductor. Examples are carbon, silicon and germanium.

Free electrons:

- The valence electrons of different material possess different energies. The greater the energy of a valence electron, the lesser it is bound to the nucleus.
- In certain substances, particularly metals, the valence electrons possess so much energy that they are very loosely attached to the nucleus.
- The loosely attached valence electrons move at random within the material and are called free electrons.

The valence electrons, which are loosely attached to the nucleus, are known as free electrons. Energy bands:

- In case of a single isolated atom an electron in any orbit has definite energy.
- When atoms are brought together as in solids, an atom is influenced by the forces from other atoms. Hence an electron in any orbit can have a range of energies rather than single energy. These range of energy levels are known as Energy bands.
- Within any material there are two distinct energy bands in which electrons may exist viz , Valence band and conduction band.

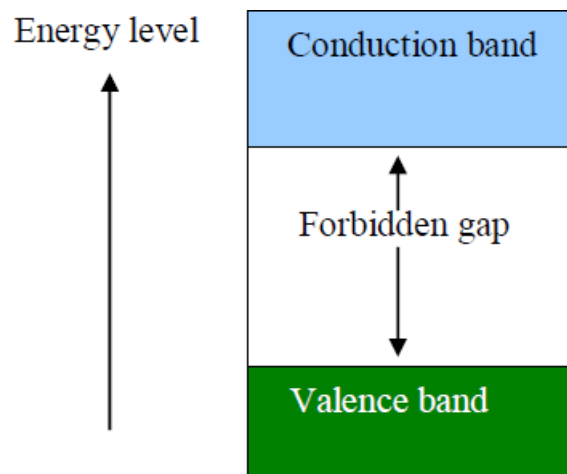


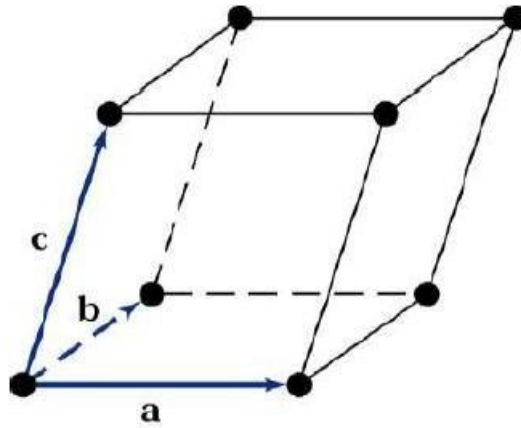
Figure 1.2: Energy level diagram

- **The range of energies possessed by valence electrons is called valence band.**
- **The range of energies possessed by free electrons is called conduction band.**
- **Valence band and conduction band are separated by an energy gap in which no electrons normally exist this gap is called forbidden gap.**

Electrons in conduction band are either escaped from their atoms (free electrons) or only weakly held to the nucleus. Thereby by the electrons in conduction band may be easily moved around within the material by applying relatively small amount of energy. (either by increasing the temperature or by focusing light on the material etc.) This is the reason why the conductivity of the material increases with increase in temperature.

But much larger amount of energy must be applied in order to extract an electron from the valence band because electrons in valence band are usually in the normal orbit around a nucleus. For any given material, the forbidden gap may be large, small or non-existent.

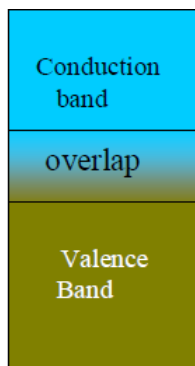
The periodic arrangement of atoms is called lattice. A unit cell of a material represents the Entire lattice. By repeating the unit cell throughout the crystal, one can generate the entire lattice.



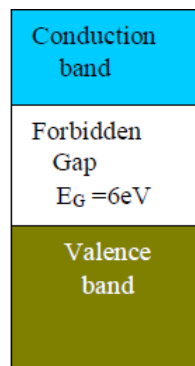
Classification of materials based on Energy band theory:

Based on the width of the forbidden gap, materials are broadly classified as

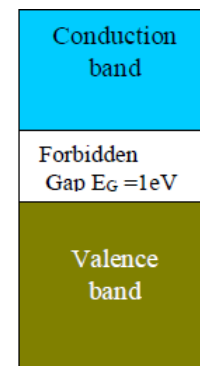
- Conductors
- Insulators
- Semiconductors.



(a) Conductor



(b) Insulator



(c) Semiconductor

Conductors:

- Conductors are those substances, which allow electric current to pass through them. Example: Copper, Al, salt solutions, etc.
- In terms of energy bands, conductors are those substances in which there is no forbidden gap. Valence and conduction band overlap as shown in fig (a).
- For this reason, very large number of electrons are available for conduction even at extremely low temperatures. Thus, conduction is possible even by a very weak electric field.

Insulators:

- Insulators are those substances, which do not allow electric current to pass through them. Example: Rubber, glass, wood etc.
- In terms of energy bands, insulators are those substances in which the forbidden gap is very large.
- Thus valence and conduction band are widely separated as shown in fig (b). Therefore insulators do not conduct electricity even with the application of a large electric field or by heating or at very high temperatures.

Semiconductors:

- Semiconductors are those substances whose conductivity lies in between that of a conductor and Insulator.
Example: Silicon, germanium, Cealenium, Gallium, arsenide etc.
- In terms of energy bands, semiconductors are those substances in which the forbidden gap is narrow.
- Thus valence and conduction bands are moderately separated as shown in fig(C).
- In semiconductors, the valence band is partially filled, the conduction band is also partially filled, and the energy gap between conduction band and valence band is narrow.
- Therefore, comparatively smaller electric field is required to push the electrons from valence band to conduction band . At low temperatures the valence band is completely filled and conduction band is completely empty. Therefore, at very low temperature a semi-conductor actually behaves as an insulator.

Conduction in solids:

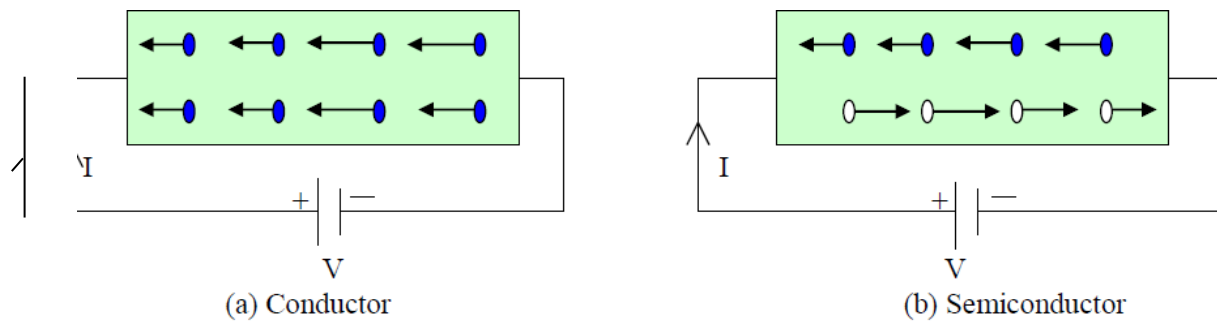
- Conduction in any given material occurs when a voltage of suitable magnitude is applied to it, which causes the charge carriers within the material to move in a desired direction.
- This may be due to electron motion or hole transfer or both.

Electron motion:

Free electrons in the conduction band are moved under the influence of the applied electric field. Since electrons have negative charge they are repelled by the negative terminal of the applied voltage and attracted towards the positive terminal.

Hole transfer:

- Hole transfer involves the movement of holes.
- Holes may be thought of positive charged particles and as such they move through an electric field in a direction opposite to that of electrons.



← Flow of electrons

→ Flow of current

← Flow of electrons

→ Flow of holes

→ Flow of current

- In a good conductor (metal) as shown in fig (a) the current flow is due to free electrons only.
- In a semiconductor as shown in fig (b). The current flow is due to both holes and electrons moving in opposite directions.
- The unit of electric current is Ampere (A) and since the flow of electric current is constituted by the movement of electrons in conduction band and holes in valence band, electrons and holes are referred as charge carriers.

Classification of semiconductors:

Semiconductors are classified into two types.

- Intrinsic semiconductors.
- Extrinsic semiconductors.

a) Intrinsic semiconductors:

- A semiconductor in an extremely pure form is known as **Intrinsic semiconductor**. Example: Silicon, germanium.
- Both silicon and Germanium are tetravalent (having 4 valence electrons).
- Each atom forms a covalent bond or electron pair bond with the electrons of neighboring atom. The structure is shown below.

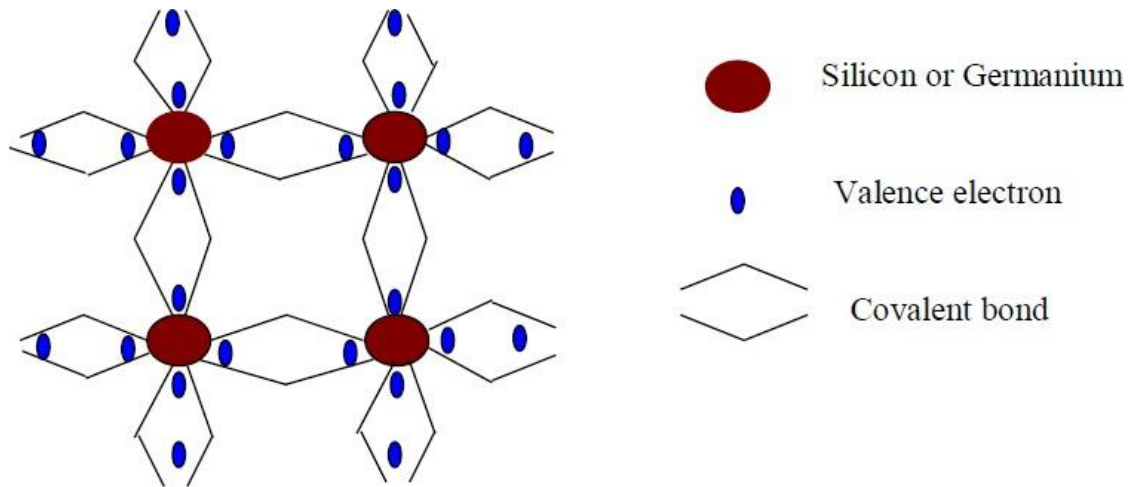


Figure 1.3: Crystalline structure of Silicon (or Germanium)

At low temperature

- At low temperature, all the valence electrons are tightly bounded the nucleus hence no free electrons are available for conduction.
- The semiconductor therefore behaves as an Insulator at absolute zero temperature.

At room temperature

- At room temperature, some of the valence electrons gain enough thermal energy to break up the covalent bonds.
- This breaking up of covalent bonds sets the electrons free and is available for conduction.
- When an electron escapes from a covalent bond and becomes free electrons a vacancy is created in a covalent bond as shown in figure above. Such a vacancy is called Hole. It carries positive charge and moves under the influence of an electric field in the direction of the electric field applied.
- Numbers of holes are equal to the number of electrons since; a hole is nothing but an absence of electrons.

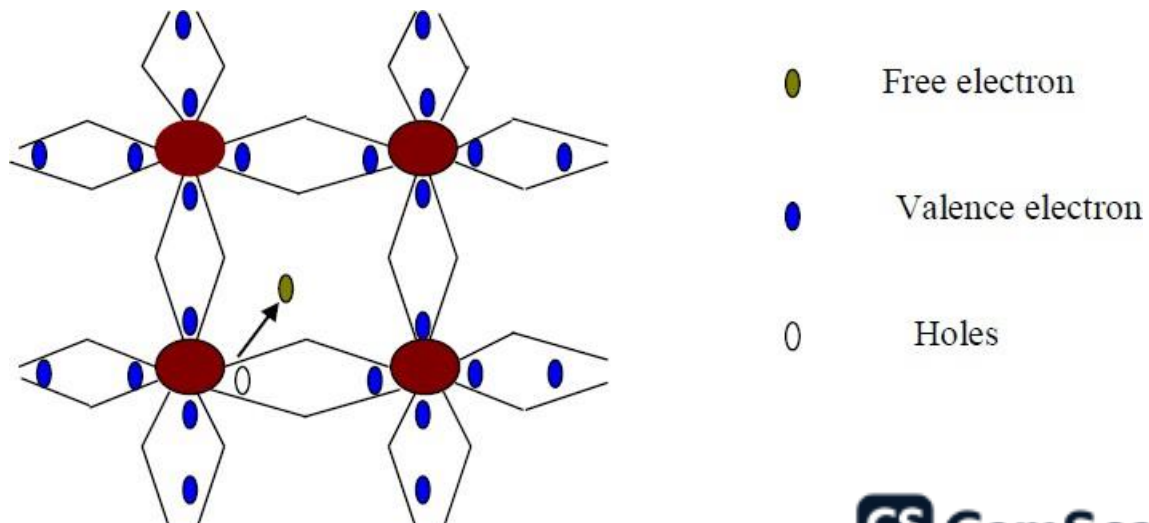


Figure 1.4: Crystalline structure of Silicon (or Germanium) at

Extrinsic Semiconductor:

- When an impurity is added to an intrinsic semiconductor its conductivity changes.
- This process of adding impurity to a semiconductor is called Doping and the impure semiconductor is called extrinsic semiconductor.
- Depending on the type of impurity added, extrinsic semiconductors are further classified as n-type and p-type semiconductor.

N-type semiconductor:

- **When a small current of Pentavalent impurity is added to a pure semiconductor it is called as n-type semiconductor.**
- Addition of Pentavalent impurity provides a large number of free electrons in a semiconductor crystal.
- Typical example for pentavalent impurities are Arsenic, Antimony and Phosphorus etc. Such impurities which produce n-type semiconductors are known as Donor impurities because they donate or provide free electrons to the semiconductor crystal.
- To understand the formation of n-type semiconductor, consider a pure silicon crystal with an impurity say arsenic added to it as shown in figure 1.5.
- We know that a silicon atom has 4 valence electrons and Arsenic has 5 valence electrons. When Arsenic is added as impurity to silicon, the 4 valence electrons of silicon make co-valent bond with 4 valence electrons of Arsenic.
- The 5th Valence electrons finds no place in the covalent bond thus, it becomes free and travels to the conduction band as shown in figure. Therefore, for each arsenic atom added, one free electron will be available in the silicon crystal. Though each arsenic atom provides one free electrons yet an extremely small amount of arsenic impurity provides enough atoms to supply millions of free electrons.

Due to thermal energy, still hole electron pairs are generated but the number of free electrons are very large in number when compared to holes. So in an n-type semiconductor electrons are majority charge carriers and holes are minority charge carriers. Since the current conduction is pre-dominantly by free electrons(-vely charges) it is called as n-type semiconductor(n- means -ve).

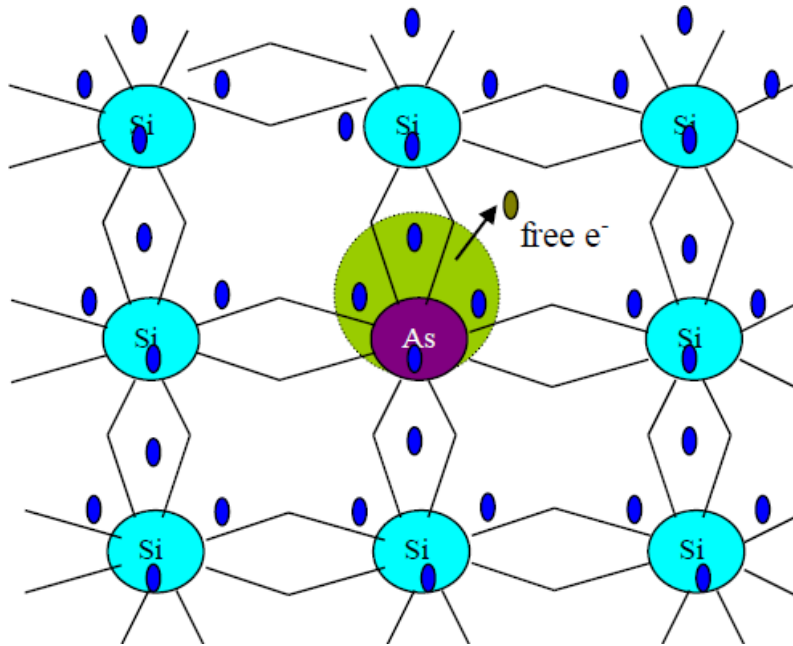


Figure 1.5: N-type semiconductors

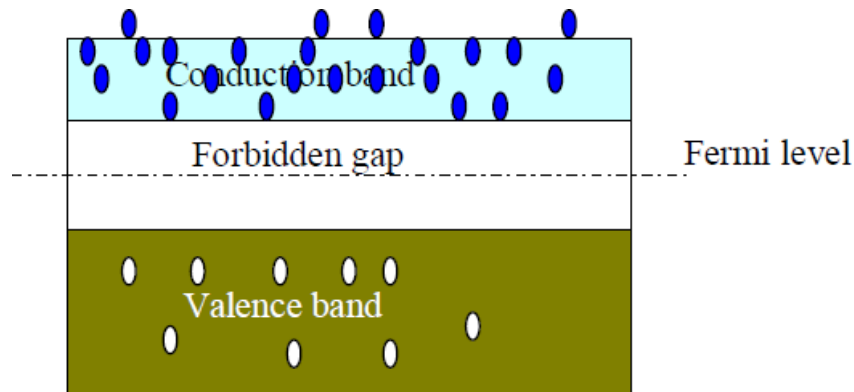


Figure 1.6: Energy band diagram for n-type semiconductor

P-type semiconductor:

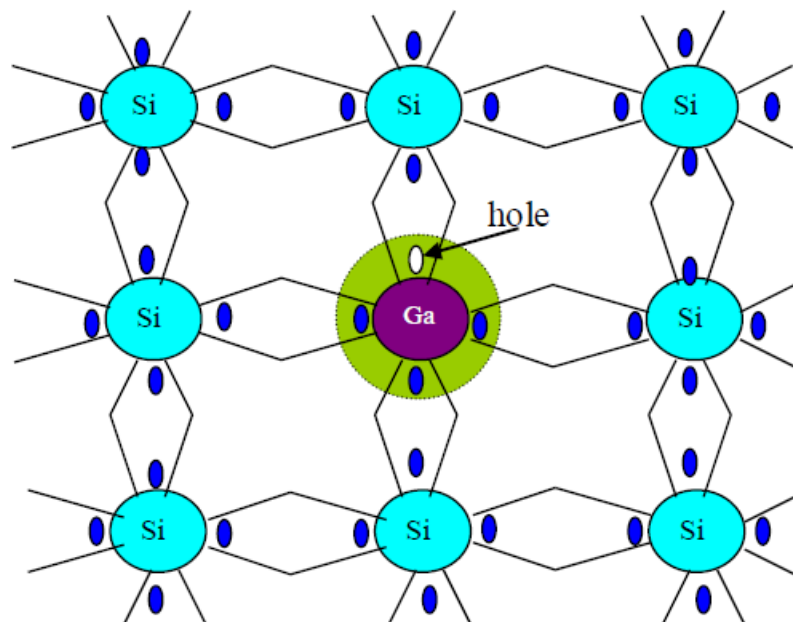


Figure 1.7: P-type semiconductor

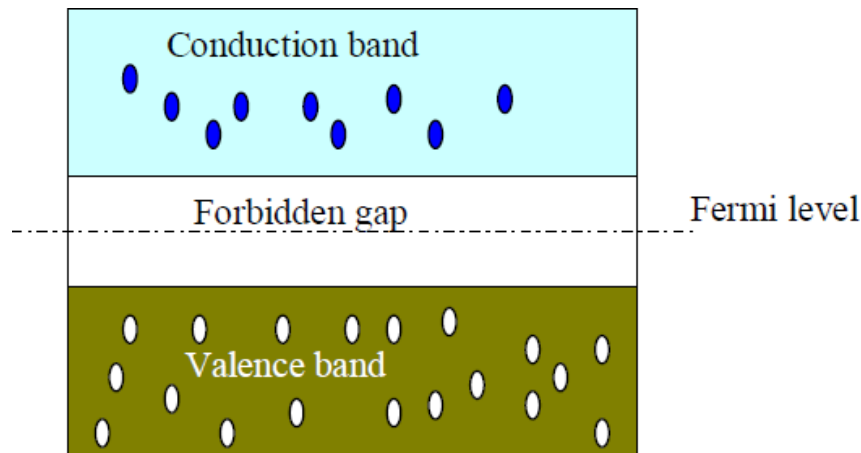


Figure 1.8: Energy band diagram for p-type semiconductor

- When a small amount of trivalent impurity is added to a pure semiconductor it is called p-type semiconductor.
- The addition of trivalent impurity provides large number of holes in the semiconductor crystals.
- Example: Gallium, Indium or Boron etc. Such impurities which produce p-type semiconductors are known as acceptor impurities because the holes created can accept the electrons in the semiconductor crystal.

To understand the formation of p-type semiconductor, consider a pure silicon crystal with an impurity say gallium added to it as shown in figure.7.

- We know that silicon atom has 4 valence electrons and Gallium has 3 electrons. When Gallium is added as impurity to silicon, the 3 valence electrons of gallium make 3 covalent bonds with 3 valence electrons of silicon.
- The 4th valence electrons of silicon cannot make a covalent bond with that of Gallium because of short of one electron as shown above. This absence of electron is called a hole. Therefore for each gallium atom added one hole is created, a small amount of Gallium provides millions of holes.

Due to thermal energy, still hole-electron pairs are generated but the number of holes is very large compared to the number of electrons. Therefore, in a p-type semiconductor holes are majority carriers and electrons are minority carriers. Since the current conduction is predominantly by hole(+ charges) it is called as p-type semiconductor (p means +ve)

PN Junction Diode:

When a p-type semiconductor material is suitably joined to n-type semiconductor the contact surface is called a p-n junction. The p-n junction is also called as semiconductor diode.

Applications of diode:

- Used as rectifier diodes in DC power suppliers
- Used as clippers and clamps
- Used as switch in logic circuit in computers
- Used as voltage multipliers.

Construction and working of a PN Junction diode: Open Circuited PN Junction:

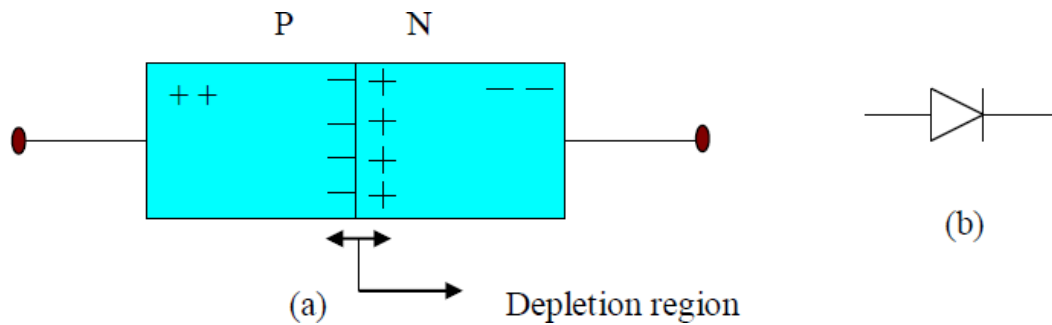


Figure 1.9 (a): p-n junction representation

Figure 1.9 (b): Symbolic representation

- The left side material is a p-type semiconductor having –ve acceptor ions and +vely charged holes. The right side material is n-type semiconductor having +ve donor ions and free electrons.
- Suppose the two pieces are suitably treated to form pn junction, then there is a tendency for the free electrons from n-type to diffuse over to the p-side and holes from p-type to the n-side . This process is called **diffusion**.
- As the free electrons move across the junction from n-type to p-type, +ve donor ions are uncovered. Hence a +ve charge is built on the n-side of the junction. At the same time, the free electrons cross the junction and uncover the –ve acceptor ions by filling in the holes. Therefore a net –ve charge is established on p-side of the junction.
- When a sufficient number of donor and acceptor ions is uncovered further diffusion is prevented.
- Thus a barrier is set up against further movement of charge carriers. This is called potential barrier or junction barrier V_0 . The potential barrier is of the order of 0.1 to 0.3V.

Note: outside this barrier on each side of the junction, the material is still neutral. Only inside the barrier, there is a +ve charge on n-side and –ve charge on p-side. This region is called depletion layer.

Biasing of a PN junction diode:

Connecting a p-n junction to an external DC voltage source is called biasing.

1. Forward biasing
2. Reverse biasing

Forward biasing

- When external voltage applied to the junction is in such a direction that it cancels the potential barrier, thus permitting current flow is called forward biasing.
- To apply forward bias, connect +ve terminal of the battery to p-type and –ve terminal to n-type as shown in fig.2.1 below.
- The applied forward potential(V_F) establishes the electric field which acts against

the field due to potential barrier. Therefore the resultant field is weakened and the barrier height is reduced at the junction as shown in fig. 2.1.

- Since the potential barrier voltage is very small, a small forward voltage (V_F) is sufficient to completely eliminate the barrier. Once the potential barrier is eliminated by the forward voltage, junction resistance (R_F) becomes almost zero and a low resistance path is established for the entire circuit. Therefore current flows in the circuit. This is called forward current (I_F).

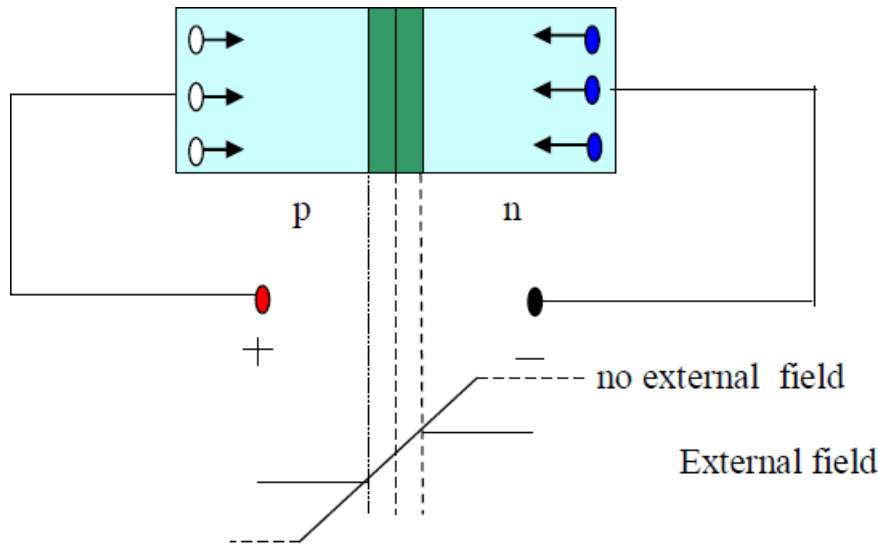


Figure 1.10: Forward biasing of p-n junction

Reverse biasing

- When the external voltage applied to the junction is in such a direction the potential barrier is increased it is called reverse biasing.

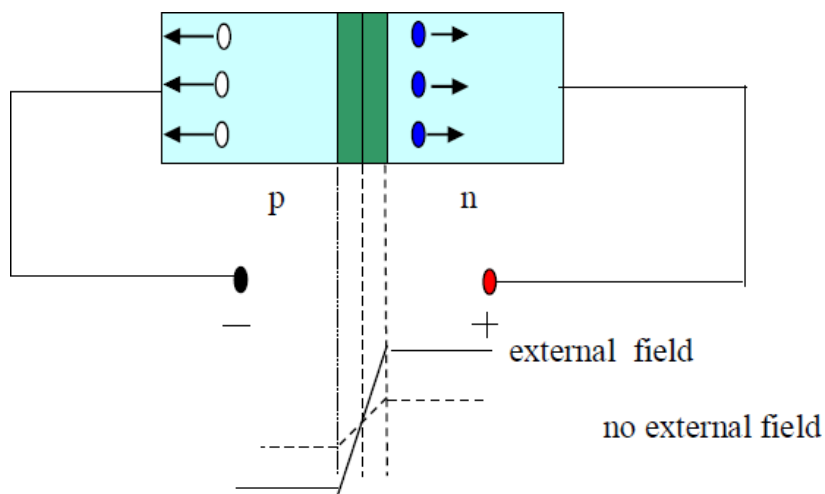


Figure 1.11: Reverse biasing of p-n junction

- To apply reverse bias, connect -ve terminal of the battery to p-type and +ve terminal to n-type as shown in figure below.

- The applied reverse voltage establishes an electric field which acts in the same direction as the field due to potential barrier. Therefore the resultant field at the junction is strengthened and the barrier height is increased as shown in fig.2.2.
- The increased potential barrier prevents the flow of charge carriers across the junction. Thus a high resistance path (R_R) is established for the entire circuit and hence current does not flow. But in practice a very little current flows due to minority charge carriers. The current is called reverse saturation current (I_S).

Volt- Ampere characteristics (V-I) of a PN junction diode:

1. **Forward Bias** - The voltage potential is connected positive, (+ve) to the P-type material and negative, (-ve) to the N-type material across the diode which has the effect of **Decreasing** the PN-junction width.

2. **Reverse Bias** - The voltage potential is connected negative, (-ve) to the P-type material and positive, (+ve) to the N-type material across the diode which has the effect of **Increasing** the PN-junction width

Forward Biased Junction Diode:

When a diode is connected in a **Forward Bias** condition, a negative voltage is applied to the N- type material and a positive voltage is applied to the P-type material.

If this external voltage (V_F) becomes greater than the value of the potential barrier (V_γ), approx.

0.7 volts for silicon and 0.3 volts for germanium, the potential barriers opposition will be overcome and current will start to flow. This current is called forward current(I_F)

This is because the negative voltage pushes or repels electrons towards the junction giving them the energy to cross over and combine with the holes being pushed in the opposite direction towards the junction by the positive voltage.

This results in a characteristics curve of zero current flowing up to this voltage point, called the "knee" on the static curves and then a high current flow through the diode with little increase in the external voltage as shown below.

Forward Characteristics:

Case 1:

When the applied forward voltage $V_F < V_\gamma$, the width of the depletion layer is increased and no current flows through the circuit.

Case 2:

When the applied forward voltage $V_F > V_\gamma$, the charge carriers move towards the junction from P to N and from N to P, due to this the width of the junction gets reduced and at one particular voltage the junction gets damaged ,and forward

resistance(R_F) offered by the diode will be very less (ideally 0) due to this there will be a flow of current due to majority charge carriers. The current is said to be forward current (I_F) which will be very large.

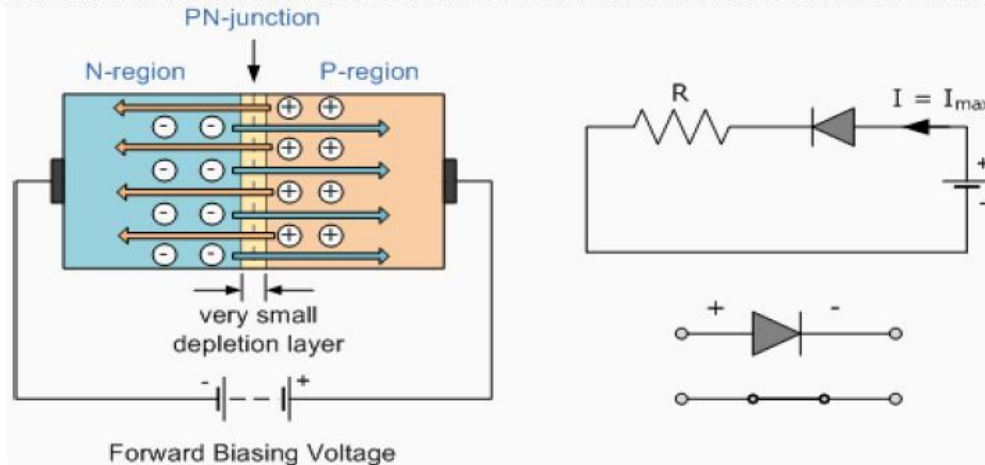


Figure 1.12: Forward Biased Junction Diode showing a Reduction in the Depletion Layer

This condition represents the low resistance path through the PN junction allowing very large currents to flow through the diode with only a small increase in bias voltage. The actual potential difference across the junction or diode is kept constant by the action of the depletion layer at approximately 0.3v for germanium and approximately 0.7v for silicon junction diodes. Since the diode can conduct "infinite" current above this knee point as it effectively becomes a short circuit, therefore resistors are used in series with the diode to limit its current flow.

Reverse Biased Junction Diode:

When a diode is connected in a **Reverse Bias** condition, a positive voltage is applied to the N-type material and a negative voltage is applied to the P-type material.

The positive voltage applied to the N-type material attracts electrons towards the positive electrode and away from the junction, while the holes in the P-type end are also attracted away from the junction towards the negative electrode.

The net result is that the depletion layer grows wider due to a lack of electrons and holes and presents a high impedance path, almost an insulator. The result is that a high potential barrier is created thus preventing current from flowing through the semiconductor material.

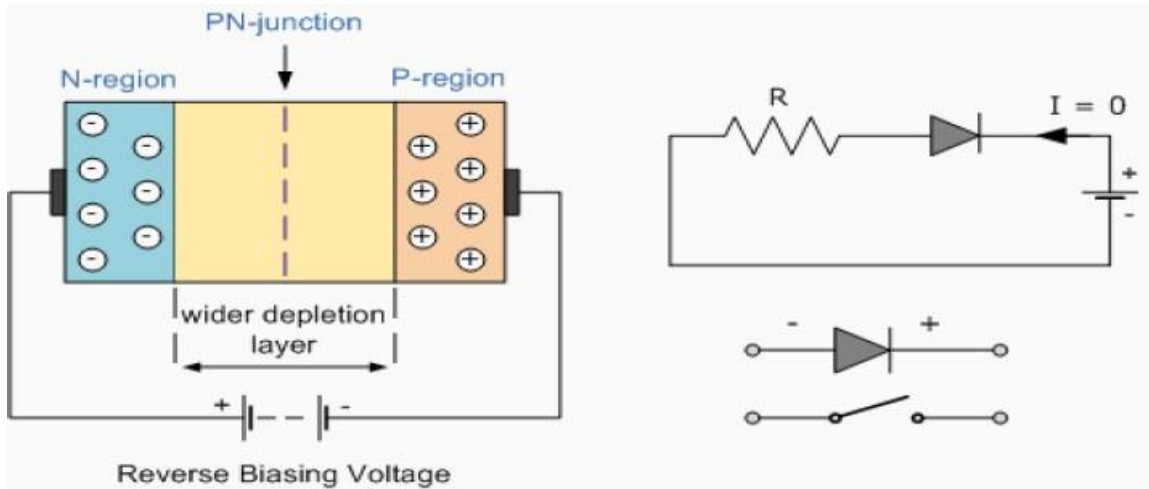


Figure 1.13: Reverse Biased Junction Diode showing an Increase in the Depletion Layer

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

Sometimes this avalanche effect has practical applications in voltage stabilising circuits where a series limiting resistor is used with the diode to limit this reverse breakdown current to a preset maximum value thereby producing a fixed voltage output across the diode. These types of diodes are commonly known as **Zener Diodes**.

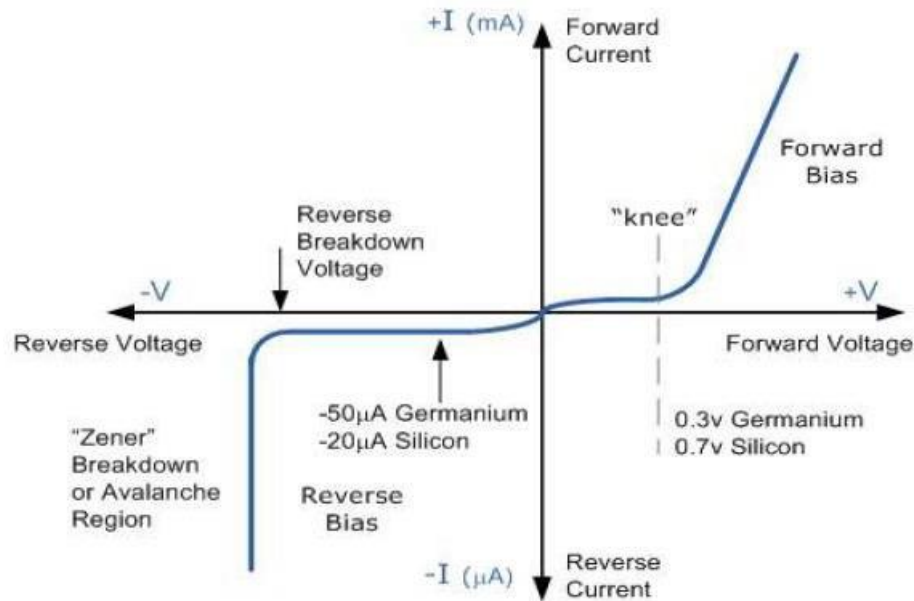


Figure 1.14: V-I Characteristics of PN Junction Diode

PN diode Currents:

The expression for the total current as a function of the applied voltage is derived below.

Here we neglect the depletion layer thickness and hence assume that the barrier width is 0. If a forward bias is applied to the diode, the holes are injected from P side to n material. The concentration of holes in the side is increased above its thermal equilibrium value P_n and given by ,

$$P_n(x) = P_{no} + P_{n(o)} e^{-x/LP} \dots\dots\dots(1)$$

Where,

L_p is called the diffusion length for holes in 'n' material and the injected or excess concentration at

$X = 0$ is

$$P_n(0) = P_n(0) - P_{n0} \dots \dots (2)$$

These several hole concentration components are indicated in figure which shows the exponential decrease of the density $P_n(x)$ with distance X into the n material.

The diffusion hole concentration components are indicated, in figure, which shows the exponential decrease the density $P_n(x)$ with distance x into the n material. The diffusion hole current in the n side is given by

$$I_{pn} = -Ae \frac{D_p d_{pn}}{dx} \dots \dots \dots (3)$$

Substituting (1) in (3) we obtain

$$I_{pn}(x) = \frac{A_e D_p P_n(0) e^{-x/L_p}}{L_p} \dots \dots \dots (4)$$

This equation verifies that the hole current decreases exponentially with distance. The dependence of I_{pn} up to applied voltage is contained implicitly in the factor $P_n(0)$ because the injected concentration is a function of voltage. We now find the dependence of $P_n(0)$ upon V .

Law of junction:

If the hole concentration at the edges of space charge region as P_n and P_p in the P and n materials, respectively and if the barrier potential across this depletion layer V_B , then

$$P_p = P_n e^{V_B/V_T} \dots \dots \dots (5)$$

This is the Boltzmann relationship of kinetic gas theory. If we apply equation (5) to the case of an open Cktd on junction, then

$$P_p = P_{p0}, P_n = P_{n0} \text{ and } V_B = V_0.$$

At the edge of depletion layer $X = 0$, $P_n = P_n(0)$.

The Boltzmann relation, i.e for this case.

$$P_{p0} = P_n(0) e^{(V_0 - V)/V_T} \dots \dots \dots (6)$$

Combining this equation with $V_0/V_T = E_0/KT$.

$$P_n(0) = P_{n0} e^{V/V_T} \dots \dots \dots (7)$$

This boundary condition is called the law of junction.

The hole concentration $P_n(0)$ injected into n side at the junction is obtained by subs equation (6) in equation (2)

$$P_n(0) = P_{n0} \left(e^{V/V_T} - 1 \right) \dots \dots \dots (8)$$

The Forward Currents:-

The hole current $I_{pn}(0)$ crossing the junction into n side is given by equation (4) we obtain

$$I_{pn}(0) = \frac{A_e D_p P_{n0}}{L_p} (e^{V/V_T} - 1) \dots \dots \dots (9)$$

The electron current $I_{np}(0)$ crossing junction into P side is obtained from equation 9 by interchanging n & P or

$$I_{np}(0) = \frac{A_e D_p n_{p0}}{L_n} (e^{V/V_T} - 1) \dots \dots \dots (10)$$

Total diode current is given by

$$I = I_{Ph}(O) + I_{NP}(O)$$

$$I = I_0 \left(e^{V/V_T} - 1 \right)$$

DIODE CURRENT

Diode current equation expresses the relationship between the current flowing through the diode as a function of the voltage applied across it. Mathematically it is given as

$$I = I_0 \left(e^{V/\eta V_T} - 1 \right)$$

where:

I = the net current flowing through the diode;

I_0 = "dark saturation current", the diode leakage current density in the absence of light;

V = applied voltage across the terminals of the diode;

q = absolute value of electron charge;

k = Boltzmann's constant; and

T = absolute temperature (K) & η = between 1 and 2, ideality factor

Temperature dependence of V-I characteristics of p-n junction diode:

The temperature has following effects on the diode parameters,

1. The cut-in voltage decreases as the temperature increases. The diode conducts at smaller voltage at large temperature.
2. The reverse saturation current increases as temperature increases.
This increases in reverse current I_0 is such that it doubles at every 10 °C rise in temperature.
Mathematically,

$$I_{o2} = 2^{(\Delta T/10)} I_{o1}$$

Where I_{o2} = Reverse current at T_2 °C, I_{o1} = Reverse current at T_1 °C

$\Delta T = (T_2 - T_1)$

3. The voltage equivalent of temperature V_T also increases as temperature increases.
4. The reverse breakdown voltage increases as temperature increases as shown.

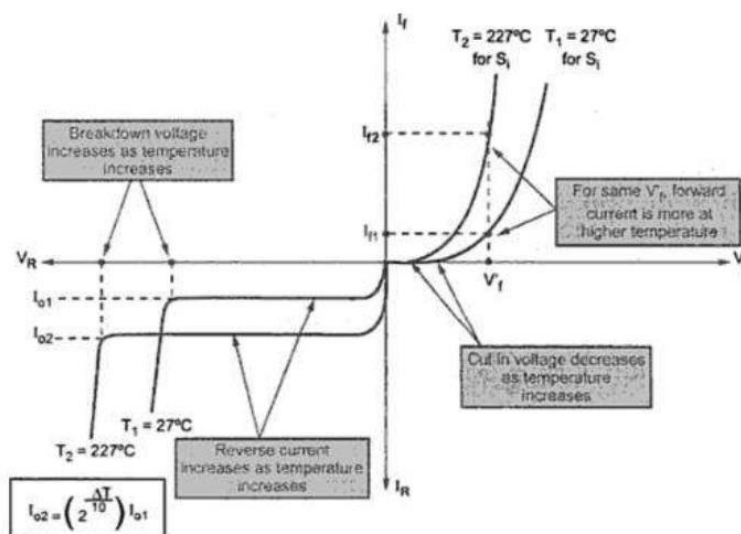


Figure 1.15: Effect of Temperature on PN Diode characteristics

(a) Effect of Temperature on forward voltage drop:

Most of the times, the drop across the diode is assumed constant. But in few situations, it is necessary to consider the effect of temperature on forward voltage drop.

It is seen that cut-in voltage decreases as temperature increases.

Note : The diode forward voltage drop decrease as temperature increases.

The rate at which it increases is - 2.3 mV/ °C for silicon while - 2.12 mV/ °C for germanium. The negative sign shows decrease in forward voltage drop as temperature increases. This is shown in figure below.

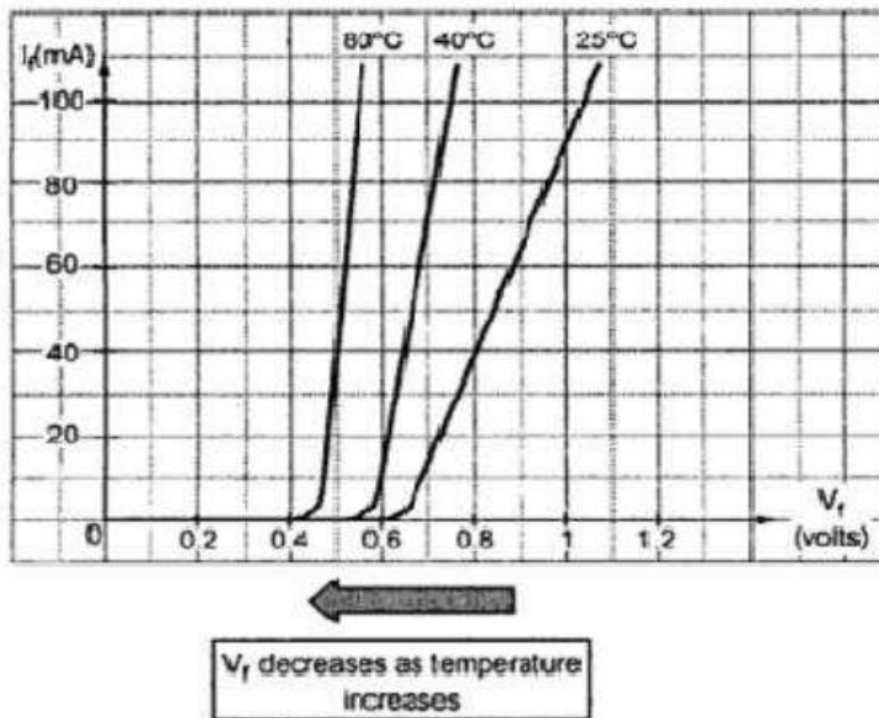


Figure 1.15: Effect of Temperature on forward voltage drop

coefficient $\Delta V_f / ^\circ\text{C}$ is called voltage/temperature coefficient of diode. Knowing this and V_{f1} at T_1 , any V_{f2} at T_2 can be obtained as,

$$V_{f2} = V_{f1} + [\Delta T \times \text{Voltage / temperature coefficient}]$$

(a) Effect of Temperature on Dynamic Resistance:

The dynamic resistance of the diode is obtained as,

$$r'_f = \frac{26 \text{ mV}}{I_f} = \frac{V_T}{I_f} = \frac{kT}{I_f}$$

where k = Boltzman's constant and T in $^\circ\text{K}$ constant.

The value 26 mV is temperature dependent and the above equation is applicable only at 25 °C. For higher temperatures, it gets changed as,

$$r'_f = \frac{kT'}{I_f}$$

where T' is new temperature in $^\circ\text{K}$. The above equation can be expressed

$$r'_f = \frac{26 \text{ mV}}{I_f} \left[\frac{T' \text{ in } ^\circ\text{K}}{298 ^\circ\text{C}} \right]$$

Note : As temperature increases, V_T increase hence dynamic forward resistance increases.

Diode Junction capacitance:

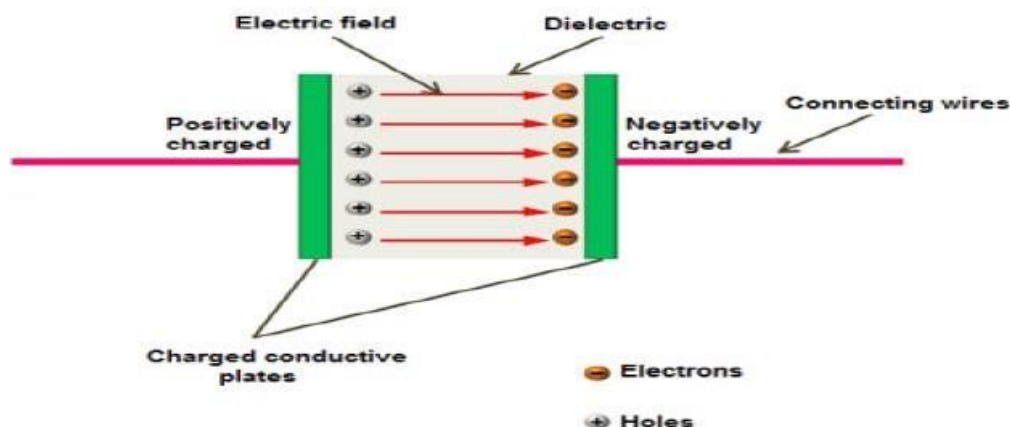
In a p-n junction diode, two types of capacitance take place. They are,

- Transition capacitance (C_T)
- Diffusion capacitance (C_D)

(a) Transition capacitance (C_T):

We know that capacitors store electric charge in the form of electric field. This charge storage is done by using two electrically conducting plates (placed close to each other) separated by an insulating material called dielectric.

The conducting plates or electrodes of the capacitor are good conductors of electricity. Therefore, they easily allow electric current through them. On the other hand, dielectric material or medium is poor conductor of electricity. Therefore, it does not allow electric current through it. However, it efficiently allows electric field.



When voltage is applied to the capacitor, charge carriers start flowing through the conducting wire. When these charge carriers reach the electrodes of the capacitor, they experience a strong opposition from the dielectric or insulating material. As a result, a large number of charge carriers are trapped at the electrodes of the capacitor. These charge carriers cannot move between the plates. However, they exert an electric field between the plates. The charge carriers which are trapped near the dielectric material will store electric charge. The ability of the material to store electric charge is called capacitance.

In a basic capacitor, the capacitance is directly proportional to the size of electrodes or plates and inversely proportional to the distance between two plates.

Just like the capacitors, a reverse biased p-n junction diode also stores electric charge at the depletion region. The depletion region is made of immobile positive and negative ions.

In a reverse biased p-n junction diode, the p-type and n-type regions have low resistance. Hence, p-type and n-type regions act like the electrodes or conducting plates of the capacitor. The depletion region of the p-n junction diode has high resistance. Hence, the depletion region acts like the dielectric or insulating material. Thus, p-n junction diode can be considered as a parallel

plate capacitor.

In depletion region, the electric charges (positive and negative ions) do not move from one place to another place. However, they exert electric field or electric force. Therefore, charge is stored at the depletion region in the form of electric field. The ability of a material to store electric charge is called capacitance. Thus, there exists a capacitance at the depletion region.

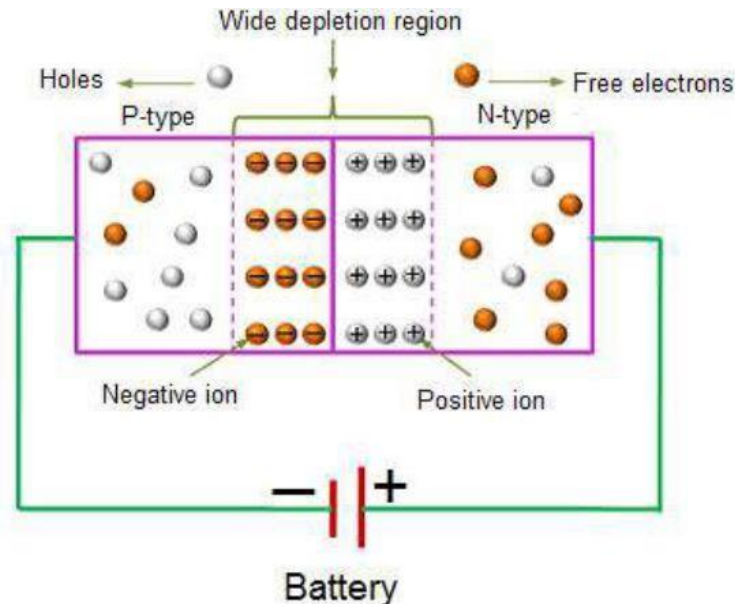


Figure 1.16: Transition Capacitance

The capacitance at the depletion region changes with the change in applied voltage. When reverse bias voltage applied to the p-n junction diode is increased, a large number of holes (majority carriers) from p-side and electrons (majority carriers) from n-side are moved away from the p-n junction. As a result, the width of depletion region increases whereas the size of p-type and n-type regions (plates) decreases.

We know that capacitance means the ability to store electric charge. The p-n junction diode with narrow depletion width and large p-type and n-type regions will store large amount of electric charge whereas the p-n junction diode with wide depletion width and small p-type and n-type regions will store only a small amount of electric charge. Therefore, the capacitance of the reverse bias p-n junction diode decreases when voltage increases.

In a forward biased diode, the transition capacitance exist. However, the transition capacitance is very small compared to the diffusion capacitance. Hence, transition capacitance is neglected in forward biased diode.

The amount of capacitance changed with increase in voltage is called transition capacitance. The transition capacitance is also known as depletion region capacitance, junction capacitance or barrier capacitance. Transition capacitance is denoted as C_T .

The change of capacitance at the depletion region can be defined as the change in electric charge per change in voltage.

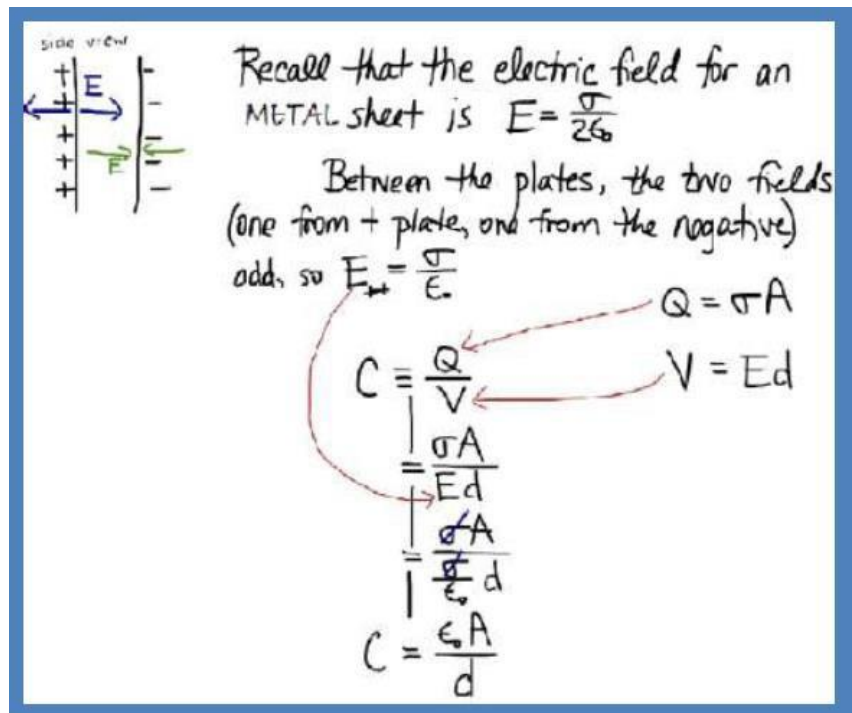
$$C_T = dQ / dV$$

Where,

C_T = Transition capacitance

dQ = Change in electric charge

dV = Change in voltage



The transition capacitance can be mathematically written as

$$C_T = \epsilon A / W$$

Where,

ϵ = Permittivity of the semiconductor

A = Area of plates or p-type and n-type regions W = Width of depletion region

Diffusion capacitance (C_D):

Diffusion capacitance occurs in a forward biased p-n junction diode. Diffusion capacitance is also sometimes referred as storage capacitance. It is denoted as C_D .

In a forward biased diode, diffusion capacitance is much larger than the transition capacitance. Hence, diffusion capacitance is considered in forward biased diode.

The diffusion capacitance occurs due to stored charge of minority electrons and minority holes near the depletion region.

When forward bias voltage is applied to the p-n junction diode, electrons (majority carriers) in the n-region will move into the p-region and recombines with the holes. In the similar way, holes in the p-region will move into the n-region and recombines with electrons. As a result, the width of depletion region decreases.

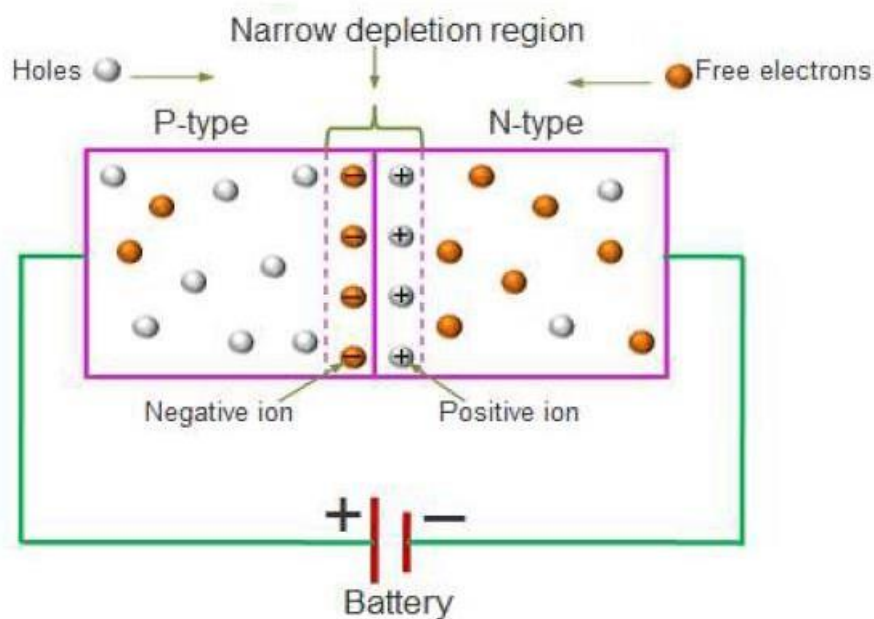


Figure 1.17: Diffusion Capacitance

The electrons (majority carriers) which cross the depletion region and enter into the p-region will become minority carriers of the p-region similarly; the holes (majority carriers) which cross the depletion region and enter into the n-region will become minority carriers of the n-region.

A large number of charge carriers, which try to move into another region will be accumulated near the depletion region before they recombine with the majority carriers. As a result, a large amount of charge is stored at both sides of the depletion region.

The accumulation of holes in the n-region and electrons in the p-region is separated by a very thin depletion region or depletion layer. This depletion region acts like dielectric or insulator of the capacitor and charge stored at both sides of the depletion layer acts like conducting plates of the capacitor.

Diffusion capacitance is directly proportional to the electric current or applied voltage. If large electric current flows through the diode, a large amount of charge is accumulated near the depletion layer. As a result, large diffusion capacitance occurs.

In the similar way, if small electric current flows through the diode, only a small amount of charge is accumulated near the depletion layer. As a result, small diffusion capacitance occurs.

When the width of depletion region decreases, the diffusion capacitance increases. The diffusion capacitance value will be in the range of nano farads (nF) to micro farads (μF).

The formula for diffusion capacitance is

$$C_D = dQ / dV$$

Where,

C_D = Diffusion capacitance

dQ = Change in number of minority carriers stored outside the depletion region dV = Change in voltage applied across diode

FORMULA FOR DIFFUSION CAPACITANCE

The formula for C_D is given by $C_D = dQ/dV$ where dQ represents the change in the number of minority carriers stored outside the depletion region when a change in voltage across diode, dV , applied.

If τ is mean lifetime of charge carriers, then a flow of charge Q yields a diode current I is given as

$$I = Q / \tau$$

In case of forward bias current is given by- $I = I_0 (e^{V/\eta V_T} - 1)$

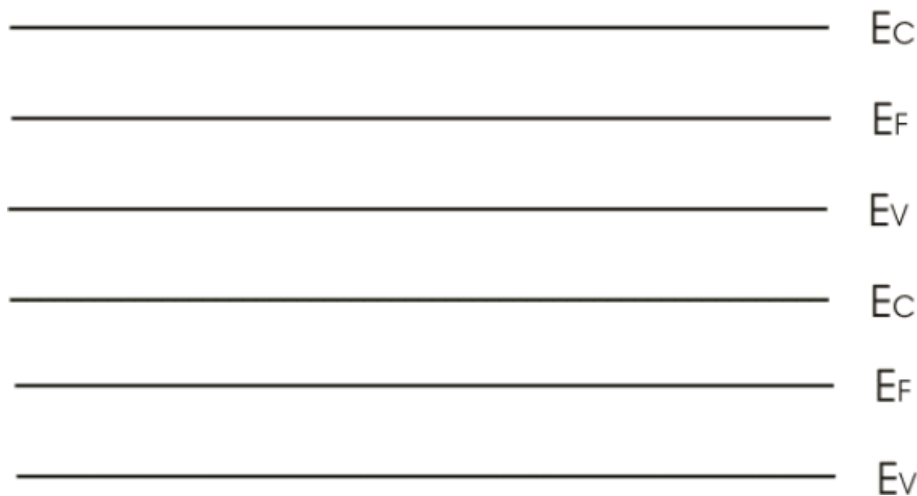
$$Q = \tau I_0 (e^{V/\eta V_T} - 1) = \tau I_0 e^{V/\eta V_T} \quad \text{as } e^{V/\eta V_T} \gg 1$$

$$\begin{aligned} \text{So diffusion capacitance } C_D &= \frac{dQ}{dV} = \frac{d(\tau I_0 e^{V/\eta V_T})}{dV} \\ &= \frac{\tau (I + I_0)}{\eta V_T} \quad \text{here } I \gg I_0 \end{aligned}$$

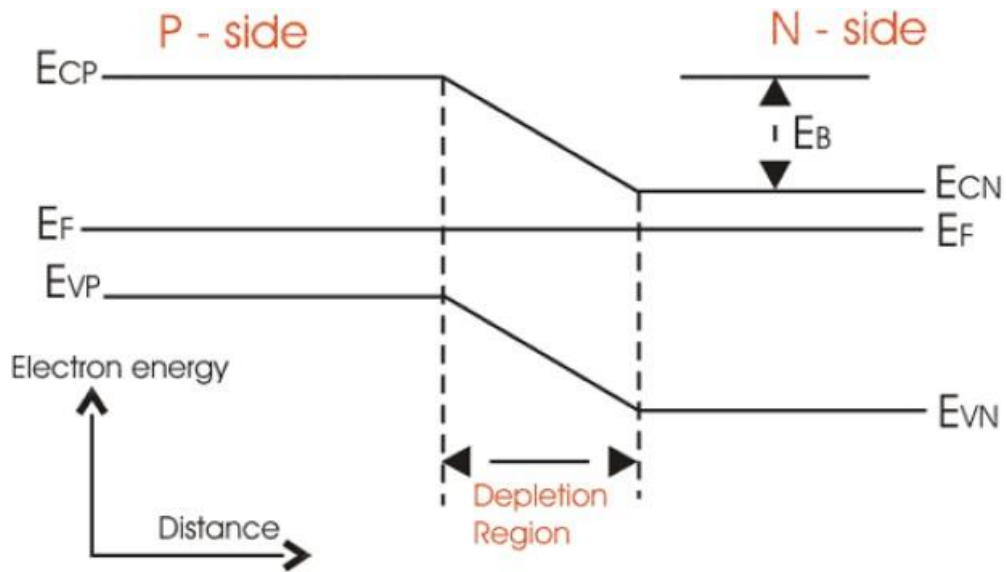
$$\text{So } C_D = \frac{\tau I}{\eta V_T}$$

Energy Band Diagram of P-N Diode:

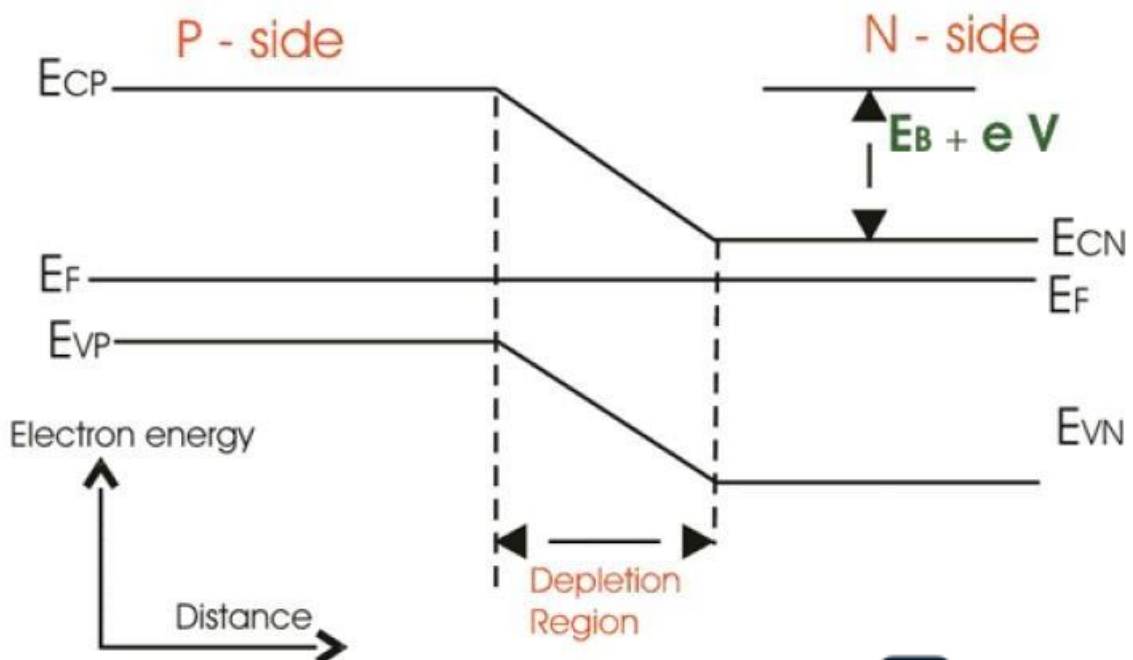
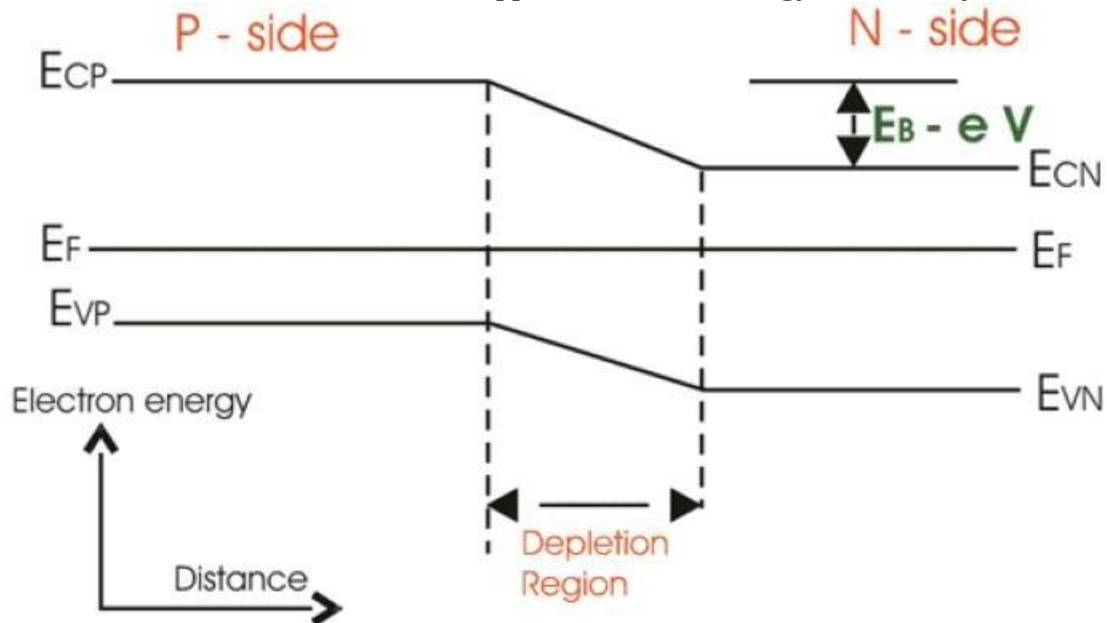
For an n-type semiconductor, the Fermi level E_F lies near the conduction band edge. E_C but for an p - type semiconductor, E_F lies near the valance band edge E_V .



Now, when a p-n junction is built, the Fermi energy E_F attains a constant value. In this scenario the p-sides conduction band edge. Similarly n-side valance band edge will be at higher level than E_{cn} , n-sides conduction band edge of p - side. This energy difference is known as barrier energy. The barrier energy is E_B
 $= E_{cp} - E_{cn} = E_{vp} - E_{vn}$



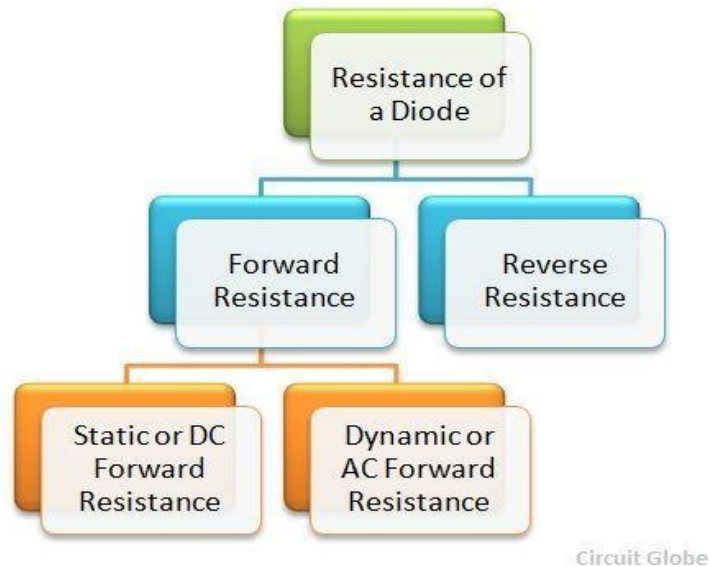
If we apply forward bias voltage V , across junction then the barrier energy decreases by an amount of eV and if V is reverse bias is applied the barrier energy increases by eV .



Diode Resistances:

In practice, no diode is an Ideal diode, this means neither it acts as a perfect conductor when forward biased nor it acts as an insulator when it is reverse biased. In other words an actual diode offers a very small resistance (not zero) when forward biased and is called a **forward resistance**. Whereas, it offers a very high resistance (not infinite) when reverse biased and is called as a **reverse resistance**.

The various resistances of a diode are as follows:



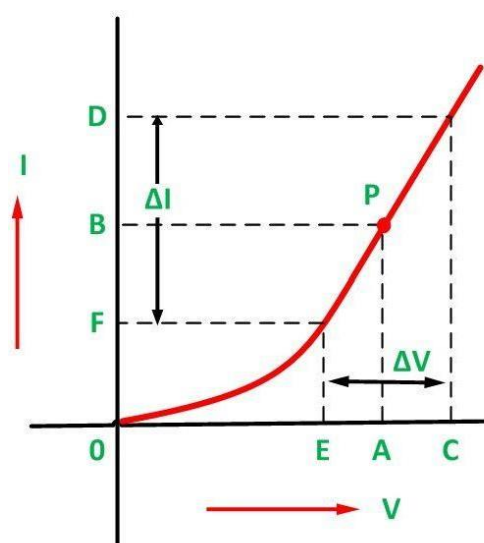
Forward Resistance

Under the Forward biased condition, the opposition offered by a diode to the forward current is known as Forward Resistance. The forward current flowing through a diode may be constant, i.e., direct current or changing i.e., alternating current. The forward resistance is classified as **Static Forward Resistance** and **Dynamic Forward Resistance**.

1) Static or DC Forward Resistance

The opposition offered by a diode to the direct current flowing forward bias condition is known as its **DC forward resistance** or Static Resistance. It is measured by taking the ratio of DC voltage across the diode to the DC current flowing through it.

The **forward characteristic** of a diode is shown below.



It is clear from the graph that for the operating point P, the forward voltage is OA and the corresponding forward current is OB. Therefore, the static forward resistance of the diode is given as

$$R_F = \frac{OA}{OB}$$

2) Dynamic or AC Forward Resistance

The opposition offered by a diode to the changing current flow I forward bias condition is known as its **AC Forward Resistance**. It is measured by a ratio of change in voltage across the diode to the resulting change in current through it. From the figure A above it is clear that for an operating point P the AC forward resistance is determined by varying the forward voltage (CE) on both the sides of the operating point equally and measuring the corresponding forward current (DF).

The Dynamic or AC Forward Resistance is represented as shown below.

$$r_f = \frac{CE}{DF} = \frac{\Delta V}{\Delta I}$$

Derivation of Dynamic Resistance

$$I = I_0 \left(e^{v/\eta v_T} - 1 \right)$$

$$\frac{dI}{dv} = \frac{d}{dv} \left[I_0 \left(e^{v/\eta v_T} - 1 \right) \right]$$

$$= \frac{I_0 e^{v/\eta v_T}}{\eta v_T} = \frac{I + I_0}{\eta v_T} = \frac{I}{\eta v_T}$$

$$r_f = \frac{dv}{dI} = \frac{\eta v_T}{I} = \frac{26 \text{ mV}}{I}$$

The value of the forward resistance of a crystal diode is very small, ranging from 1 to 25 Ohms.

Reverse Resistance (RR)

Under the Reverse biasing condition, the opposition offered by the diode to the reverse current is known as **Reverse Resistance**. Ideally, the reverse resistance of a diode is considered to be infinite. However, in actual practice the reverse resistance is not infinite because diode conducts a small leakage current (due to minority carriers) when reverse biased.

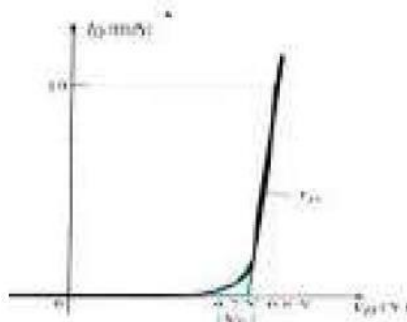
The value of reverse resistance is very large as compared to forward resistance. The ratio of reverse to forward resistance is 1 00 000 : 1 for silicon diodes, whereas it is 40 000 : 1 for germanium diode.

Diode Equivalent Circuits:

An equivalent circuit is a combination of elements properly chosen to best represent the actual terminal characteristics of a device, system, or such in a particular operating region. In other words, once the equivalent circuit is defined, the device symbol can be removed from a schematic and the equivalent circuit inserted in its place without severely affecting the actual behavior of the system. The result is often a network that can be solved using traditional circuit analysis techniques.

(a) Piecewise-Linear Equivalent Circuit :

One technique for obtaining an equivalent circuit for a diode is to approximate the characteristics of the device by straight-line segments, as shown in Fig. 1.21. The resulting equivalent circuit is naturally called the *piecewise-linear equivalent circuit*. It should be obvious from Fig. 1.21 that the straight-line segments do not result in an exact duplication of the actual characteristics, especially in the knee region. However, the resulting segments are sufficiently close to the actual curve to establish an equivalent circuit that will provide an excellent first approximation to the actual behavior of the device. The ideal diode is included to establish that there is only one direction of conduction through the device, and a reverse-bias condition will result in the open circuit state for the device. Since a silicon semiconductor diode does not reach the conduction state until V_D reaches 0.7 V with a forward bias (as shown in Fig. 1.21), a battery V_T opposing the conduction direction must appear in the equivalent circuit as shown in Fig. 1.21. The battery simply specifies that the voltage across the device must be greater than the threshold battery voltage before conduction through the device in the direction dictated by the ideal diode can be established. When conduction is established the resistance of the diode will be the specified value of r_{av} .



piecewise-linear equivalent circuit using straight-line segments to approximate the characteristic curve

The approximate level of r_{av} can usually be determined from a specified operating point on the specification sheet. For instance, for a silicon semiconductor diode, if $I_F = 10 \text{ mA}$ (a forward conduction current for the diode) at $V_D = 0.8 \text{ V}$, we know for silicon that a shift of 0.7 V is required before the characteristics rise.

$$r_{av} = \left. \frac{\Delta V_d}{\Delta I_d} \right|_{\text{pt. to pt.}} = \frac{0.8 \text{ V} - 0.7 \text{ V}}{10 \text{ mA} - 0 \text{ mA}} = \frac{0.1 \text{ V}}{10 \text{ mA}} = 10 \Omega$$

(b) Simplified Equivalent Circuits :

For most applications, the resistance r_{av} is sufficiently small to be ignored in comparison to the other elements of the network. The removal of r_{av} from the equivalent circuit is the same as implying that the characteristics of the diode. Under dc conditions has a drop of 0.7 V across it in the conduction state at any level of diode current.

(c) Ideal Equivalent Circuits:

Now that r_{av} has been removed from the equivalent circuit let us take it a step further and establish that a 0.7-V level can often be ignored in comparison to the applied voltage level. In this case the equivalent circuit will be reduced to that of an ideal diode as shown in Fig. 1.22 with its characteristics.

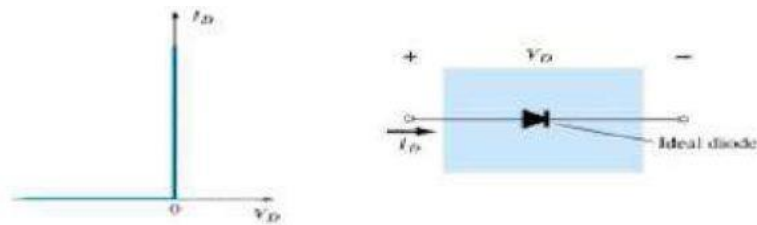


Figure 1.22: Ideal diode and its characteristics

Type	Conditions	Model	Characteristics
Piecewise-linear model			
Simplified model	$R_{\text{network}} \gg r_{av}$		
Ideal device	$R_{\text{network}} \gg r_{av}$ $E_{\text{network}} \gg V_F$		

Figure 1.22: Diode Equivalent Circuits

Rectifiers and Filters:

For the operation of most of the electronics devices and circuits, a d.c. source is required. So it is advantageous to convert domestic a.c. supply into d.c. voltages. The process of converting a.c. voltage into d.c. voltage is called as rectification.

This is achieved with i) Step-down Transformer, ii) Rectifier, iii) Filter and iv) Voltage regulator circuits.

These elements constitute d.c. regulated power supply shown in the fig 1 below.

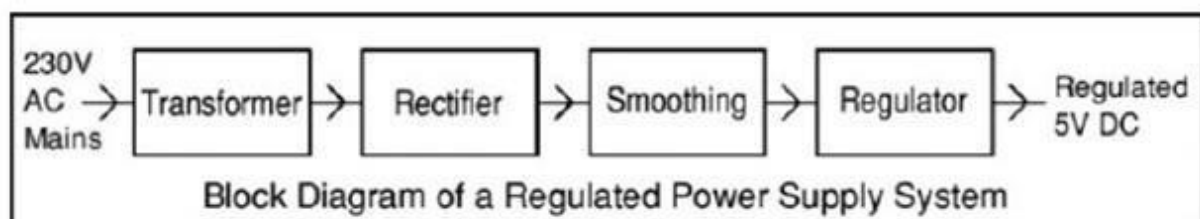


fig1 . Block diagram of Regulated D.C. Power Supply

- Transformer – steps down 230V AC mains to low voltage AC.
- Rectifier – converts AC to DC, but the DC output is varying.
- Smoothing – smooth the DC from varying greatly to a small ripple.

- Regulator – eliminates ripple by setting DC output to a fixed voltage.

The block diagram of a regulated D.C. power supply consists of step-down transformer, rectifier, filter, voltage regulator and load. An ideal regulated power supply is an electronics circuit designed to provide a predetermined d.c. voltage V_o which is independent of the load current and variations in the input voltage and temperature. If the output of a regulator circuit is a AC voltage then it is termed as voltage stabilizer, whereas if the output is a DC voltage then it is termed as voltage regulator.

RECTIFIER:

Any electrical device which offers a low resistance to the current in one direction but a high resistance to the current in the opposite direction is called rectifier. Such a device is capable of converting a sinusoidal input waveform, whose average value is zero, into a unidirectional Waveform, with a non-zero average component. A rectifier is a device, which converts a.c. voltage (bi-directional) to pulsating d.c. voltage (Unidirectional).

Characteristics of a Rectifier Circuit:

Any electrical device which offers a low resistance to the current in one direction but a high resistance to the current in the opposite direction is called rectifier. Such a device is capable of converting a sinusoidal input waveform, whose average value is zero, into a unidirectional waveform, with a non-zero average component.

A rectifier is a device, which converts a.c. voltage (bi-directional) to pulsating d.c.. Load currents: They are two types of output current. They are average or d.c. current and RMS currents.

Average or DC current: The average current of a periodic function is defined as the area of one cycle of the curve divided by the base.

It is expressed mathematically as

i. **AVERAGE VALUE/DC VALUE/MEAN VALUE** = $\frac{\text{Area over one period}}{\text{Total time period}}$

$$V_{dc} = \frac{1}{T} \int_0^T V d(wt)$$

ii. EFFECTIVE (OR) R.M.S CURRENT:

The effective (or) R.M.S. current squared of a periodic function of time is given by the area of one cycle of the curve, which represents the square of the function divided by the base.

$$V_{rms} = \sqrt{\frac{1}{T} \int_0^T V^2 d(wt)}$$

iii. PEAK FACTOR:

It is the ratio of peak value to Rms value

$$\text{Peak factor} = \frac{\text{peakvalue}}{\text{rmsvalue}}$$

iv. FORM FACTOR:

It is the ratio of Rms value to average value

$$\text{Form factor} = \frac{\text{Rmsvalue}}{\text{averagevalue}}$$

v. **RIPPLE FACTOR (Γ)**: It is defined as ration of R.M.S. value of a.c. component to the d.c. component in the output is known as “Ripple Factor

$$\Gamma = \frac{V_{ac}}{V_{dc}}$$

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

vi. EFFICIENCY:

It is the ratio of d.c output power to the a.c. input power. It signifies, how efficiently the rectifier circuit converts a.c. power into d.c. power.

$$\eta = \frac{o/p \text{ power}}{i/p \text{ power}}$$

vii. PEAK INVERSE VOLTAGE (PIV):

It is defined as the maximum reverse voltage that a diode can withstand without destroying the junction.

viii. TRANSFORMER UTILIZATION FACTOR (TUF):

The d.c. power to be delivered to the load in a rectifier circuit decides the rating of the Transformer used in the circuit. So, transformer utilization factor is defined as

$$TUF = \frac{P_{dc}}{P_{ac(rated)}}$$

ix. % REGULATION:

The variation of the d.c. output voltage as a function of d.c. load current is called regulation. The percentage regulation is defined as

$$\% \text{ Regulation} = \frac{V_{NL} - V_{FL}}{V_{FL}} * 100$$

For an ideal power supply, % Regulation is zero.

2.0 CLASSIFICATION OF RECTIFIERS:

Using one or more diodes in the circuit, following rectifier circuits can be designed.

- 1) Half - Wave Rectifier
- 2) Full – Wave Rectifier
- 3) Bridge Rectifier

2.0.1) HALF-WAVE RECTIFIER:

A Half – wave rectifier as shown in **fig 2** is one, which converts a.c. voltage into a pulsating voltage using only one half cycle of the applied a.c. voltage.

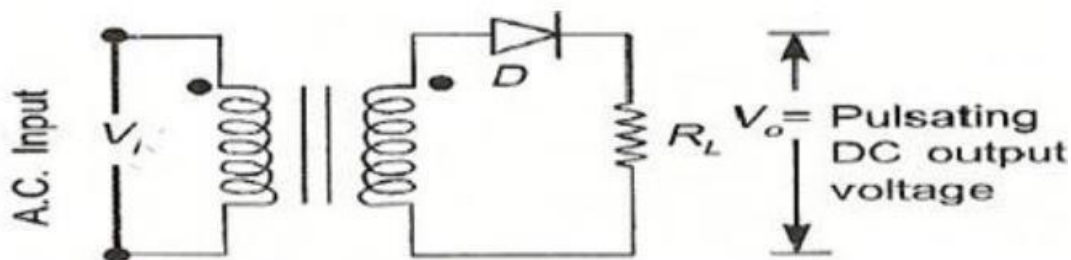


fig 2 Basic structure of Half-Wave Rectifier

The a.c. voltage is applied to the rectifier circuit using step-down transformer-rectifying element i.e., p-n junction diode and the source of a.c. voltage, all connected in series. The a.c. voltage is applied to the rectifier circuit using step-down transformer.

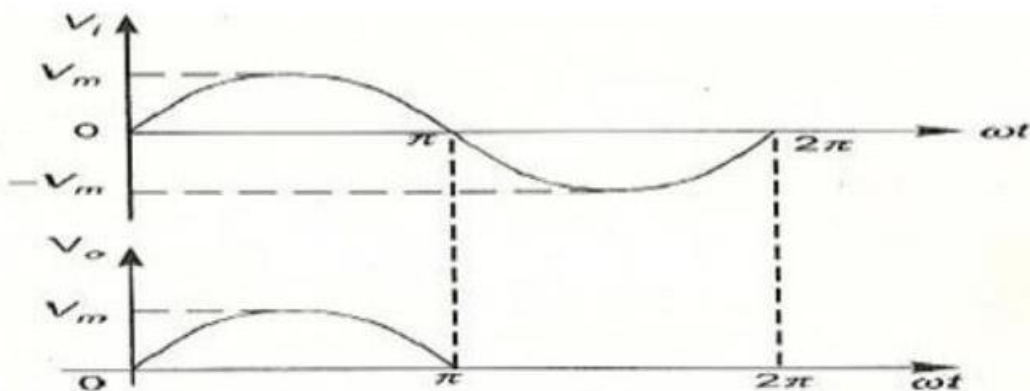


fig 3 Input and output waveforms of a Half wave rectifier

$$V = V_m \sin(\omega t)$$

The input to the rectifier circuit, Where V_m is the peak value of secondary a.c. voltage.

Operation:

For the positive half-cycle of input a.c. voltage, the diode D is forward biased and hence it

conducts. Now a current flows in the circuit and there is a voltage drop across RL. The waveform of the diode current (or) load current is shown in **fig 3**.

For the negative half-cycle of input, the diode D is reverse biased and hence it does not Conduct. Now no current flows in the circuit i.e., $i=0$ and $V_o=0$. Thus for the negative half-cycle no power is delivered to the load.

Analysis:

In the analysis of a HWR, the following parameters are to be analyzed.

1. DC output current
2. DC Output voltage
3. R.M.S. Current
4. R.M.S. voltage
5. Rectifier Efficiency (η)
6. Ripple factor (γ)
7. Peak Factor
8. % Regulation
9. Transformer Utilization Factor (TUF)
10. form factor
11. o/p frequency

Let a sinusoidal voltage V_i be applied to the input of the rectifier.

Then $V=V_m \sin (wt)$ Where V_m is the maximum value of the secondary voltage. Let the diode be idealized to piece-wise linear approximation with resistance R_f in the forward direction i.e., in the ON state and $R_r (= \infty)$ in the reverse direction i.e., in the OFF state. Now the current 'i' in the diode (or) in the load resistance RL is given by $V=V_m \sin (wt)$

i) AVERAGE VOLTAGE

$$V_{dc} = \frac{1}{T} \int_0^T V d(wt)$$

$$V_{dc} = \frac{1}{T} \int_0^{2\pi} V(\alpha) d\alpha$$

$$V_{dc} = \frac{1}{2\pi} \int_{\pi}^{2\pi} V(\alpha) d\alpha$$

$$V_{dc} = \frac{1}{2\pi} \int_0^{\pi} V_m \sin(wt) d\alpha$$

$$V_{dc} = \frac{V_m}{\pi}$$

ii).AVERAGE CURRENT:

$$I_{dc} = \frac{I_m}{\Pi}$$

iii) RMS VOLTAGE:

$$V_{rms} = \sqrt{\frac{1}{T} \int_0^T V^2 d(wt)}$$

$$V_{rms} = \sqrt{\frac{1}{2\Pi} \int_0^{2\Pi} (V_m \sin(wt))^2 d(wt)}$$

$$V_{rms} = \frac{V_m}{2}$$

(iv) RMS CURRENT:

$$I_{rms} = \frac{I_m}{\Pi}$$

(v) PEAK FACTOR:

$$\text{Peak factor} = \frac{\text{peakvalue}}{\text{rmsvalue}}$$

$$\text{Peak Factor} = \frac{V_m}{(V_m / 2)}$$

$$\text{Peak Factor} = 2$$

(vi) FORM FACTOR:

$$\text{Form factor} = \frac{\text{Rmsvalue}}{\text{averagevalue}}$$

$$\text{Form factor} = \frac{(V_m / 2)}{V_m / \Pi}$$

$$\text{Form Factor} = 1.57$$

(vii) RIPPLE FACTOR:

$$\Gamma = \frac{V_{ac}}{V_{dc}}$$

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

$$\Gamma = \frac{\sqrt{V_{rms}^2 - V_{dc}^2}}{V_{ac}}$$

$$\Gamma = \sqrt{\frac{V_{rms}^2}{V_{dc}^2} - 1}$$

$$\Gamma = 1.21$$

(viii) **EFFICIENCY(η):**

$$\eta = \frac{o / p_{power}}{i / p_{power}} * 100$$

$$\eta = \frac{P_{ac}}{P_{dc}} * 100$$

$$\eta = 40.8$$

(ix) **TRANSFORMER UTILIZATION FACTOR (TUF):**

The d.c. power to be delivered to the load in a rectifier circuit decides the rating of the transformer used in the circuit. Therefore, transformer utilization factor is defined as

$$TUF = \frac{P_{dc}}{P_{ac(rated)}}$$

$$= 0.286.$$

$$TUF$$

The value of TUF is low which shows that in half-wave circuit, the transformer is not fully utilized. If the transformer rating is 1 KVA (1000VA) then the half-wave rectifier can deliver $1000 \times 0.287 = 287$ watts to resistance load.

(x) **PEAK INVERSE VOLTAGE (PIV):**

It is defined as the maximum reverse voltage that a diode can withstand without destroying the junction. The peak inverse voltage across a diode is the peak of the negative half- cycle. For half-wave rectifier, PIV is V_m .

DISADVANTAGES OF HALF-WAVE RECTIFIER:

1. The ripple factor is high.
2. The efficiency is low.
3. The Transformer Utilization factor is low.

Because of all these disadvantages, the half-wave rectifier circuit is normally not used as a power rectifier circuit.

FULL WAVE RECTIFIER:

A full-wave rectifier converts an ac voltage into a pulsating dc voltage using both half cycles of the applied ac voltage. In order to rectify both the half cycles of ac input, two diodes are used in this circuit. The diodes feed a common load R_L with the help of a center-tap transformer. A center-tap transformer is the one, which produces two sinusoidal waveforms of same magnitude and frequency but out of phase with respect to the ground in the secondary winding of the transformer. The full wave rectifier is shown in the **fig 4** below

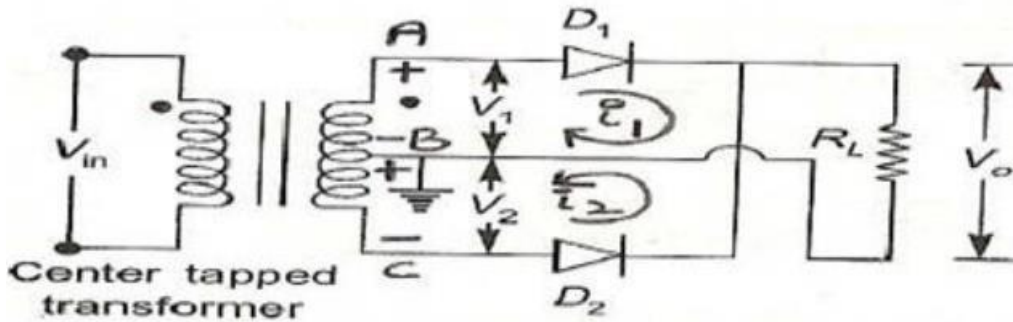


fig 4 Full-Wave Rectifier.

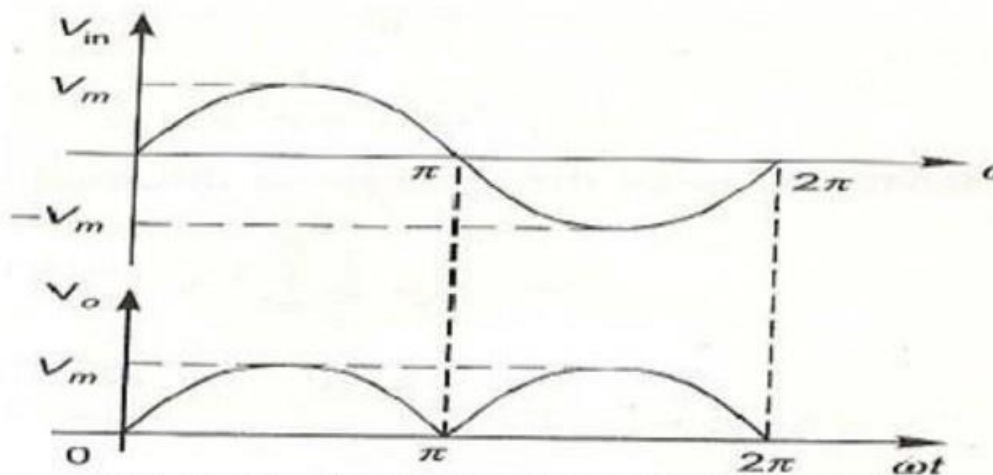


Fig. 5 input and output waveforms of Fullwave rectifier

Fig. 5 shows the input and output wave forms of the ckt.

During positive half of the input signal, anode of diode D_1 becomes positive and at the same time the anode of diode D_2 becomes negative. Hence D_1 conducts and D_2 does not conduct. The load current flows through D_1 and the voltage drop across R_L will be equal to the input voltage.

During the negative half cycle of the input, the anode of D_1 becomes negative and the anode of D_2 becomes positive. Hence, D_1 does not conduct and D_2 conducts. The load current flows through D_2 and the voltage drop across R_L will be equal to the input voltage. It is noted that the load current flows in the both the half cycles of ac voltage and in the same direction through the load resistance.

i) AVERAGE VOLTAGE

$$V_{dc} = I_{dc} \cdot R_L = \frac{2 I_m}{\pi} \cdot R_L \quad \text{We know } I_m = \frac{V_m}{R_s + R_f + R_L}$$

$$\therefore V_{dc} = \frac{2 V_m R_L}{\pi (R_s + R_f + R_L)}$$

$$\text{If } (R_s + R_f) \ll R_L$$

$$V_{dc} = \frac{2 V_m}{\pi} = 0.637 V_m.$$

ii) AVERAGE CURRENT

$$\begin{aligned} I_{dc} &= \frac{1}{2\pi} \int_0^{2\pi} i d\theta = \frac{1}{2\pi} \int_0^{2\pi} I_m \sin \theta d\theta \\ &= \frac{I_m}{2\pi} \left[\int_0^{\pi} \sin \theta d\theta - \int_{\pi}^{2\pi} \sin \theta d\theta \right] \\ &= \frac{I_m}{2\pi} [(-2)(-2)] \\ &= \frac{I_m}{2\pi} \cdot 4 = \frac{2 I_m}{\pi} = 0.637 I_m. \end{aligned}$$

$$\boxed{I_{dc} = 0.637 I_m.}$$

$$\therefore I_{DC} \text{ FWR} = 2 I_{DC} \text{ HWR.}$$

iii) RMS VOLTAGE:

$$V_{rms} = \sqrt{\frac{1}{T} \int_0^T V^2 d(wt)}$$

$$V_{rms} = \sqrt{\frac{1}{2\pi} \int_0^{2\pi} (V_m \sin(wt))^2 d(wt)}$$

$$\boxed{V_{rms} = \frac{V_m}{\sqrt{2}}}$$

iv) RMS CURRENT:

$$\boxed{I_{rms} = \frac{2 I_m}{\pi}}$$

v) PEAK FACTOR

$$\text{Peak factor} = \frac{\text{peakvalue}}{\text{rmsvalue}}$$

$$\text{Peak Factor} = \frac{V_m}{(V_m / 2)}$$

$$\text{Peak Factor} = 2$$

vi) FORM FACTOR

$$\text{Form factor} = \frac{\text{Rms value}}{\text{average value}}$$

$$\text{Form factor} = \frac{(V_m / \sqrt{2})}{2V_m / \pi}$$

$$\text{Form Factor} = 1.11$$

vii) RIPPLE FACTOR:

$$\gamma = \sqrt{\left(\frac{I_{rms}}{I_{dc}}\right)^2 - 1}$$

for FWR,

$$I_{rms} = \frac{I_m}{\sqrt{2}} \quad \& \quad I_{DC} = \frac{2 I_m}{\pi}$$

$$\therefore \gamma_{FWR} = \sqrt{\left(\frac{I_m / \sqrt{2}}{2 I_m / \pi}\right)^2 - 1}$$

$$= \sqrt{\left(\frac{\pi}{2\sqrt{2}}\right)^2 - 1}$$

$$= \sqrt{\left(\frac{3.1416}{2 \times 1.414}\right)^2 - 1} = 0.483$$

viii) Efficiency (η):

$$\eta = \frac{o / p_{power}}{i / p_{power}} * 100$$

$$\eta = \frac{P_{dc}}{P_{ac}} \times 100\%$$

$$\text{For FWR, } P_{dc} = I_{dc}^2 \cdot R_L = \left(\frac{2}{\pi} \cdot I_m \right)^2 \cdot R_L$$

$$P_{ac} = I_{rms}^2 (R_f + R_s + R_L)$$

$$\left(\frac{I_m}{\sqrt{2}} \right)^2 (R_f + R_s + R_L)$$

$$\eta = \frac{\frac{I_m^2 4}{\pi^2} \cdot R_L}{\frac{I_m^2}{2} \cdot (R_f + R_s + R_L)}$$

$$\text{If } (R_f + R_s) \ll R_L$$

$$\eta = \frac{4}{\pi^2} \cdot \frac{2}{1} = \frac{8}{\pi^2} = 0.812 = 81.2\%$$

ix) TRANSFORMER UTILIZATION FACTOR (TUF):

The d.c. power to be delivered to the load in a rectifier circuit decides the rating of the transformer used in the circuit. Therefore, transformer utilization factor is defined as

$$TUF = \frac{P_{dc}}{P_{ac(rated)}}$$

$$a) \text{ TUF (Secondary)} = \frac{P_{dc} \text{ delivered to load}}{AC \text{ power rating of transformer secondary}}$$

$$b) \text{ Since both the windings are used } TUF_{FWR} = 2 TUF_{HWR}$$

$$= 2 \times 0.287 = 0.574$$

$$c) \text{ TUF primary} = \text{Rated efficiency} = \frac{P_{dc}}{P_{ac}} \times 100 = 81.2\%$$

$$d) \text{ Average} = \frac{0.812 + 0.574}{2} = 0.693$$

x) Peak Inverse Voltage (PIV):

It is defined as the maximum reverse voltage that a diode can withstand without destroying the junction. The peak inverse voltage across a diode is the peak of the negative half-cycle. For half-wave rectifier, PIV is $2V_m$.

xi) % Regulation

$$\text{Voltage regulation} =$$

$$= \frac{I_{dc} (R_s + R_f)}{\frac{2V_m}{\pi} - I_{DC} (R_f + R_s)}$$

Advantages

- 1) Ripple factor = 0.482 (against 1.21 for HWR)
- 2) Rectification efficiency is 0.812 (against 0.405 for HWR)
- 3) Better TUF (secondary) is 0.574 (0.287 for HWR)
- 4) No core saturation problem

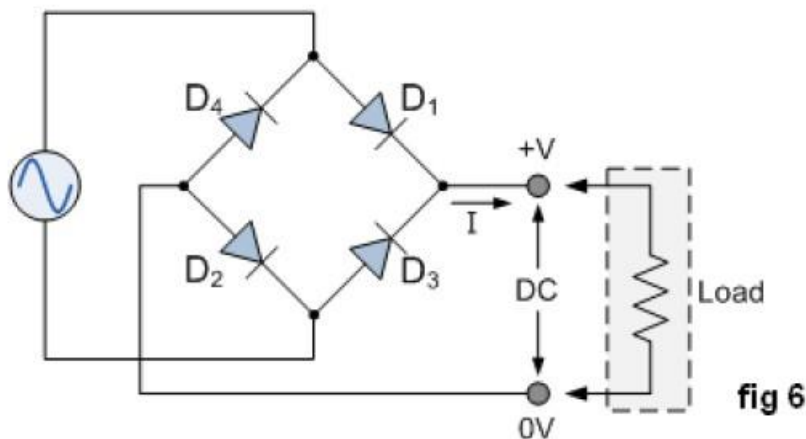
Disadvantages:

1. Requires centre tapped transformer.

BRIDGE RECTIFIER.

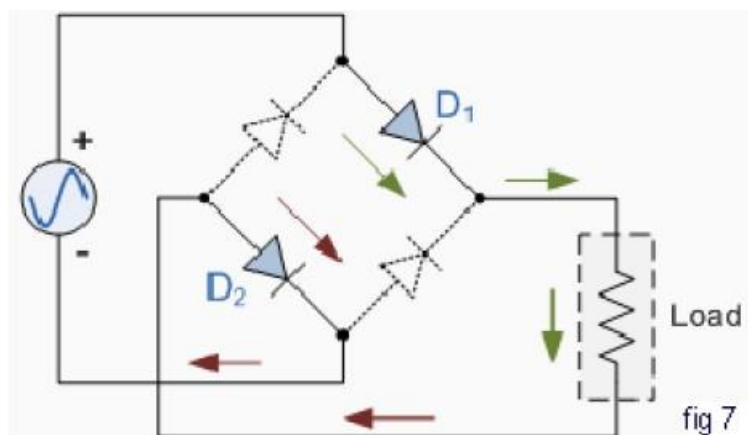
Another type of circuit that produces the same output waveform as the full wave rectifier circuit above, is that of the **Full Wave Bridge Rectifier**. This type of single phase rectifier uses four individual rectifying diodes connected in a closed loop "bridge" configuration to produce the desired output. The main advantage of this bridge circuit is that it does not require a special centre tapped transformer, thereby reducing its size and cost. The single secondary winding is connected to one side of the diode bridge network and the load to the other side as shown below.

The Diode Bridge Rectifier



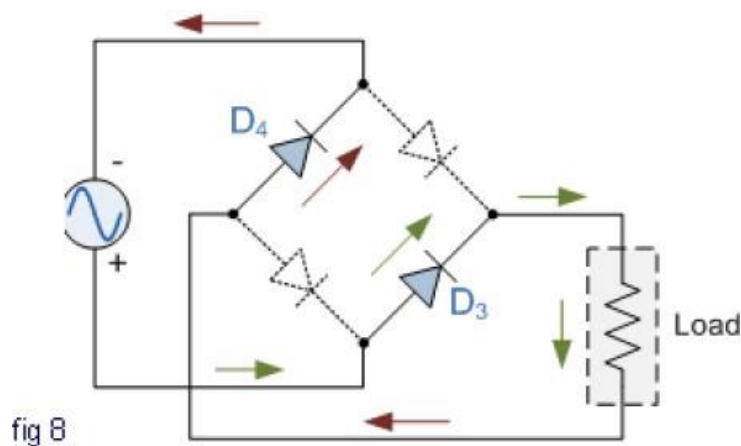
The four diodes labelled D_1 to D_4 are arranged in "series pairs" with only two diodes conducting current during each half cycle. During the positive half cycle of the supply, diodes D_1 and D_2 conduct in series while diodes D_3 and D_4 are reverse biased and the current flows through the load as shown below (fig 7).

The Positive Half-cycle



The Negative Half-cycle

During the negative half cycle of the supply, diodes D_3 and D_4 conduct in series (fig 8), but diodes D_1 and D_2 switch "OFF" as they are now reverse biased. The current flowing through the load is the same direction as before.



As the current flowing through the load is unidirectional, so the voltage developed across the load is also unidirectional the same as for the previous two diode full-wave rectifier, therefore the average DC voltage across the load is $0.637V_{\max}$. However in reality, during each half cycle the current flows through two diodes instead of just one so the amplitude of the output voltage is two voltage drops ($2 \times 0.7 = 1.4V$) less than the input $V_{\max}$ amplitude. The ripple frequency is now twice the supply frequency (e.g. 100Hz for a 50Hz supply)

Therefore, the following expressions are same as that of full wave rectifier.

a) Average current $I_{dc} = \frac{2I_m}{\pi}$

b) RMS current $I_{rms} = \frac{I_m}{\sqrt{2}}$

c) DC output voltage (no.load) $V_{DC} = \frac{2V_m}{\pi}$

d) Ripple factor $\gamma = 0.482$

e) Rectification efficiency $\eta = 0.812$

f) DC output voltage full load.

$$= V_{DCFL} = \frac{2V_m}{\pi} - I_{dc}(R_s + 2R_f); \quad \text{i.e., less by one diode loss.}$$

TUF of both primary & secondary are 0.812 therefore TUF overall is 0.812 (better than FWR with 0.693)

Comparison:

Sl No.	Parameter	HWR	FWR	BR
1	No. of diodes	1	2	4
2	PIV of diodes	V_m	$2V_m$	V_m
3	Secondary voltage (rms)	V	$V/0.7$	V
4	DC output voltage at no load	$\frac{V_m}{\pi} = 0.318 V_m$	$\frac{2V_m}{\pi} = 0.636 V_m$	$\frac{2V_m}{\pi} = 0.636 V_m$
5	Ripple factor γ	1.21	0.482	0.482
6	Ripple frequency	f	$2f$	$2f$
7	Rectification efficiency η	0.406	0.812	0.812
8	TUF	0.287	0.693	0.812

FILTERS

The output of a rectifier contains dc component as well as ac component. Filters are used to minimize the undesirable ac i.e., ripple leaving only the dc component to appear at the output.

Some important filters are:

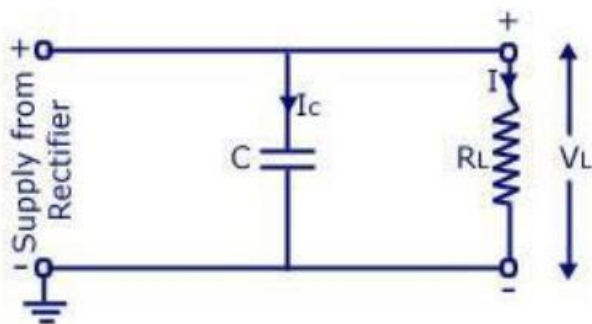
1. Inductor filter
2. Capacitor filter
3. LC or L section filter
4. CLC or Π -type filter

CAPACITOR FILTER

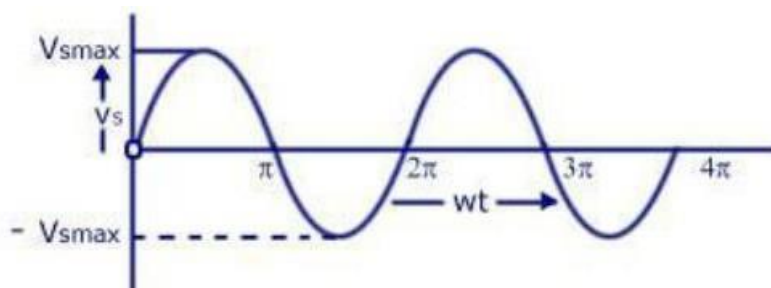
This is the most simple form of the **filter circuit** and in this arrangement a high value capacitor C is placed directly across the output terminals, as shown in figure. During the conduction period it gets charged and stores up energy to it during non-conduction period. Through this process, the time duration during which F_t is to be noted here that the capacitor C gets charged to the peak because there is no resistance (except the negligible forward resistance of diode) in the charging path. But the discharging time is quite large (roughly 100 times more than the charging time depending upon the value of R_L) because it discharges through load resistance R_L .

The function of the capacitor filter may be viewed in terms of impedances. The large value capacitor C offers a low impedance shunt path to the ac components or ripples but offers high impedance to the dc component. Thus ripples get bypassed through capacitor C and only dc component flows through the load resistance R_L .

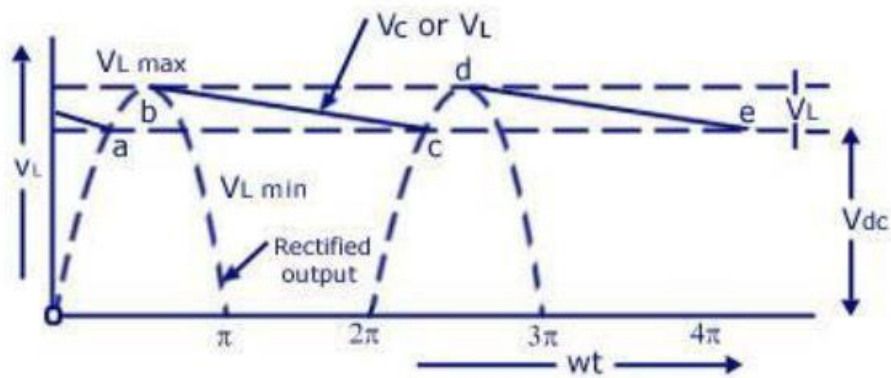
Capacitor filter is very popular because of its low cost, small size, light weight and good characteristics



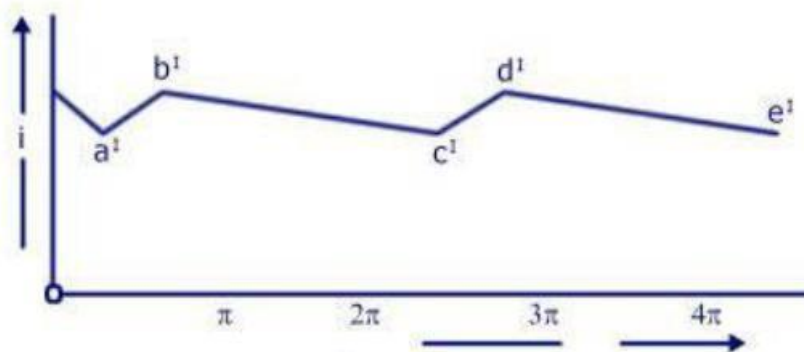
Circuit Diagram



Input voltage Waveform to Rectifier

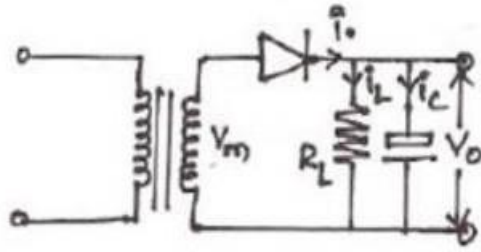


Rectified and filtered Output Voltage Waveform



Load Current Waveform
Half-wave Rectifier With Shunt Capacitor Filter

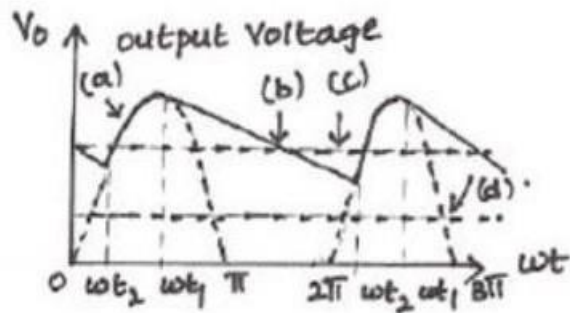
CAPACITOR FILTER WITH HWR:



Cut In angle – ωt_2

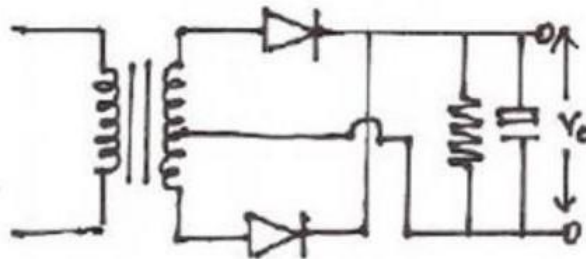
Cut out angle = ωt_1

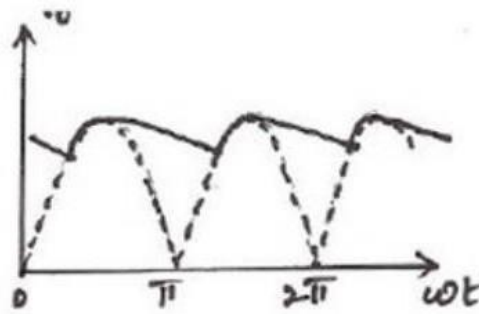
$$\omega t_1 = \pi - \tan^{-1} \omega C R_L$$



- (a) Capacitor charging through diode
($\omega t_2 - \omega t_1$)
- (b) Capacitor discharging through R_L
(ωt_1 to ωt_2)
- (c) Average (DC) voltage with filter
- (d) Average (DC) voltage without filter.

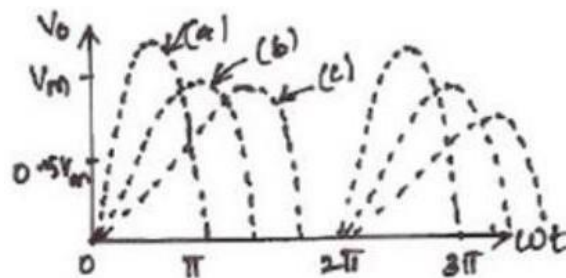
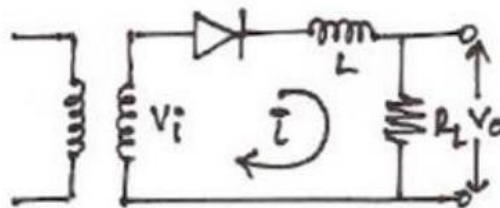
CAPACITOR FILTER WITH FWR:





$$\text{Ripple factor } r = \frac{1}{4\sqrt{3}fCR_L}$$

Ripple freq_{FWR} = 2 ripple freq_{HWR}.

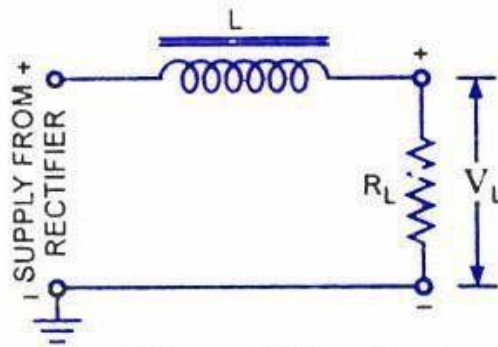


$$(a) \frac{W_L}{P_L} = 0 \quad (b) \frac{W_L}{P_L} = 1 \quad (c) \frac{W_L}{P_L} = 5.$$

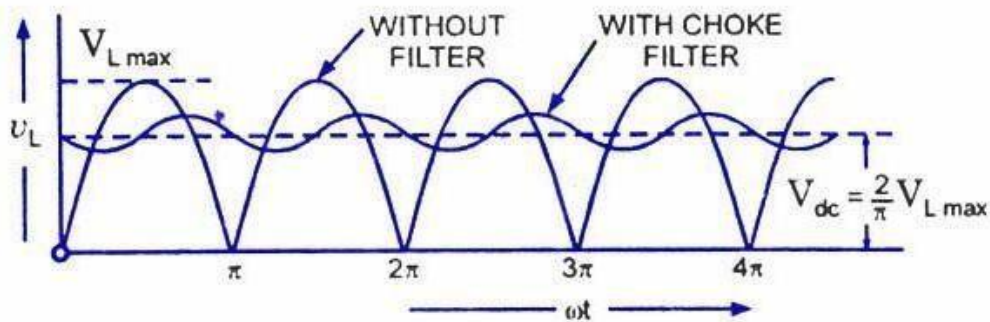
The worth noting points about shunt capacitor filter are:

1. For a fixed-value filter capacitance larger the load resistance R_L larger will be the discharge time constant $C R_L$ and therefore, lower the ripples and more the output voltage. On the other hand lower the load resistance (or more the load current), lower will be the output voltage.
2. Similarly smaller the filter capacitor, the less charge it can hold and more it will discharge. Thus the peak-to-peak value of the ripple will increase, and the average dc level will decrease. Larger the filter capacitor, the more charge it can hold and the less it will discharge. Hence the peak-to-peak value of the ripple will be less, and the average dc level will increase. But, the maximum value of the capacitance that can be employed is limited by another factor. The larger the capacitance value, the greater is the current required to charge the capacitor to a given voltage. The maximum current that can be handled by a diode is limited by the figure quoted by the manufacturer. Thus the maximum value of the capacitance that can be used in the shunt filter capacitor is limited.

SERIES INDUCTOR FILTER.



Circuit Diagram



Output Voltage Waveforms

Full-Wave Rectifier With Series Inductor Filter

In this arrangement a high value inductor or choke L is connected in series with the rectifier element and the load, as illustrated in figure. The filtering action of an inductor filter depends upon its property of opposing any change in the current flowing through it. When the output current of the rectifier increases above a certain value, energy is stored in it in the form of magnetic field and this energy is given up when the output current falls below the average value. Thus by placing a choke coil in series with the rectifier output and load, any sudden change in current that might have occurred in the circuit without an inductor is smoothed out by the presence of the inductor L .

The function of the inductor filter may be viewed in terms of impedances. The choke offers high impedance to the ac components but offers almost zero resistance to the desired dc components. Thus ripples are removed to a large extent. Nature of the output voltage without filter and with choke filter is shown in figure.

For dc (zero frequency), the choke resistance R_c in series with the load resistance R_L forms a voltage divider and dc voltage across the load is given as where V_{dc} is dc voltage output from a full-wave rectifier. Usually choke coil resistance R_c , is much small than R_L and, therefore, almost entire of the dc voltage is available across the load resistance R_L .

Since the reactance of inductor increases with the increase in frequency, better filtering of the higher harmonic components takes place, so effect of third and higher harmonic voltages can be neglected.

As obvious from equation , if choke coil resistance R_c is negligible in comparison to load resistance R_L , then the entire dc component of rectifier output is available across $2 R_L$ and is equal to $\frac{1}{2} V_{L max}$. The ac voltage partly drops across X_L and partly over R_L .

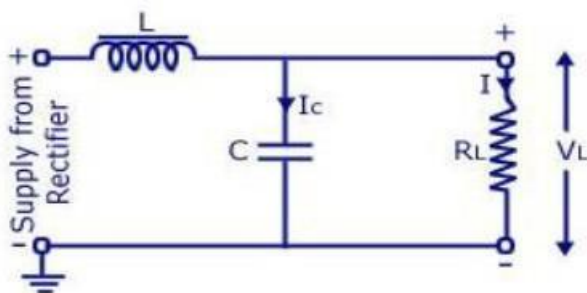
L-SECTION FILTER:

A simple series inductor reduces both the peak and effective values of the output current and output voltage. On the other hand a simple shunt capacitor filter reduces the ripple voltage but increases the diode current. The diode may get damaged due to large current and at the same time it causes greater heating of supply transformer resulting in reduced efficiency.

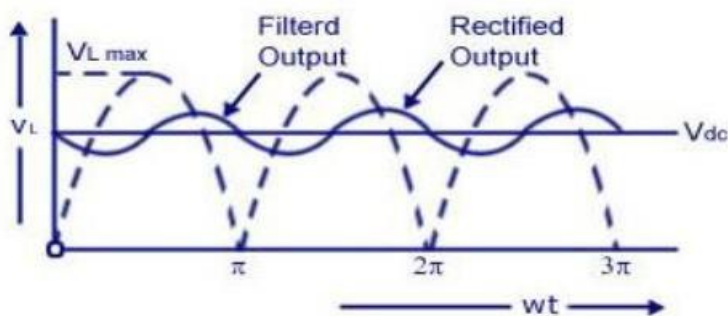
In an inductor filter, ripple factor increases with the increase in load resistance R_L while in a capacitor filter it varies inversely with load resistance R_L .

From economical point of view also, neither series inductor nor shunt capacitor type filters are suitable.

Practical filter-circuits are derived by combining the voltage stabilizing action of shunt capacitor with the current smoothing action of series choke coil. By using combination of inductor and capacitor ripple factor can be lowered, diode current can be restricted and simultaneously ripple factor can be made almost independent of load resistance (or load current). Two types of most commonly used combinations are choke-input or L-section filter and capacitor-input or Pi-Filter.



Circuit Diagram



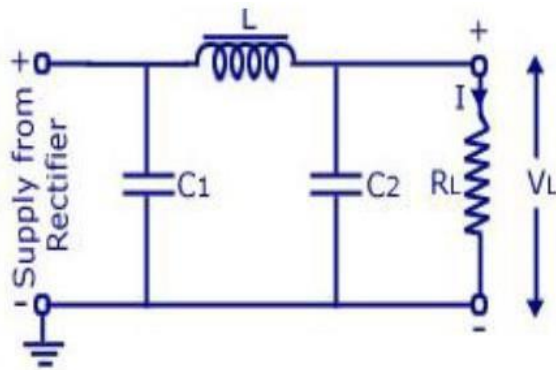
Rectified and Filtered Output Voltage Waveform
Full-wave Rectifier With Choke-Input Filter

CHOKE-INPUT FILTER IS EXPLAINED BELOW:

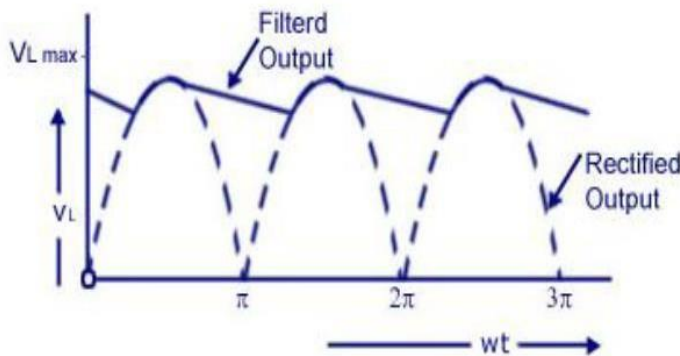
Choke-input filter consists of a choke L connected in series with the rectifier and a capacitor C connected across the load. This is also sometimes called the L-section filter because in this arrangement inductor and capacitor are connected, as an inverted L. In figure only one filter section is shown. But several identical sections are often employed to improve the smoothing action. (The choke L on the input side of the filter readily allows dc to pass but opposes the flow of ac components because its dc resistance is negligibly small but ac inductance is large. Any

fluctuation that remains in the current even after passing through the choke are largely by-passed around the load by the shunt capacitor because X_c is much smaller than R_L . Ripples can be reduced effectively by making X_L greater than X_c at ripple frequency. However, a small ripple still remains in the filtered output and this is considered negligible if it is less than 1%. The rectified and filtered output voltage waveforms from a full-wave rectifier with a choke-input filter are shown in figure.

II-SECTION FILTER:



Circuit Diagram



Rectified and Filtered Output Voltage Waveform
Full-wave Rectifier With capacitor Input Filter

CAPACITOR-INPUT OR PI-FILTER.

Such a filter consists of a shunt capacitor C_1 at the input followed by an L-section filter formed by series inductor L and shunt capacitor C_2 . This is also called the π -filter because the shape of the circuit diagram for this filter appears like Greek letter π (pi). Since the rectifier feeds directly into the capacitor so it is also called *capacitor input filter*.

As the rectified output is fed directly into a capacitor C_1 . Such a filter can be used with a half-wave rectifier (series inductor and L-section filters cannot be used with half-wave rectifiers). Usually electrolytic capacitors are used even though their capacitances are large but they occupy minimum space. Usually both capacitors C_1 and C_2 are enclosed in one metal container. The metal container serves as the common ground for the two capacitors.

A capacitor-input or π -filter is characterized by a high voltage output at low current drains. Such a filter is used, if, for a given transformer, higher voltage than that can be obtained from an L-section filter is required and if low ripple than that can be obtained from a shunt capacitor filter or L-section filter is desired. In this filter, the input capacitor C_1 is selected to offer very low

reactance to the ripple frequency. Hence major part of filtering is accomplished by the input capacitor C_1 . Most of the remaining ripple is removed by the L-section filter consisting of a choke L and capacitor C_2 .)

The action of this filter can *best* be understood by considering the action of L-section filter, formed by L and C_2 , upon the triangular output voltage wave from the input capacitor C_1 . The charging and discharging action of input capacitor C_1 has already been discussed. The output voltage is roughly the same as across input capacitor C_1 less the dc voltage drop in inductor. The ripples contained in this output are reduced further by L-section filter. The output voltage of pi-filter falls off rapidly with the increase in load-current and, therefore, the voltage regulation with this filter is very poor.

SALIENT FEATURES OF L-SECTION AND PI-FILTERS.

1. In pi-filter the dc output voltage is much larger than that can be had from an L-section filter with the same input voltage.
2. In pi-filter ripples are less in comparison to those in shunt capacitor or L-section filter. So smaller valued choke is required in a pi-filter in comparison to that required in L-section filter.
3. In pi-filter, the capacitor is to be charged to the peak value hence the rms current in supply transformer is larger as compared in case of L-section filter.
4. Voltage regulation in case of pi-filter is very poor, as already mentioned. So π -filters are suitable for fixed loads whereas L-section filters can work satisfactorily with varying loads provided a minimum current is maintained.
5. In case of a pi-filter PIV is larger than that in case of an L-section filter.

COMPARISON OF FILTERS

1. A capacitor filter provides V_m volts at less load current. But regulation is poor.
2. An Inductor filter gives high ripple voltage for low load currents. It is used for high load currents
3. L – Section filter gives a ripple factor independent of load current. Voltage Regulation can be improved by use of bleeder resistance
4. Multiple L – Section filter or π filters give much less ripple than the single L – Section Filter.

CLIPPERS AND CLAMPERS

We've been talking about one application of the humble diode – rectification. These simple devices are also powerful tools in other applications. Specifically, this section of our studies looks at signal modification in terms of **clipping** and **clamping**.

CLIPPERS

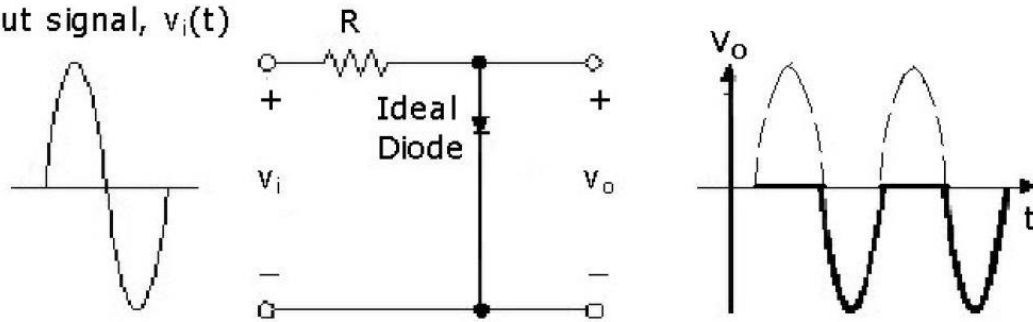
Clipping circuits (also known as limiters, amplitude selectors, or slicers), are used to remove the part of a signal that is above or below some defined reference level. We've already seen an example of a clipper in the half-wave rectifier – that circuit basically clips everything at the

reference level of zero and let only the positive-going (or negative-going) portion of the input waveform through.

To clip to a reference level other than zero, a dc source (shown as a battery in your text) is put in series with the diode. Depending on the direction of the diode and the polarity of the battery, the circuit will either clip the input waveform above or below the reference level (the battery voltage for an ideal diode; i.e., for $V_{on}=0$). This process is illustrated in the four parts of Figure 3:

Without the battery, the output of the circuit below would be the negative portion of the input wave (assuming the bottom node is grounded). When $v_i > 0$, the diode is on (short-circuited), v_i is dropped across R and $v_o=0$. When $v_i < 0$, the diode is off (open-circuited), the voltage across R is zero and $v_o=v_i$. (Don't worry; we won't be doing this for all the circuits!) Anyway, the reference level would be zero.

Input signal, $v_i(t)$



With the battery in the orientation shown in Figure 3a (and below), the diode doesn't turn on until $v_i > V_B$ (If this looks strange, revisit the definition of forward bias). This shifts the reference level up and clips the input at $+V_B$ and passes everything for $v_i < V_B$.

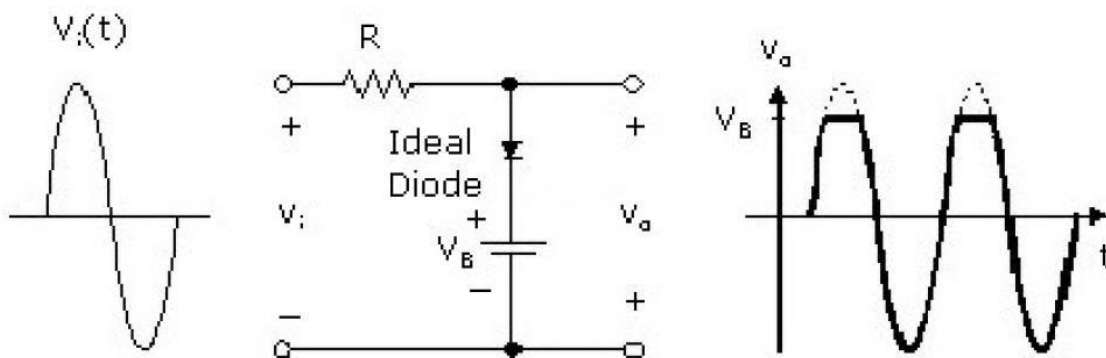


Figure 3.a

Figure 3b has the battery with the same orientation as in part (a), but the diode has been flipped. Without the battery, the positive portion of the input waveform would be passed (i.e., a reference level of zero). With the battery, the diode conducts for $v_i < V_B$. This means that the reference level is shifted to $+V_B$ and only $v_i > V_B$ appears at the output.

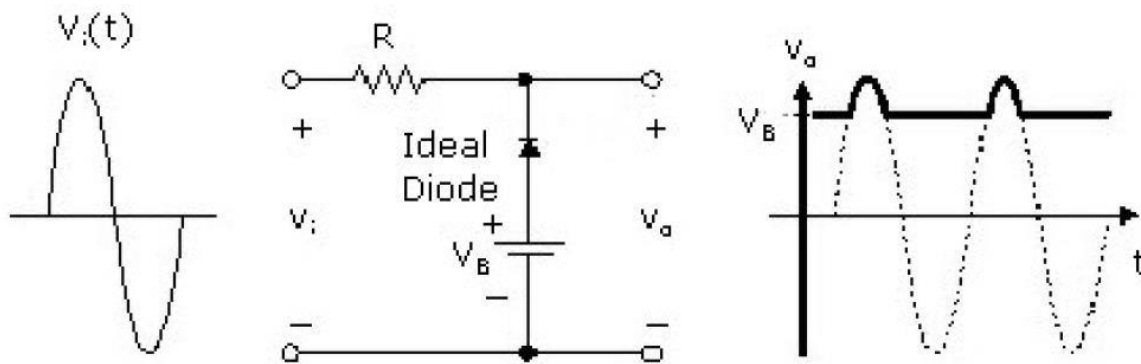


Figure 3.b

Again referencing part (a), the diode is in the original position but the polarities on the battery have been switched in Figure 3c. The discussion follows the same logic as earlier, but now the reference level has been shifted to $-V_B$. The final result is that $v_o = v_i$ for $v_i < -V_B$.

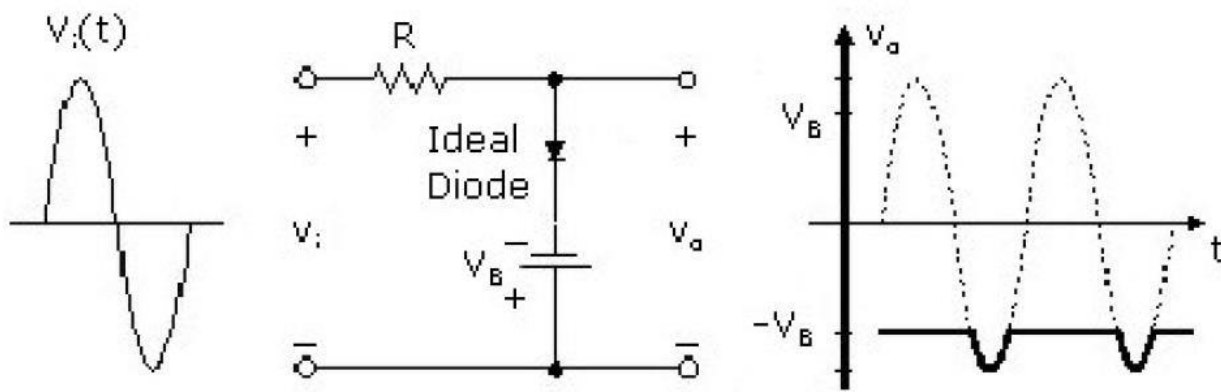


Figure 3.c

Finally, Figure 3d behaves the same as part (b), but the polarity on the battery has been switched, shifting the reference level to $-V_B$. The signal that appears as the output is V_i as long as $V_i > -V_B$.

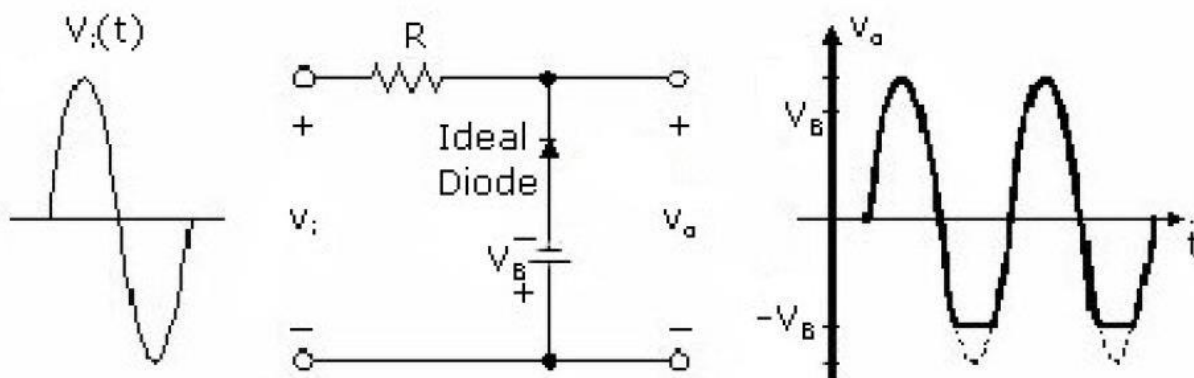
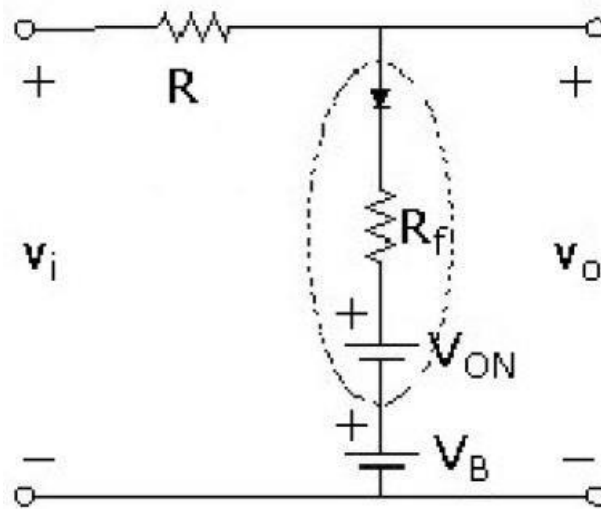


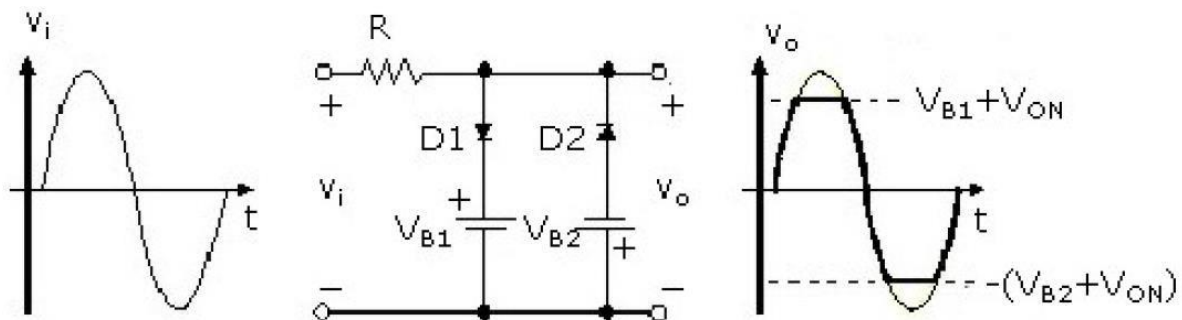
Figure 3.d

To accommodate a practical diode, the turn-on voltage ($V_{ON}=0.7V$ for silicon) and forward resistance (R_f) are included, along with the ideal diode, in the model (as shown in below Figure reproduced to the right). The effective reference level will either have a magnitude of V_B+V_{ON} or V_B-V_{ON} , depending on the relative polarities of the two sources (review combining voltage sources in series if necessary). Including R_f in the diode path creates a voltage divider when the diode is forward biased. The result of this slight drop across R_f (remember that the forward

resistance is generally pretty small), is a slight distortion in the output waveform – it is no longer strictly “limited” or “clipped” to the reference level, as is illustrated in below Figure in your text. The four possible configurations of below Figure are still valid, with the effective reference level ranging in magnitude from 0.7V (if $V_B=0$) on up.



A **parallel-biased clipper** is a circuit that clips the positive and negative-going portions of the input signal simultaneously. This is designed by using two parallel diodes oriented in opposite directions – note that it is very important that the diodes are oppositely oriented (think voltage sources in parallel – a big no-no!). Just as in our previous discussion, the path containing diode D1 will provide the upper limit with reference level $V_{B1}+V_{ON}$ (with the V_{B1} polarity shown) and the path containing D2 will provide the lower limit with reference level $V_{B2}+V_{ON}$ (with the V_{B2} polarity shown). An example of this type of clipper, with the resulting output waveform is shown below (Figure 3.45 of your text, where it looks like they assumed R_f was negligible):

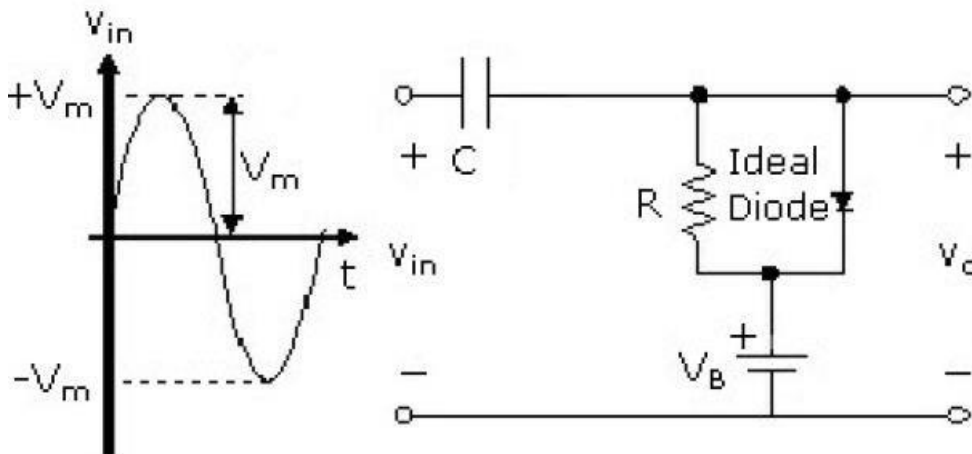


A **series-biased clipper** involves placing a battery in series with the input. The result of this modification is that the input signal is no longer symmetric about the zero axis, but instead shifts by an amount defined by the magnitude and polarity of V_B . In Figure 3.46, the four possible permutations and the resulting output waveforms are shown for an ideal diode. When you look at this figure, keep in mind that the input signal is swinging between +2 and –2 volts and the battery magnitude V_B is 1 volt. This may avoid some confusion when looking at the output waveforms – the $(2V+V_B)$ is just 3 volts and the $-(2V-V_B)$ may be replaced by –1 volt. The series placement of the battery is not changing the input waveform in any way, it is simply affecting when the diode turns on. To include the effects of a **practical diode**, include V_{ON} and R_f in the diode path and crunch the math...

CLAMPERS

Clamping circuits, also known as **dc restorers** or **clamped capacitors**, shift an input signal by an amount defined by an independent voltage source. While clippers limit the part of the input signal that reaches the output according to some reference level(s), the entire input reaches the output in a clamping circuit – it is just shifted so that the maximum (or minimum) value of the input is “clamped” to the independent source.

something that makes sense! Basically, we have a sinusoidal input of magnitude V_m with zero offset (i.e., symmetric signal) fed to the clamper circuit.

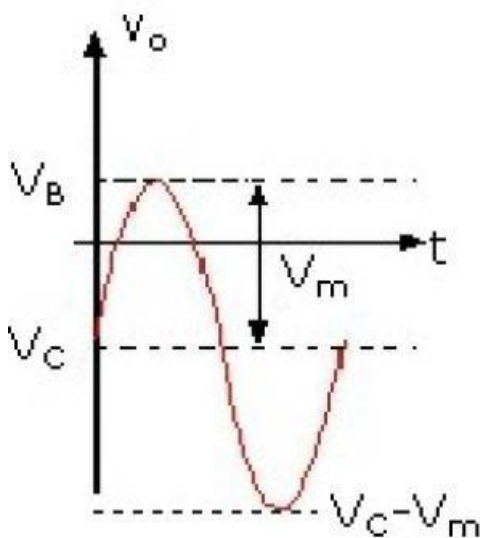


Taking the input by sections to build the output...

- If $V_{in} > V_B$, the diode is on, R is shorted, and the output is V_B .
- When $V_{in} < V_B$, the diode is off, current flows through the resistor, and the capacitor charges to a voltage $V_C = V_B - V_{in}$. The maximum voltage on the capacitor will be related to the maximum swing of the input by: . After fully charging, the capacitor acts like series source in the circuit (with the RC time constant conditions discussed below). After steady state is reached, the output voltage is found by loop analysis as

$$V_o = V_{in} + V_C.$$

In the circuit shown above, $V_C = -V_m + V_B$ (KVL at maximum negative input, $V_{in} = -V_m$) or $V_o = V_{in} - V_m + V_B$. Specifically, for the extreme values of the input signal: if $V_{in} = +V_m$, $V_o = V_B$ and if $V_{in} = -V_m$, $V_o = -V_m - V_m + V_B = -V_m + V_C = V_C - V_m$ (Whew! I went through all that to get something that looked like it belonged to the figure to the right!).

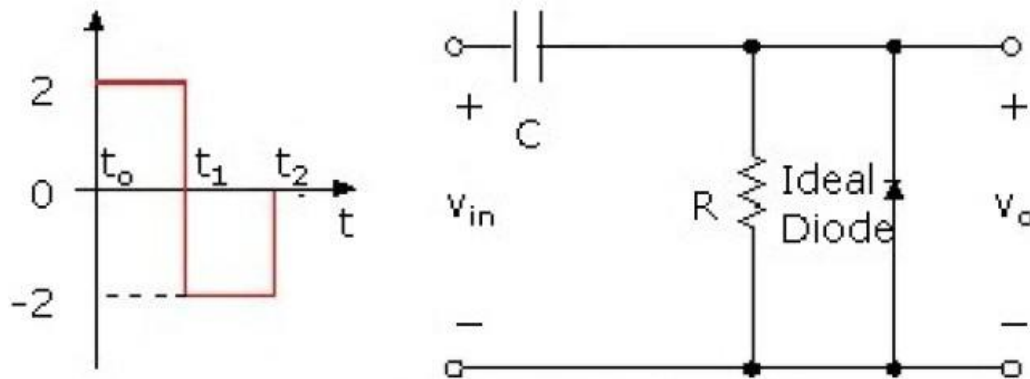


Keep in mind that the above analysis was for the diode orientation and V_B polarity shown. If the diode had been flipped, the minimum rather than the maximum of the input would have been clamped to V_B .

A clamping circuit has to have an independent source, a diode, a resistor and a capacitor. To keep a

constant voltage on the capacitor over the period of the input, the RC time constant must be large. A **design rule of thumb** is to make the RC time constant at least five times the half-period of the input signal, which results in approximately an 18% error over a half-period due to capacitor discharge. If this error is too large, the RC time constant may be increased but, as with everything in design, there comes a point where factors such as size and power dissipation may make any further improvements impractical.

Following our discussion above for Figure. Below Figure illustrates a circuit that will clamp a square wave to zero ($V_B=0$).



$$V_C = +V_m = 2V$$

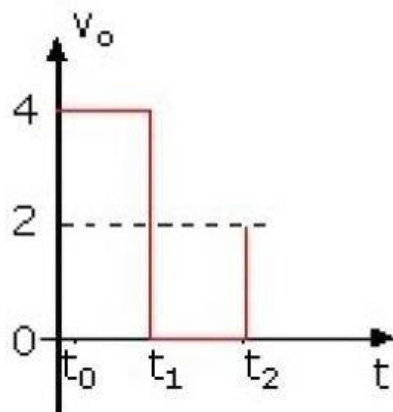
(because of diode orientation and $V_B=0$)

$$v_o = v_{in} + V_C$$

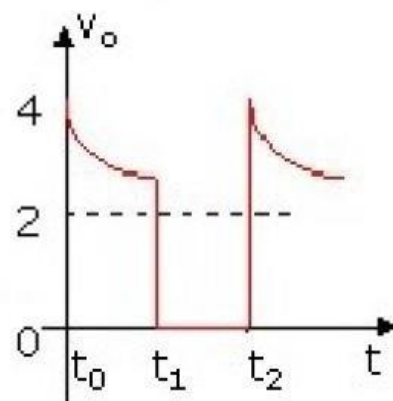
$$v_o = 2 V_m = 4V \text{ for } v_{in} = V_m$$

$$v_o = 0 \text{ for } v_{in} = -V_m$$

Parts (b) and (c) of below figure, shown below, demonstrate the output of the clamper with a long time constant and the distortion introduced by the capacitor discharging for a short time constant. It is noted in your book that this square wave may be considered a worst-case situation since it places the greatest demands on a clamping circuit due to the instantaneous changes in the waveform. (Remember from your discussions of harmonics and signal composition, that instantaneous change requires infinite frequencies.)



(b) Long time constant



(c) Short time constant

UNIT II

BIPOLAR JUNCTION TRANSISTOR

INTRODUCTION

A bipolar junction transistor (BJT) is a three terminal device in which operation depends on the interaction of both majority and minority carriers and hence the name bipolar. The BJT is analogous to vacuum triode and is comparatively smaller in size. It is used in amplifier and oscillator circuits, and as a switch in digital circuits. It has wide applications in computers, satellites and other modern communication systems.

After knowing the details about a single PN junction, or simply a diode, let us try to go for the two PN junction connection. If another P-type material or N-type material is added to a single PN junction, another junction will be formed. Such a formation is simply called as a **Transistor**.

A **Transistor** is a three terminal semiconductor device that regulates current or voltage flow and acts as a switch or gate for signals.

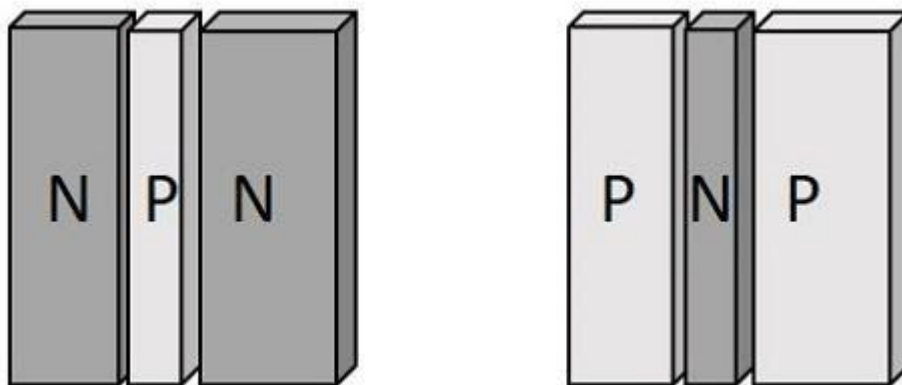
Uses of a transistor

- A transistor acts as **an Amplifier**, where the signal strength has to be increased.
- A transistor also acts as a **switch** to choose between available options.
- It also **regulates** the incoming **current and voltage** of the signals.

CONSTRUCTIONAL DETAILS OF A TRANSISTOR

The Transistor is a three terminal solid state device which is formed by connecting two diodes back to back. Hence it has got **two PN junctions**. Three terminals are drawn out of the three semiconductor materials present in it. This type of connection offers two types of transistors. They are **PNP** and **NPN** which means an N-type material between two Ptypes and the other is a P-type material between two N-types respectively.

The following illustration shows the basic construction of transistors



The three terminals drawn from the transistor indicate **Emitter**, **Base** and **Collector** terminals. They have their functionality as discussed below.

Emitter

- The left-hand side of the above shown structure can be understood as **Emitter**.
- This has a **moderate size** and is **heavily doped** as its main function is to **supply** a number of **majority carriers**, i.e. either electrons or holes.
- As this emits electrons, it is called as an Emitter.
- This is simply indicated with the letter **E**.

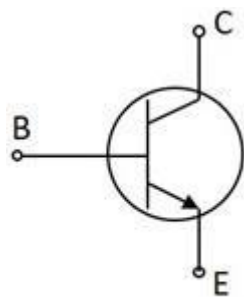
Base

- The middle material in the above figure is the **Base**.
- This is **thin** and **lightly doped**.
- Its main function is to **pass** the majority carriers from the emitter to the collector.
- This is indicated by the letter **B**.

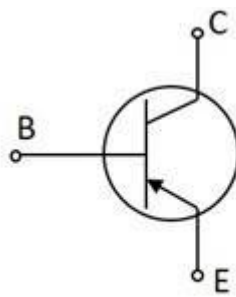
Collector

- The right side material in the above figure can be understood as a **Collector**.
- Its name implies its function of **collecting the carriers**.
- This is a **bit larger** in size than emitter and base. It is **moderately doped**.
- This is indicated by the letter **C**.

The symbols of PNP and NPN transistors are as shown below.



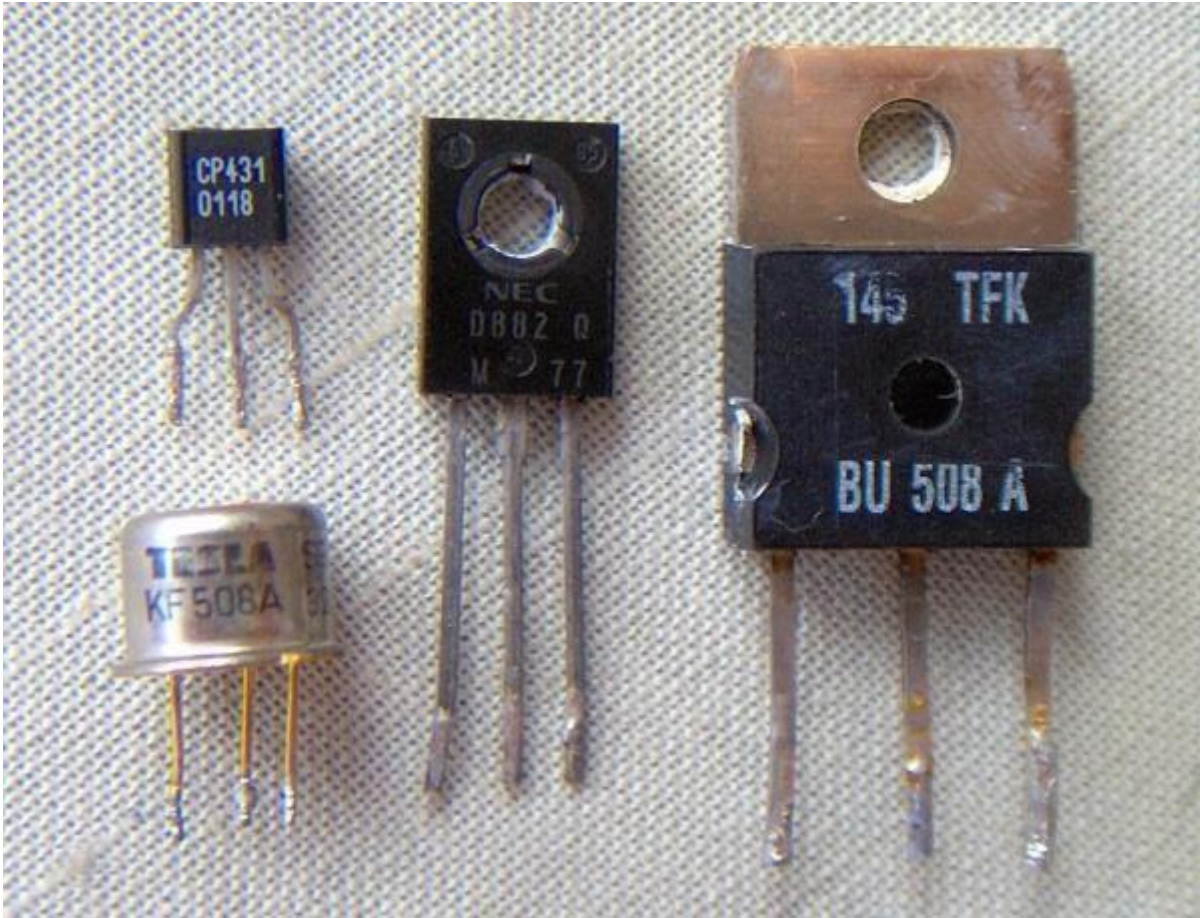
NPN transistor



PNP transistor

The **arrow-head** in the above figures indicated the **emitter** of a transistor. As the collector of a transistor has to dissipate much greater power, it is made large. Due to the specific functions of emitter and collector, they are **not interchangeable**. Hence the terminals are always to be kept in mind while using a transistor.

In a Practical transistor, there is a notch present near the emitter lead for identification. The PNP and NPN transistors can be differentiated using a Multimeter. The following image shows how different practical transistors look like.

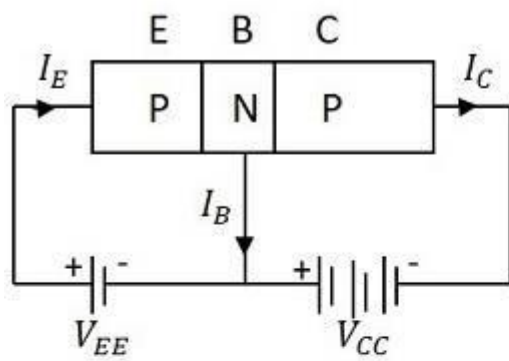


We have so far discussed the constructional details of a transistor, but to understand the operation of a transistor, first we need to know about the biasing.

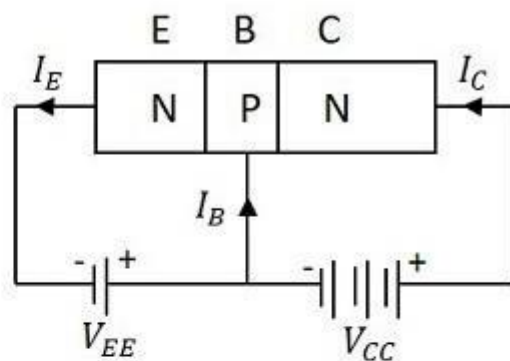
Transistor Biasing

As we know that a transistor is a combination of two diodes, we have two junctions here. As one junction is between the emitter and base, that is called as **Emitter-Base junction** and likewise, the other is **Collector-Base junction**.

Biasing is controlling the operation of the circuit by providing power supply. The function of both the PN junctions is controlled by providing bias to the circuit through some dc supply. The figure below shows how a transistor is biased.



P-N-P Transistor biasing



N-P-N Transistor biasing

By having a look at the above figure, it is understood that

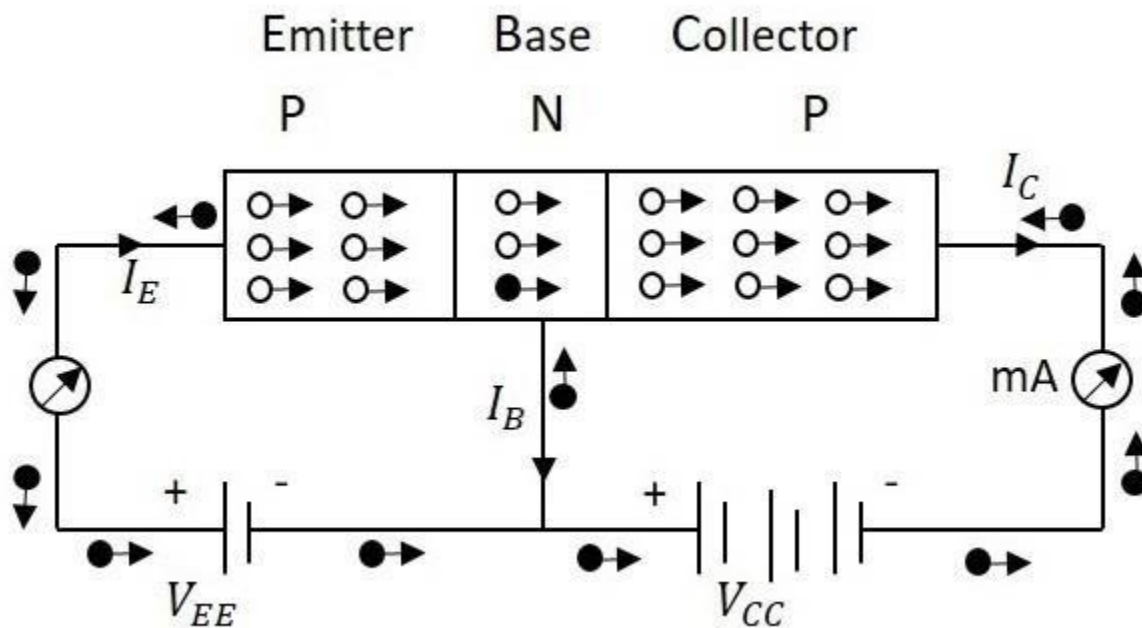
- The N-type material is provided negative supply and P-type material is given positive supply to make the circuit **Forward bias**.
- The N-type material is provided positive supply and P-type material is given negative supply to make the circuit **Reverse bias**.

By applying the power, the **emitter base junction** is always **forward biased** as the emitter resistance is very small. The **collector base junction** is **reverse biased** and its resistance is a bit higher. A small forward bias is sufficient at the emitter junction whereas a high reverse bias has to be applied at the collector junction.

The direction of current indicated in the circuits above, also called as the **Conventional Current**, is the movement of hole current which is **opposite to the electron current**.

Operation of PNP Transistor

The operation of a PNP transistor can be explained by having a look at the following figure, in which emitter-base junction is forward biased and collector-base junction is reverse biased.



The voltage V_{EE} provides a positive potential at the emitter which repels the holes in the P-type material and these holes cross the emitter-base junction, to reach the base region. There a very low percent of holes re-combine with free electrons of N-region. This provides very low current which constitutes the base current I_B . The remaining holes cross the collector-base junction, to constitute collector current I_C , which is the hole current.

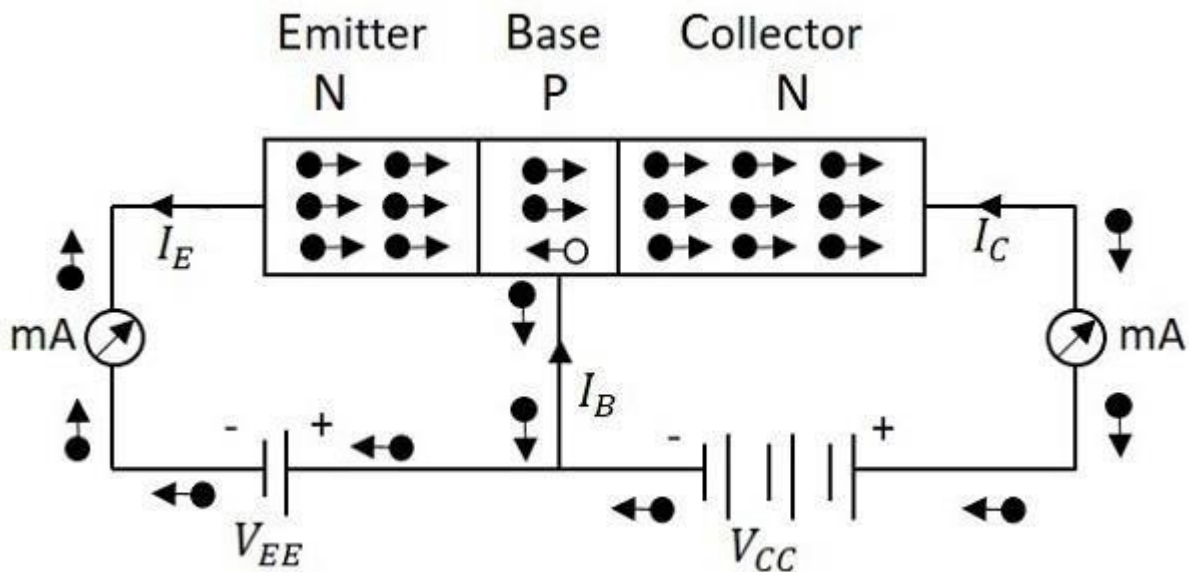
As a hole reaches the collector terminal, an electron from the battery negative terminal fills the space in the collector. This flow slowly increases and the electron minority current flows through the emitter, where each electron entering the positive terminal of V_{EE} , is replaced by a hole by moving towards the emitter junction. This constitutes emitter current I_E .

Hence we can understand that –

- The conduction in a PNP transistor takes place through holes.
- The collector current is slightly less than the emitter current.
- The increase or decrease in the emitter current affects the collector current.

Operation of NPN Transistor

The operation of an NPN transistor can be explained by having a look at the following figure, in which emitter-base junction is forward biased and collector-base junction is reverse biased.



The voltage V_{EE} provides a negative potential at the emitter which repels the electrons in the N-type material and these electrons cross the emitter-base junction, to reach the base region. There, a very low percent of electrons re-combine with free holes of P-region. This provides very low current which constitutes the base current I_B . The remaining holes cross the collector-base junction, to constitute the collector current I_C .

As an electron reaches out of the collector terminal, and enters the positive terminal of the battery, an electron from the negative terminal of the battery V_{EE} enters the emitter region. This flow slowly increases and the electron current flows through the transistor.

Hence we can understand that –

- The conduction in a NPN transistor takes place through electrons.
- The collector current is higher than the emitter current.
- The increase or decrease in the emitter current affects the collector current.

Advantages of Transistors

There are many advantages of using a transistor, such as –

- High voltage gain.
- Lower supply voltage is sufficient.
- Most suitable for low power applications.
- Smaller and lighter in weight.
- Mechanically stronger than vacuum tubes.
- No external heating required like vacuum tubes.
- Very suitable to integrate with resistors and diodes to produce ICs.

There are few disadvantages such as they cannot be used for high power applications due to lower power dissipation. They have lower input impedance and they are temperature dependent.

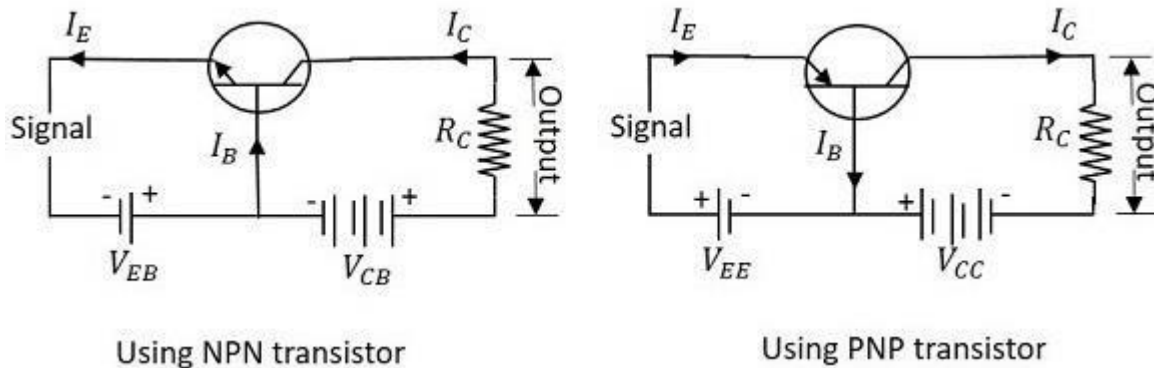
BIPOLAR TRANSISTOR CONFIGURATIONS

Any transistor has three terminals, the **emitter**, the **base**, and the **collector**. Using these 3 terminals the transistor can be connected in a circuit with one terminal common to both input and output in three different possible configurations.

The three types of configurations are **Common Base**, **Common Emitter** and **Common Collector** configurations. In every configuration, the emitter junction is forward biased and the collector junction is reverse biased.

Common Base (CB) Configuration

The name itself implies that the **Base** terminal is taken as common terminal for both input and output of the transistor. The common base connection for both NPN and PNP transistors is as shown in the following figure.



For the sake of understanding, let us consider NPN transistor in CB configuration. When the emitter voltage is applied, as it is forward biased, the electrons from the negative terminal repel the emitter electrons and current flows through the emitter and base to the collector to contribute collector current. The collector voltage V_{CB} is kept constant throughout this.

In the CB configuration, the input current is the emitter current I_E and the output current is the collector current I_C .

Current Amplification Factor (α)

The ratio of change in collector current (ΔI_C) to the change in emitter current (ΔI_E) when collector voltage V_{CB} is kept constant, is called as **Current amplification factor**. It is denoted by α .

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

Expression for Collector current

With the above idea, let us try to draw some expression for collector current.

Along with the emitter current flowing, there is some amount of base current I_B which flows through the base terminal due to electron hole recombination. As collector-base junction is reverse biased, there is another current which is flown due to minority charge carriers. This is the leakage current which can be understood as $I_{leakage}$. This is due to minority charge carriers and hence very small.

The emitter current that reaches the collector terminal is

$$\alpha I_E$$

Total collector current

$$I_C = \alpha I_E + I_{leakage}$$

If the emitter-base voltage $V_{EB} = 0$, even then, there flows a small leakage current, which can be termed as I_{CBO} (collector-base current with output open).

The collector current therefore can be expressed as

$$I_C = \alpha I_E + I_{CBO}$$

$$I_E = I_C + I_B$$

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

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

$$I_C = [\alpha / (1 - \alpha)] I_B + I_{CBO}$$

$$I_C = [\alpha / (1 - \alpha)] I_B + [1 / (1 - \alpha)] I_{CBO}$$

Hence the above derived is the expression for collector current. The value of collector current depends on base current and leakage current along with the current amplification factor of that transistor in use.

Characteristics of CB configuration

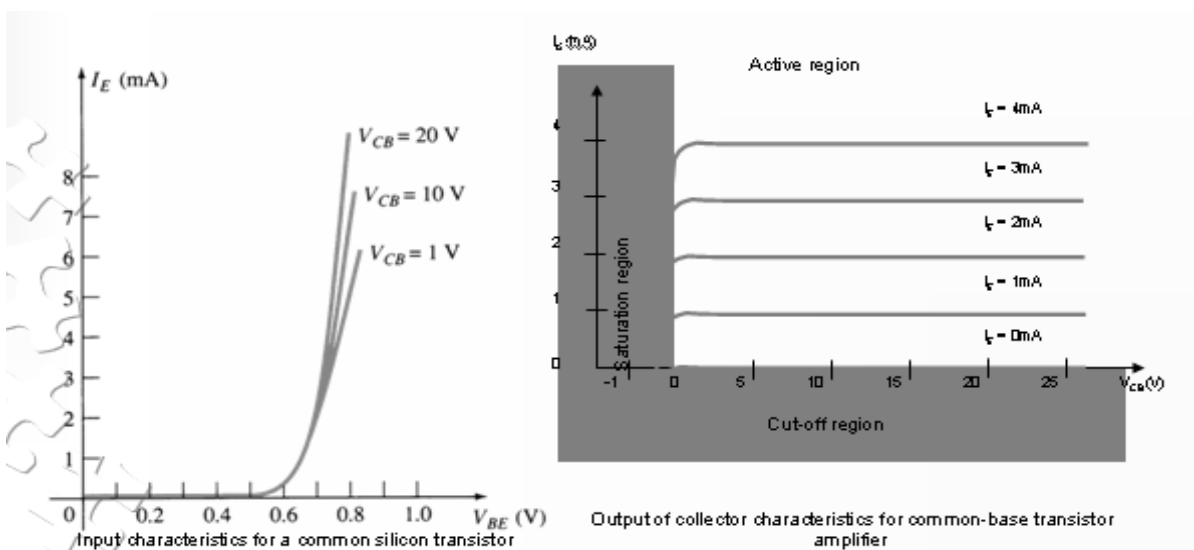
- This configuration provides voltage gain but no current gain.
- Being V_{CB} constant, with a small increase in the Emitter-base voltage V_{EB} , Emitter current I_E gets increased.
- Emitter Current I_E is independent of Collector voltage V_{CB} .
- Collector Voltage V_{CB} can affect the collector current I_C only at low voltages, when V_{EB} is kept constant.
- The input resistance R_i is the ratio of change in emitter-base voltage (ΔV_{EB}) to the change in emitter current (ΔI_E) at constant collector base voltage V_{CB} .

$$R_i = \Delta V_{EB} / \Delta I_E \text{ at constant } V_{CB}$$

- As the input resistance is of very low value, a small value of V_{EB} is enough to produce a large current flow of emitter current I_E .
- The output resistance R_o is the ratio of change in the collector base voltage (ΔV_{CB}) to the change in collector current (ΔI_C) at constant emitter current I_E .

$$R_o = \Delta V_{CB} / \Delta I_C \text{ at constant } I_E$$

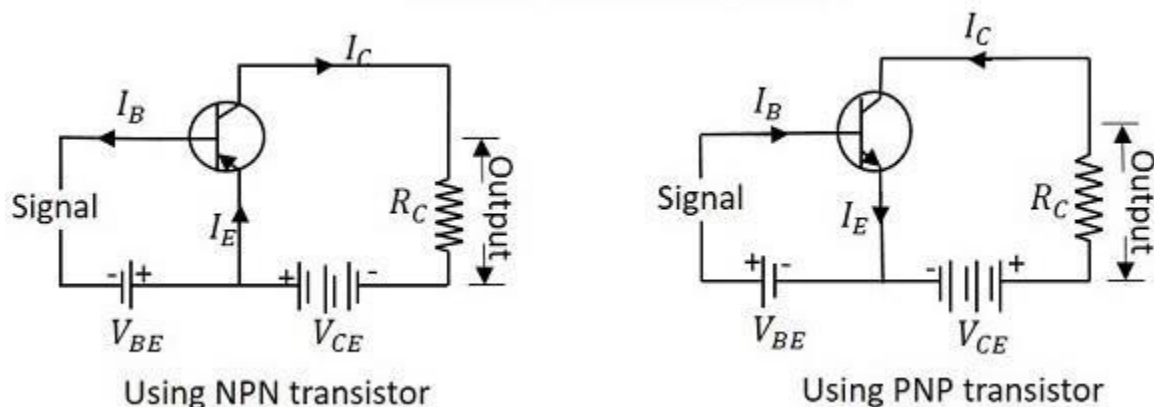
- As the output resistance is of very high value, a large change in V_{CB} produces a very little change in collector current I_C .
- This Configuration provides good stability against increase in temperature.
- The CB configuration is used for high frequency applications



Active region	Saturation region	Cut-off region
• I_E increased, I_C increased • BE junction forward bias and CB junction reverse bias • Refer to the graf, $I_C \approx I_E$ • I_C not depends on V_{CB} • Suitable region for the transistor working as amplifier	• BE and CB junction is forward bias • Small changes in V_{CB} will cause big different to I_C • The allocation for this region is to the left of $V_{CB} = 0$ V.	• Region below the line of $I_E = 0$ A • BE and CB is reverse bias • no current flow at collector, only leakage current

COMMON-EMITTER CONFIGURATION

The name itself implies that the **Emitter** terminal is taken as common terminal for both input and output of the transistor. The common emitter connection for both NPN and PNP transistors is as shown in the following figure.



Just as in CB configuration, the emitter junction is forward biased and the collector junction is reverse biased. The flow of electrons is controlled in the same manner. The input current is the base current I_B and the output current is the collector current I_C here.

Base Current Amplification factor (β)

The ratio of change in collector current (ΔI_C) to the change in base current (ΔI_B) is known as **Base Current Amplification Factor**. It is denoted by β .

$$\beta = \Delta I_C / \Delta I_B$$

Relation between β and α

Let us try to derive the relation between base current amplification factor and emitter current amplification factor.

$$\begin{aligned}\beta &= \Delta I_C / \Delta I_B \\ \alpha &= \Delta I_C / \Delta I_E \\ I_E &= I_B + I_C \\ \Delta I_E &= \Delta I_B + \Delta I_C\end{aligned}$$

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

We can write

$$\beta = \Delta I_C / (\Delta I_E - \Delta I_C)$$

Dividing by ΔI_E

$$\beta = (\Delta I_C / \Delta I_E) / [(\Delta I_E / \Delta I_E) - (\Delta I_C / \Delta I_E)]$$

We have

$$\alpha = \Delta I_C / \Delta I_E$$

Therefore,

$$\beta = \alpha / (1 - \alpha)$$

From the above equation, it is evident that, as α approaches 1, β reaches infinity.

Hence, **the current gain in Common Emitter connection is very high**. This is the reason this circuit connection is mostly used in all transistor applications.

Expression for Collector Current

In the Common Emitter configuration, I_B is the input current and I_C is the output current.

We know

$$I_E = I_B + I_C$$

And

$$\begin{aligned} I_C &= \alpha I_E + I_{CBO} \\ &= \alpha (I_B + I_C) + I_{CBO} \\ I_C (1 - \alpha) &= \alpha I_B + I_{CBO} \\ I_C &= [\alpha / (1 - \alpha)] I_B + [1 / (1 - \alpha)] I_{CBO} \end{aligned}$$

If base circuit is open, i.e. if $I_B = 0$,

The collector emitter current with base open is I_{CEO}

$$I_{CEO} = [1 / (1 - \alpha)] I_{CBO}$$

Substituting the value of this in the previous equation, we get

$$\begin{aligned} I_C &= [\alpha / (1 - \alpha)] I_B + I_{CEO} \\ I_C &= \beta I_B + I_{CEO} \end{aligned}$$

Hence the equation for collector current is obtained.

Knee Voltage

In CE configuration, by keeping the base current I_B constant, if V_{CE} is varied, I_C increases nearly to 1V of V_{CE} and stays constant thereafter. This value of V_{CE} up to which collector current I_C changes with V_{CE} is called the **Knee Voltage**. The transistors while operating in CE configuration, they are operated above this knee voltage.

Characteristics of CE Configuration

- This configuration provides good current gain and voltage gain.
- Keeping V_{CE} constant, with a small increase in V_{BE} the base current I_B increases rapidly than in CB configurations.

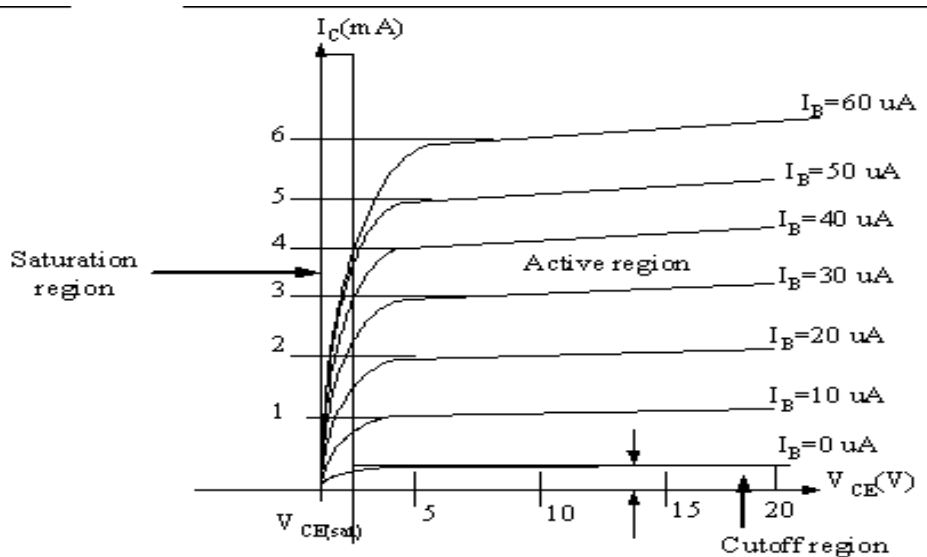
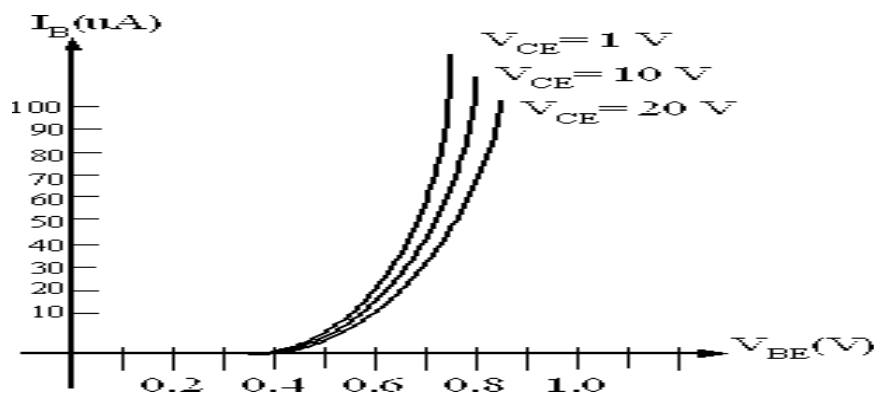
- For any value of V_{CE} above knee voltage, I_C is approximately equal to βI_B .
- The input resistance R_i is the ratio of change in base emitter voltage (ΔV_{BE}) to the change in base current (ΔI_B) at constant collector emitter voltage V_{CE} .

$$R_i = \Delta V_{BE} / \Delta I_B \text{ at constant } V_{CE}$$

- As the input resistance is of very low value, a small value of V_{BE} is enough to produce a large current flow of base current I_B .
- The output resistance R_o is the ratio of change in collector emitter voltage (ΔV_{CE}) to the change in collector current (ΔI_C) at constant I_B .

$$R_o = \Delta V_{CE} / \Delta I_C \text{ at constant } I_B$$

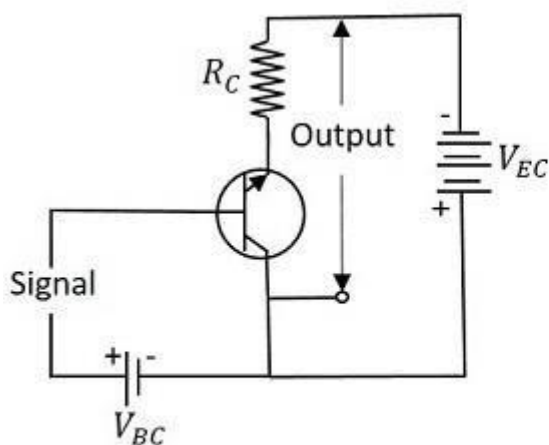
- As the output resistance of CE circuit is less than that of CB circuit.
- This configuration is usually used for bias stabilization methods and audio frequency applications



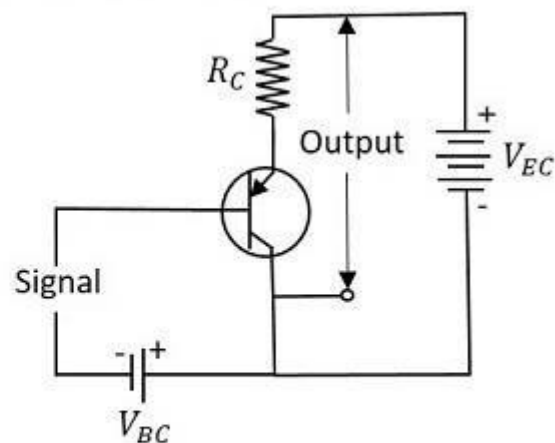
Active region	Saturation region	Cut-off region
• B-E junction is forward bias • C-B junction is reverse bias • can be employed for voltage, current and power amplification	• B-E and C-B junction is forward bias, thus the values of I_B and I_C is too big. • The value of V_{CE} is so small. • Suitable region when the transistor as a logic switch. • NOT and avoid this region when the transistor as an amplifier.	• region below $I_B=0\mu A$ is to be avoided if an undistorted o/p signal is required • B-E junction and C-B junction is reverse bias • $I_B=0$, I_C not zero, during this condition $I_C=I_{CEO}$ where is this current flow when B-E is reverse bias.

COMMON – COLLECTOR CONFIGURATION

The name itself implies that the **Collector** terminal is taken as common terminal for both input and output of the transistor. The common collector connection for both NPN and PNP transistors is as shown in the following figure.



Using NPN transistor



Using PNP transistor

Just as in CB and CE configurations, the emitter junction is forward biased and the collector junction is reverse biased. The flow of electrons is controlled in the same manner. The input current is the base current I_B and the output current is the emitter current I_E here.

Current Amplification Factor (γ)

The ratio of change in emitter current (ΔI_E) to the change in base current (ΔI_B) is known as **Current Amplification factor** in common collector (CC) configuration. It is denoted by γ .

$$\gamma = \Delta I_E / \Delta I_B$$

- The current gain in CC configuration is same as in CE configuration.
- The voltage gain in CC configuration is always less than 1.

Relation between γ and α

Let us try to draw some relation between γ and α

$$\gamma = \Delta I_E / \Delta I_B$$

$$\alpha = \Delta I_C / \Delta I_E$$

$$I_E = I_B + I_C$$

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

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

Substituting the value of I_B , we get

$$\gamma = \Delta I_E / (\Delta I_E - \Delta I_C)$$

Dividing by ΔI_E

$$\begin{aligned} \gamma &= (\Delta I_E / \Delta I_E) / [(\Delta I_E / \Delta I_E) - (\Delta I_C / \Delta I_E)] \\ &= 1 / (1 - \alpha) \\ \gamma &= 1 / (1 - \alpha) \end{aligned}$$

Expression for collector current

We know

$$\begin{aligned} I_C &= \alpha I_E + I_{CBO} \\ I_E &= I_B + I_C = I_B + (\alpha I_E + I_{CBO}) \\ I_E(1 - \alpha) &= I_B + I_{CBO} \\ I_E &= [I_B / (1 - \alpha)] + [I_{CBO} / (1 - \alpha)] \\ I_C &\cong I_E = (\beta + 1)I_B + (\beta + 1)I_{CBO} \end{aligned}$$

The above is the expression for collector current.

Characteristics of CC Configuration

- This configuration provides current gain but no voltage gain.
- In CC configuration, the input resistance is high and the output resistance is low.
- The voltage gain provided by this circuit is less than 1.
- The sum of collector current and base current equals emitter current.
- The input and output signals are in phase.
- This configuration works as non-inverting amplifier output.
- This circuit is mostly used for impedance matching. That means, to drive a low impedance load from a high impedance source.

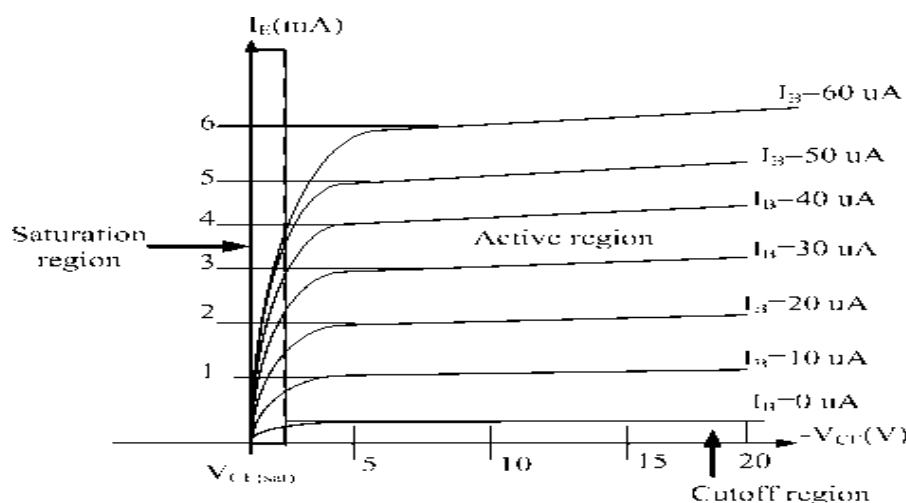


Fig 4.9 : Output characteristic in CC configuration for npn transistor

LIMITS OF OPERATION

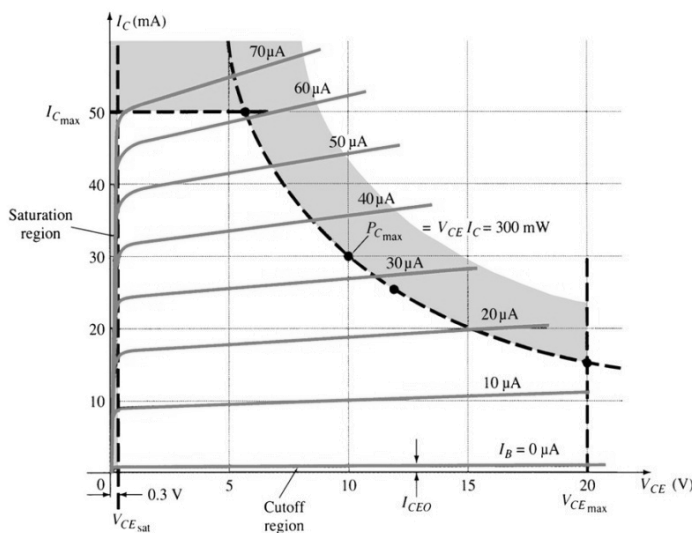
Many BJT transistors used as an amplifier. Thus it is important to notice the limits of operations. At least 3 maximum values is mentioned in data sheet.

There are:

- Maximum power dissipation at collector: P_{Cmax} or P_D
- Maximum collector-emitter voltage: V_{CEmax} sometimes named as $V_{BR(CEO)}$ or V_{CEO} .
- Maximum collector current: I_{Cmax}

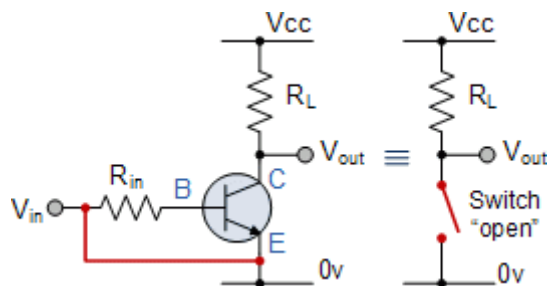
There are few rules that need to be followed for BJT transistor used as an amplifier. The rules are: transistor needs to be operated in active region!

$$I_C < I_{Cmax} ; P_C < P_{Cmax}$$



TRANSISTOR AS A SWITCH

Transistor switches can be used to switch a low voltage DC device (e.g. LED's) ON or OFF by using a transistor in its saturated or cut-off state



When used as an AC signal amplifier, the transistors Base biasing voltage is applied in such a way that it always operates within its “active” region, that is the linear part of the output characteristics curves are used.

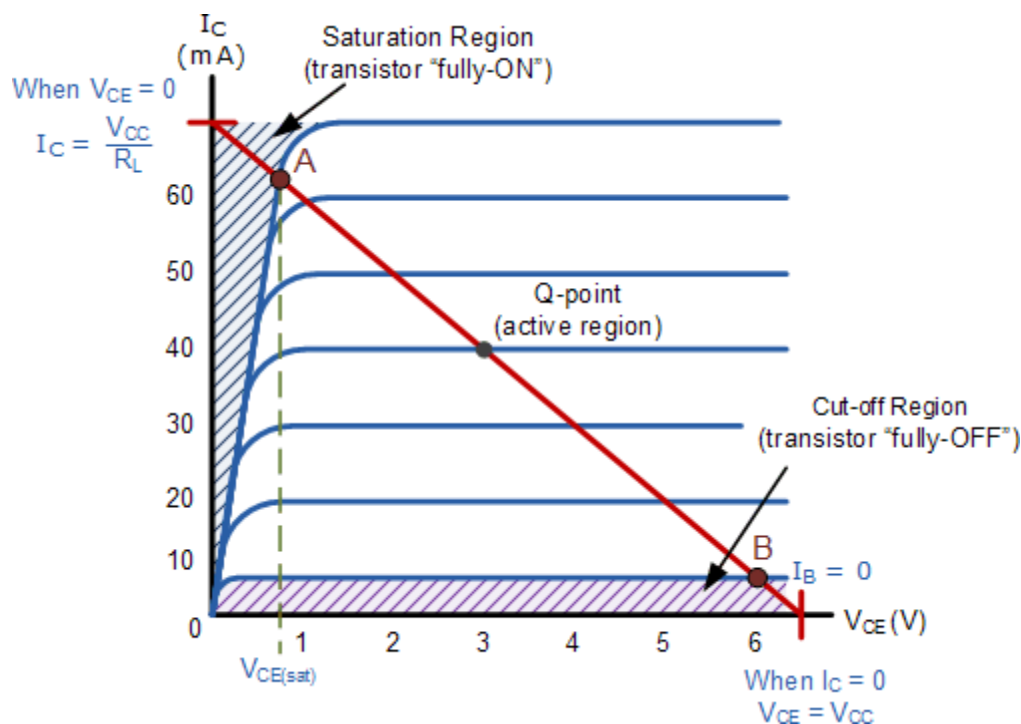
However, both the NPN & PNP type bipolar transistors can be made to operate as “ON/OFF” type solid state switch by biasing the transistors Base terminal differently to that for a signal amplifier.

Solid state switches are one of the main applications for the use of transistor to switch a DC output “ON” or “OFF”. Some output devices, such as LED’s only require a few milliamps at logic level DC voltages and can therefore be driven directly by the output of a logic gate. However, high power devices such as motors, solenoids or lamps, often require more power than that supplied by an ordinary logic gate so transistor switches are used.

If the circuit uses the **Bipolar Transistor as a Switch**, then the biasing of the transistor, either NPN or PNP is arranged to operate the transistor at both sides of the “ I-V ” characteristics curves we have seen previously.

The areas of operation for a transistor switch are known as the **Saturation Region** and the **Cut-off Region**. This means then that we can ignore the operating Q-point biasing and voltage divider circuitry required for amplification, and use the transistor as a switch by driving it back and forth between its “fully-OFF” (cut-off) and “fully-ON” (saturation) regions as shown below.

Operating Regions

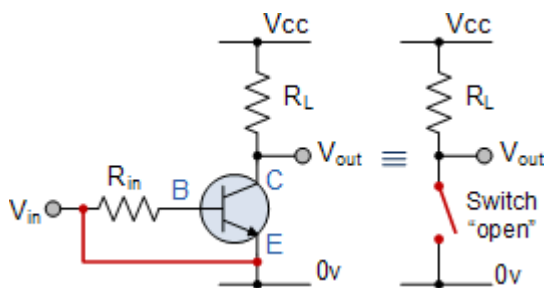


The pink shaded area at the bottom of the curves represents the “Cut-off” region while the blue area to the left represents the “Saturation” region of the transistor. Both these transistor regions are defined as:

1. Cut-off Region

Here the operating conditions of the transistor are zero input base current (I_B), zero output collector current (I_C) and maximum collector voltage (V_{CE}) which results in a large depletion layer and no current flowing through the device. Therefore the transistor is switched “Fully-OFF”.

Cut-off Characteristics



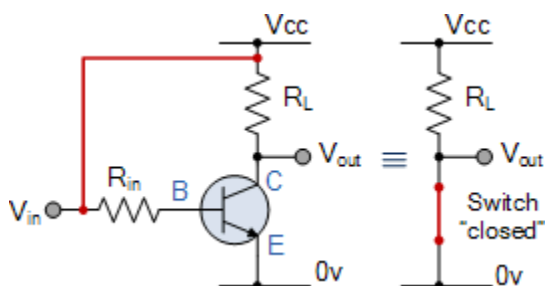
- The input and Base are grounded (0v)
- Base-Emitter voltage $V_{BE} < 0.7v$
- Base-Emitter junction is reverse biased
- Base-Collector junction is reverse biased
- Transistor is “fully-OFF” (Cut-off region)
- No Collector current flows ($I_C = 0$)
- $V_{OUT} = V_{CE} = V_{CC} = "1"$
- Transistor operates as an “open switch”

Then we can define the “cut-off region” or “OFF mode” when using a bipolar transistor as a switch as being, both junctions reverse biased, $V_B < 0.7v$ and $I_C = 0$. For a PNP transistor, the Emitter potential must be negative with respect to the Base.

2. Saturation Region

Here the transistor will be biased so that the maximum amount of base current is applied, resulting in maximum collector current resulting in the minimum collector emitter voltage drop which results in the depletion layer being as small as possible and maximum current flowing through the transistor. Therefore the transistor is switched “Fully-ON”.

Saturation Characteristics



- The input and Base are connected to V_{CC}
- Base-Emitter voltage $V_{BE} > 0.7v$
- Base-Emitter junction is forward biased
- Base-Collector junction is forward biased
- Transistor is “fully-ON” (saturation region)
- Max Collector current flows ($I_C = V_{CC}/R_L$)
- $V_{CE} = 0$ (ideal saturation)
- $V_{OUT} = V_{CE} = "0"$

- Transistor operates as a “closed switch”

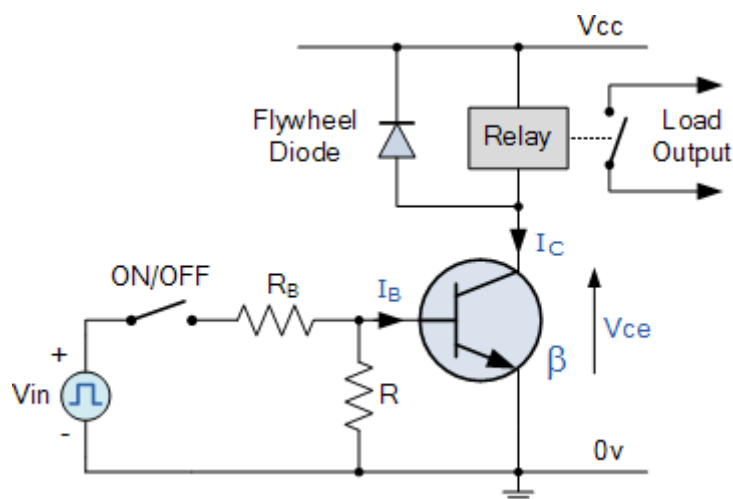
Then we can define the “saturation region” or “ON mode” when using a bipolar transistor as a switch as being, both junctions forward biased, $V_B > 0.7\text{V}$ and $I_C = \text{Maximum}$. For a PNP transistor, the Emitter potential must be positive with respect to the Base.

Then the transistor operates as a “single-pole single-throw” (SPST) solid state switch. With a zero signal applied to the Base of the transistor it turns “OFF” acting like an open switch and zero collector current flows. With a positive signal applied to the Base of the transistor it turns “ON” acting like a closed switch and maximum circuit current flows through the device.

The simplest way to switch moderate to high amounts of power is to use the transistor with an open-collector output and the transistors Emitter terminal connected directly to ground. When used in this way, the transistors open collector output can thus “sink” an externally supplied voltage to ground thereby controlling any connected load.

An example of an NPN Transistor as a switch being used to operate a relay is given below. With inductive loads such as relays or solenoids a flywheel diode is placed across the load to dissipate the back EMF generated by the inductive load when the transistor switches “OFF” and so protect the transistor from damage. If the load is of a very high current or voltage nature, such as motors, heaters etc, then the load current can be controlled via a suitable relay as shown.

Basic NPN Transistor Switching Circuit



The circuit resembles that of the *Common Emitter* circuit we looked at in the previous tutorials. The difference this time is that to operate the transistor as a switch the transistor needs to be turned either fully “OFF” (cut-off) or fully “ON” (saturated). An ideal transistor switch would have infinite circuit resistance between the Collector and Emitter when turned “fully-OFF” resulting in zero current flowing through it and zero resistance between the Collector and Emitter when turned “fully-ON”, resulting in maximum current flow.

In practice when the transistor is turned “OFF”, small leakage currents flow through the transistor and when fully “ON” the device has a low resistance value causing a small saturation voltage (V_{CE}) across it. Even though the transistor is not a perfect switch, in both the cut-off and saturation regions the power dissipated by the transistor is at its minimum.

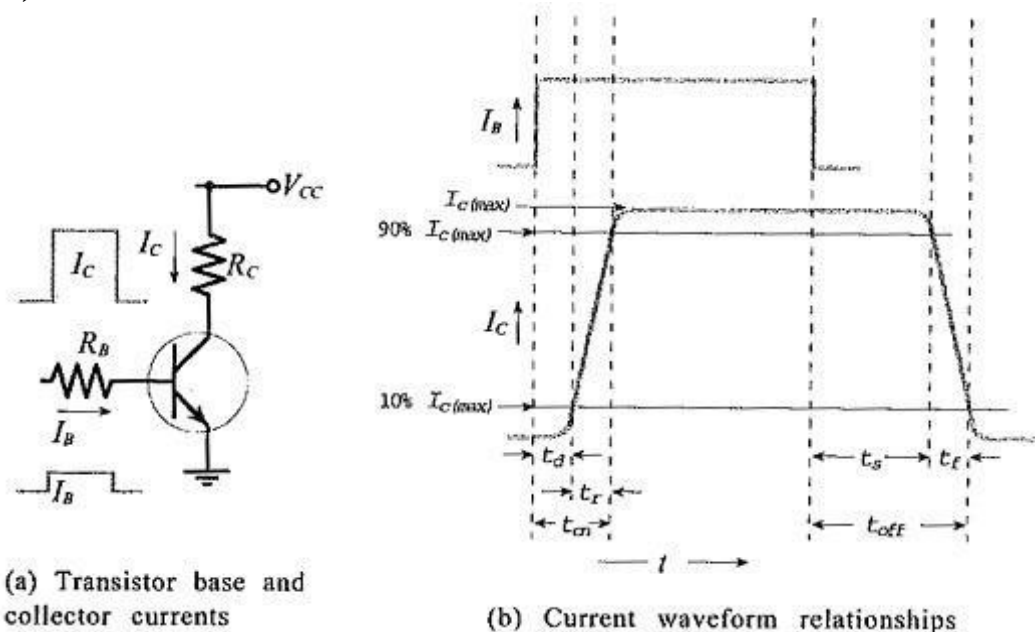
In order for the Base current to flow, the Base input terminal must be made more positive than the Emitter by increasing it above the 0.7 volts needed for a silicon device by varying this Base-

Emitter voltage V_{BE} , the Base current is also altered and which in turn controls the amount of Collector current flowing through the transistor as previously discussed.

When maximum Collector current flows the transistor is said to be Saturated. The value of the Base resistor determines how much input voltage is required and corresponding Base current to switch the transistor fully “ON”.

TRANSISTOR SWITCHING TIMES

For Transistor Switching Times, the switching speed of the device can be an important quantity. Consider the circuit in Fig. When the base input current is applied, the transistor does not switch on immediately. Like frequency response, the Transistor Switching Times is affected by junction capacitance and the transit time of electrons across the junctions. The time between the application of the input pulse and commencement of collector current flow is termed the **delay time** (t_d), [see Fig. Even when the transistor begins to switch on, a finite time elapses before I_C reaches its maximum level. This is known as the **rise time** (t_r). The rise time is specified as the time required for I_C to go from 10% to 90% of its maximum level. As illustrated, the turn-on time (t_{on}), is the sum of t_d and t_r .



The turn-on time for a switching transistor is the sum of the delay time (t_d) and the rise time (t_r). The transistor turn-off time is the sum of the storage time (t_s) and the fall time (t_f).

When the input current is switched off, I_C does not go to zero until after a turn-off time t_{off} , made up of a storage time (t_s), and a fall time (t_f), as illustrated. The fall time is specified as the time required for I_C to go from 90% to 10% of its maximum level. The storage time is the result of charge carriers being trapped in the depletion region when a junction polarity is reversed.

When a transistor is in a saturated on condition, both the collector-base and emitter-base junctions are forward biased. At switch-off both junctions are reverse biased, and before I_C begins to fall, the stored charge carriers must be withdrawn or made to recombine with opposite-type charge carriers.

Characteristic	Symbol	Min	Max	Unit
Switching Characteristics (2N3904)				
Delay time	$V_{CC} = 3\text{ V}, V_{BE} = 0.5\text{ V},$	t_d	—	35 ns
Rise time	$I_C = 10\text{ mA}, I_{B1} = 1\text{ mA}$	t_r	—	35 ns
Storage time	$V_{CC} = 3\text{ V}, I_C = 10\text{ mA},$	t_s	—	200 ns
Fall time	$I_{B1} = I_{B2} = 1\text{ mA}$	t_f	—	50 ns

Transistor switching times, as specified on a data sheet.

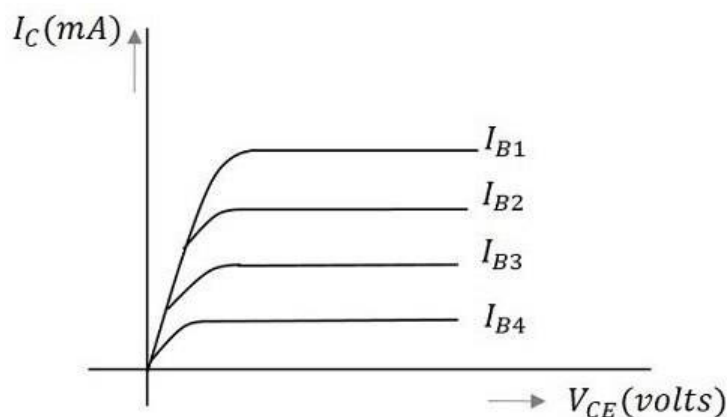
For a fast-switching transistor, t_{on} and t_{off} must be of the order of nanoseconds. The 2N3904 transistor data sheet portion in Fig. specifies the following Transistor Switching Times: $t_d = 35\text{ ns}$, $t_r = 35\text{ ns}$, $t_s = 200\text{ ns}$, and $t_f = 50\text{ ns}$.

TRANSISTOR BIASING AND STABILIZATION

Till now we have discussed different regions of operation for a transistor. But among all these regions, we have found that the transistor operates well in active region and hence it is also called as **linear region**. The outputs of the transistor are the collector current and collector voltages.

Output Characteristics

When the output characteristics of a transistor are considered, the curve looks as below for different input values.



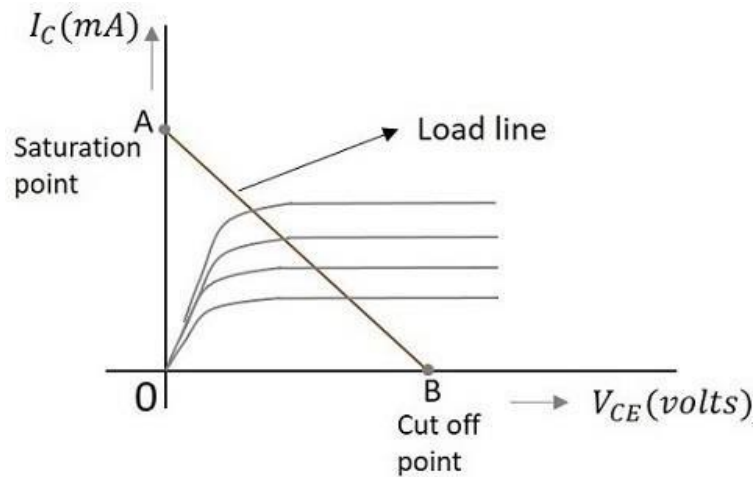
In the above figure, the output characteristics are drawn between collector current I_C and collector voltage V_{CE} for different values of base current I_B . These are considered here for different input values to obtain different output curves.

Load Line

When a value for the maximum possible collector current is considered, that point will be present on the Y-axis, which is nothing but the **Saturation point**. As well, when a value for the maximum possible collector emitter voltage is considered, that point will be present on the X-axis, which is the **Cutoff point**.

When a line is drawn joining these two points, such a line can be called as **Load line**. This is called so as it symbolizes the output at the load. This line, when drawn over the output characteristic curve, makes contact at a point called as **Operating point** or **quiescent point** or simply **Q-point**.

The concept of load line can be understood from the following graph.

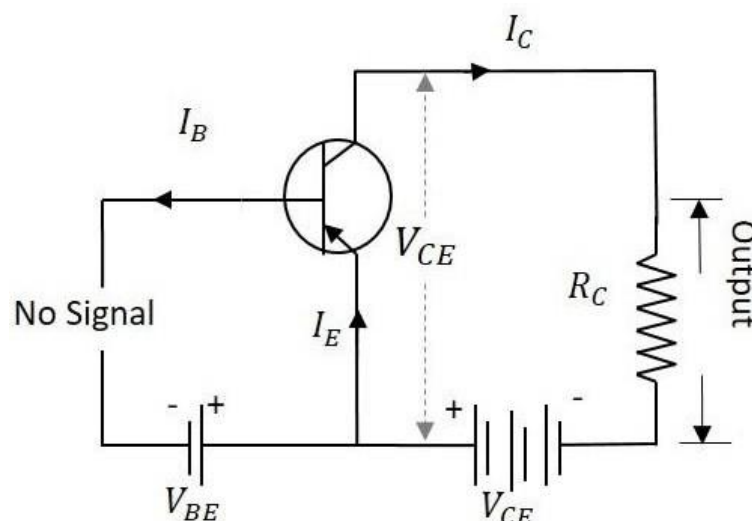


The load line is drawn by joining the saturation and cut off points. The region that lies between these two is the **linear region**. A transistor acts as a good amplifier in this linear region.

If this load line is drawn only when DC biasing is given to the transistor, but **no input** signal is applied, then such a load line is called as **DC load line**. Whereas the load line drawn under the conditions when an **input signal** along with the DC voltages are applied, such a line is called as an **AC load line**.

DC Load Line

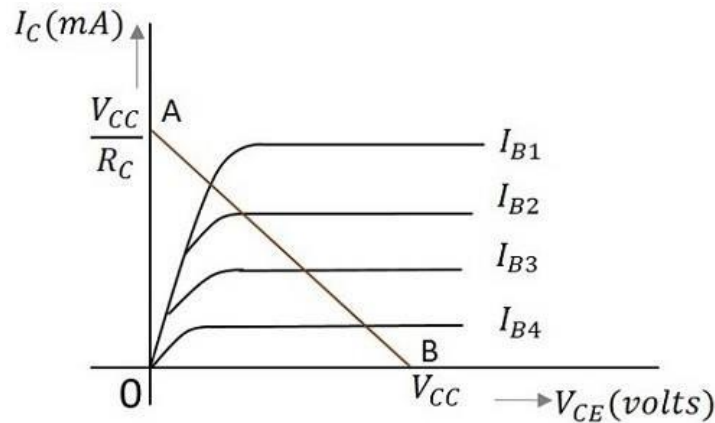
When the transistor is given the bias and no signal is applied at its input, the load line drawn under such conditions, can be understood as **DC condition**. Here there will be no amplification as the **signal is absent**. The circuit will be as shown below.



The value of collector emitter voltage at any given time will be

$$V_{CE} = V_{CC} - I_C R_C$$

As V_{CC} and R_C are fixed values, the above one is a first degree equation and hence will be a straight line on the output characteristics. This line is called as **D.C. Load line**. The figure below shows the DC load line.



To obtain the load line, the two end points of the straight line are to be determined. Let those two points be A and B.

To obtain A

When collector emitter voltage $V_{CE} = 0$, the collector current is maximum and is equal to V_{CC}/R_C . This gives the maximum value of V_{CE} . This is shown as

$$\begin{aligned} V_{CE} &= V_{CC} - I_C R_C \\ 0 &= V_{CC} - I_C R_C \\ I_C &= V_{CC}/R_C \end{aligned}$$

This gives the point A ($OA = V_{CC}/R_C$) on collector current axis, shown in the above figure.

To obtain B

When the collector current $I_C = 0$, then collector emitter voltage is maximum and will be equal to the V_{CC} . This gives the maximum value of I_C . This is shown as

$$\begin{aligned} V_{CE} &= V_{CC} - I_C R_C \\ &= V_{CC} \\ (\text{AS } I_C &= 0) \end{aligned}$$

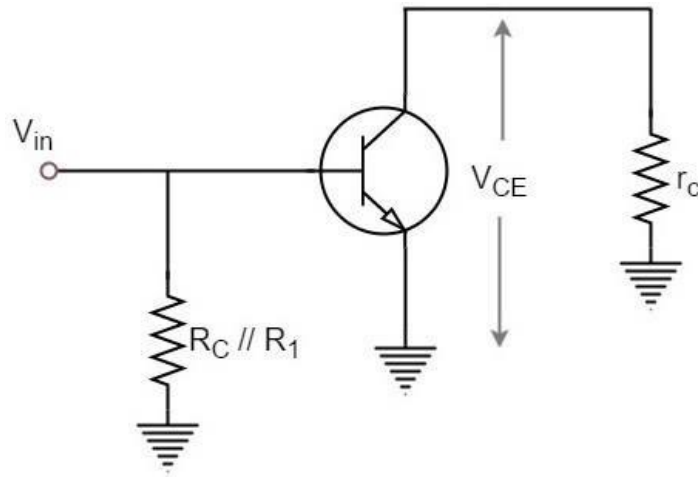
This gives the point B, which means ($OB = V_{CC}$) on the collector emitter voltage axis shown in the above figure.

Hence we got both the saturation and cutoff point determined and learnt that the load line is a straight line. So, a DC load line can be drawn.

AC Load Line

The DC load line discussed previously, analyzes the variation of collector currents and voltages, when no AC voltage is applied. Whereas the AC load line gives the peak-to-peak voltage, or the maximum possible output swing for a given amplifier.

We shall consider an AC equivalent circuit of a CE amplifier for our understanding.



From the above figure,

$$V_{CE} = (R_C // R_1) \times I_C$$

$$r_c = R_C // R_1$$

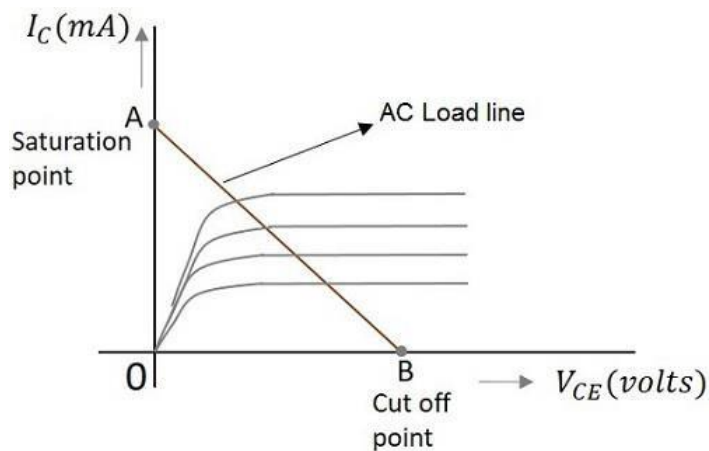
For a transistor to operate as an amplifier, it should stay in active region. The quiescent point is so chosen in such a way that the maximum input signal excursion is symmetrical on both negative and positive half cycles.

Hence,

$$V_{\max} = V_{CEQ} \text{ and } V_{\min} = -V_{CEQ}$$

Where V_{CEQ} is the emitter-collector voltage at quiescent point

The following graph represents the AC load line which is drawn between saturation and cut off points.



From the graph above, the current I_C at the saturation point is

$$I_{C(\text{sat})} = I_{CQ} + (V_{CEQ} / r_c)$$

The voltage V_{CE} at the cutoff point is

$$V_{CE(\text{off})} = V_{CEQ} + I_{CQ} r_c$$

Hence the maximum current for that corresponding $V_{CEQ} = V_{CEQ} / (R_C // R_1)$ is

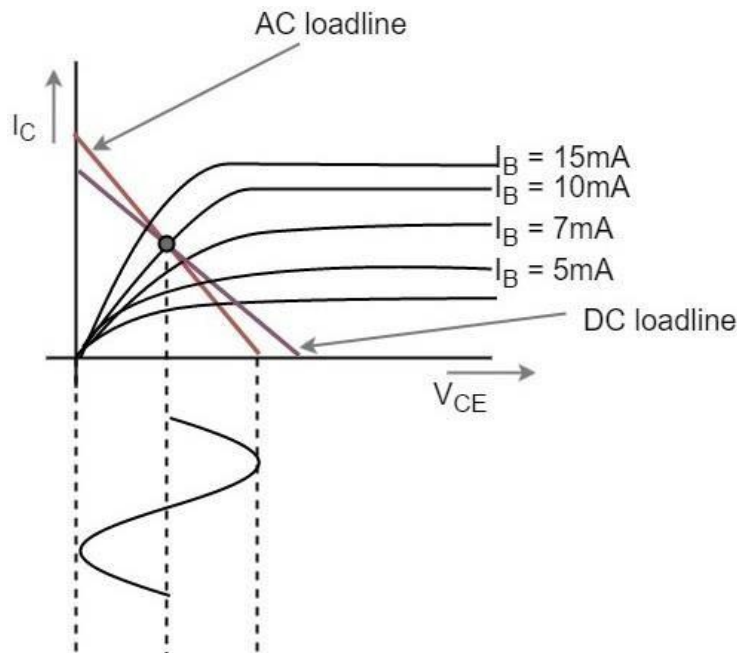
$$I_{CQ} = I_{CQ} * (R_C // R_1) / I_{CQ} = I_{CQ} * (R_C // R_1)$$

Hence by adding quiescent currents the end points of AC load line are

$$I_{C(sat)} = I_{CQ} + V_{CEQ} / (R_C // R_1)$$
$$V_{CE(off)} = V_{CEQ} + I_{CQ} * (R_C // R_1)$$

AC and DC Load Line

When AC and DC Load lines are represented in a graph, it can be understood that they are not identical. Both of these lines intersect at the **Q-point** or **quiescent point**. The endpoints of AC load line are saturation and cut off points. This is understood from the figure below.



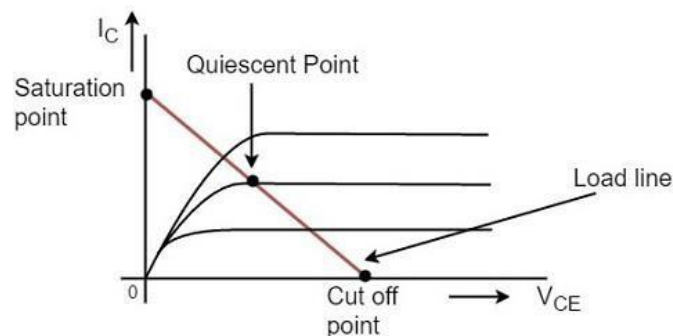
From the above figure, it is understood that the quiescent point (the dark dot) is obtained when the value of base current I_B is 10mA . This is the point where both the AC and DC load lines intersect.

OPERATING POINT

When a line is drawn joining the saturation and cut off points, such a line can be called as **Load line**. This line, when drawn over the output characteristic curve, makes contact at a point called as **Operating point**.

This operating point is also called as **quiescent point** or simply **Q-point**. There can be many such intersecting points, but the Q-point is selected in such a way that irrespective of AC signal swing, the transistor remains in the active region.

The following graph shows how to represent the operating point.

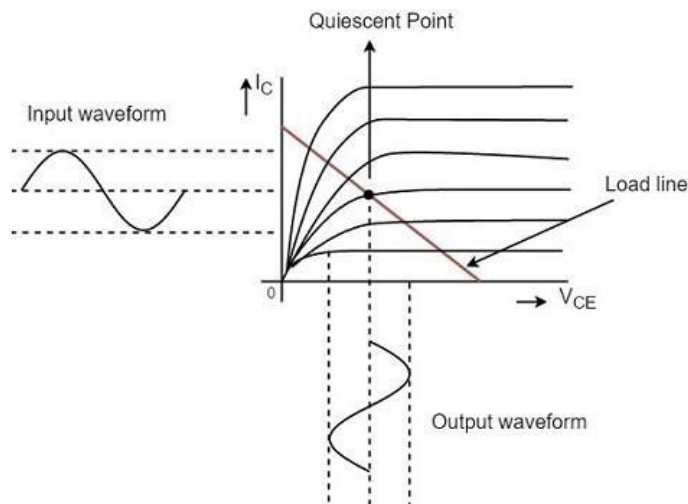


The operating point should not get disturbed as it should remain stable to achieve faithful amplification. Hence the quiescent point or Q-point is the value where the **Faithful Amplification** is achieved.

Faithful Amplification

The process of increasing the signal strength is called as **Amplification**. This amplification when done without any loss in the components of the signal, is called as **Faithful amplification**.

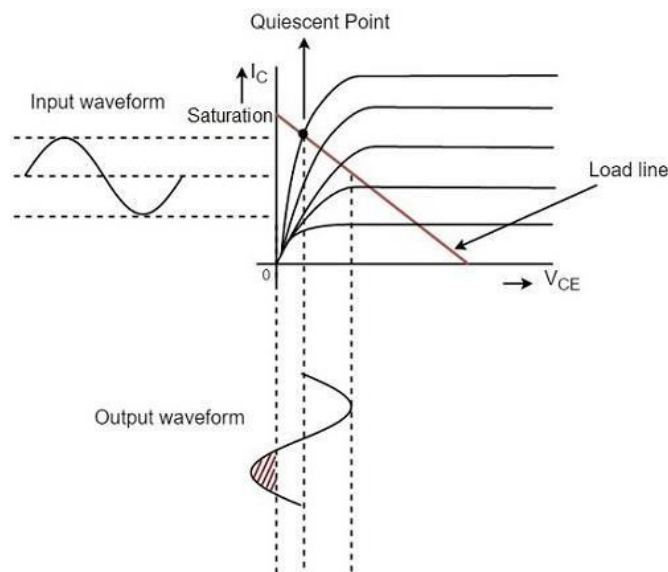
Faithful amplification is the process of obtaining complete portions of input signal by increasing the signal strength. This is done when AC signal is applied at its input.



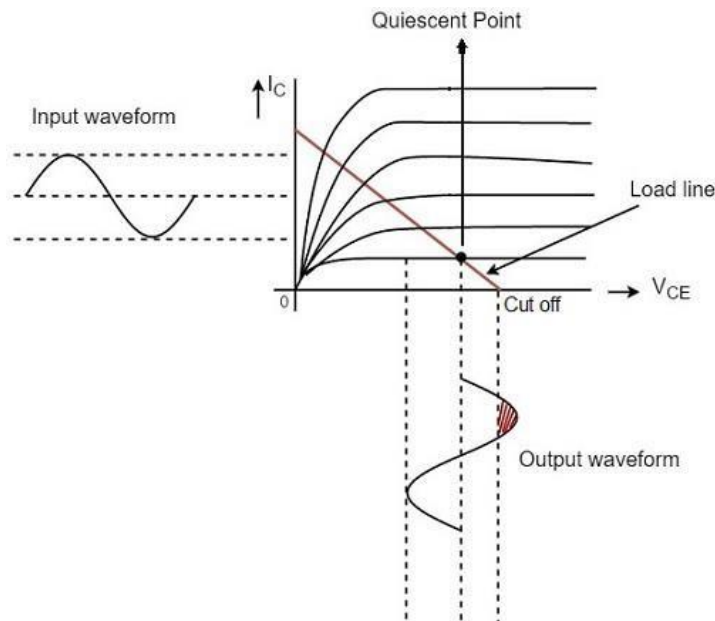
In the above graph, the input signal applied is completely amplified and reproduced without any losses. This can be understood as **Faithful Amplification**.

The operating point is so chosen such that it lies in the **active region** and it helps in the reproduction of complete signal without any loss.

If the operating point is considered near saturation point, then the amplification will be as under.



If the operation point is considered near cut off point, then the amplification will be as under.



Hence the placement of operating point is an important factor to achieve faithful amplification. But for the transistor to function properly as an amplifier, its input circuit (i.e., the base-emitter junction) remains forward biased and its output circuit (i.e., collector-base junction) remains reverse biased.

The amplified signal thus contains the same information as in the input signal whereas the strength of the signal is increased.

Key factors for Faithful Amplification

To ensure faithful amplification, the following basic conditions must be satisfied.

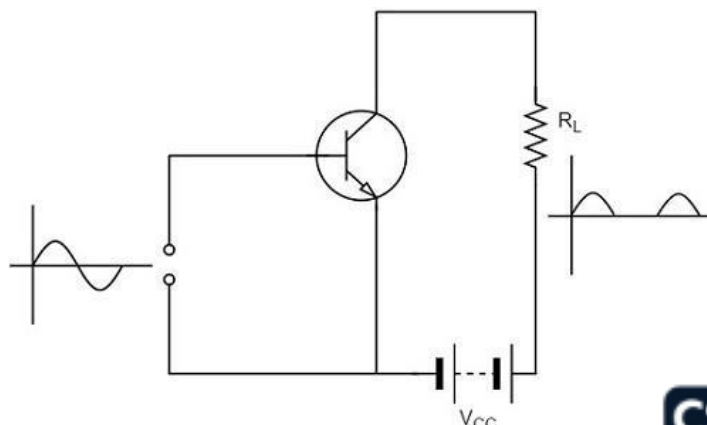
- Proper zero signal collector current
- Minimum proper base-emitter voltage (V_{BE}) at any instant.
- Minimum proper collector-emitter voltage (V_{CE}) at any instant.

The fulfillment of these conditions ensures that the transistor works over the active region having input forward biased and output reverse biased.

Proper Zero Signal Collector Current

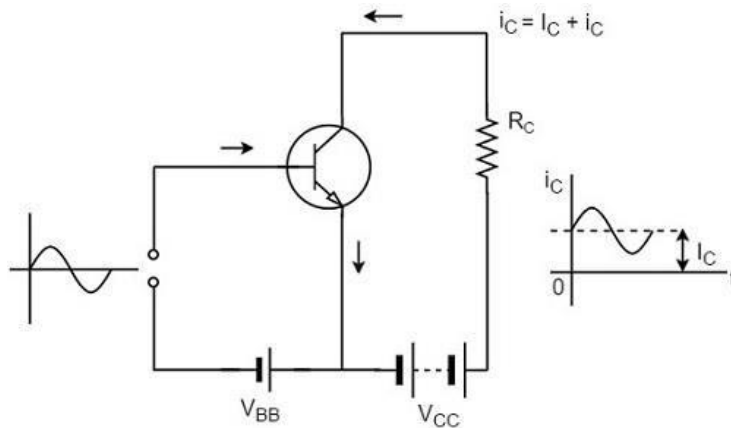
In order to understand this, let us consider a NPN transistor circuit as shown in the figure below. The base-emitter junction is forward biased and the collector-emitter junction is reverse biased. When a signal is applied at the input, the base-emitter junction of the NPN transistor gets forward biased for positive half cycle of the input and hence it appears at the output.

For negative half cycle, the same junction gets reverse biased and hence the circuit doesn't conduct. This leads to **unfaithful amplification** as shown in the figure below.



Let us now introduce a battery V_{BB} in the base circuit. The magnitude of this voltage should be such that the base-emitter junction of the transistor should remain in forward biased, even for negative half cycle of input signal. When no input signal is applied, a DC current flows in the circuit, due to V_{BB} . This is known as **zero signal collector current** I_C .

During the positive half cycle of the input, the base-emitter junction is more forward biased and hence the collector current increases. During the negative half cycle of the input, the input junction is less forward biased and hence the collector current decreases. Hence both the cycles of the input appear in the output and hence **faithful amplification** results, as shown in the below figure.



Hence for faithful amplification, proper zero signal collector current must flow. The value of zero signal collector current should be at least equal to the maximum collector current due to the signal alone.

Proper Minimum V_{BE} at any instant

The minimum base to emitter voltage V_{BE} should be greater than the cut-in voltage for the junction to be forward biased. The minimum voltage needed for a silicon transistor to conduct is 0.7v and for a germanium transistor to conduct is 0.5v. If the base-emitter voltage V_{BE} is greater than this voltage, the potential barrier is overcome and hence the base current and collector currents increase sharply.

Hence if V_{BE} falls low for any part of the input signal, that part will be amplified to a lesser extent due to the resultant small collector current, which results in unfaithful amplification.

Proper Minimum V_{CE} at any instant

To achieve a faithful amplification, the collector emitter voltage V_{CE} should not fall below the cut-in voltage, which is called as **Knee Voltage**. If V_{CE} is lesser than the knee voltage, the collector base junction will not be properly reverse biased. Then the collector cannot attract the electrons which are emitted by the emitter and they will flow towards base which increases the base current. Thus the value of β falls.

Therefore, if V_{CE} falls low for any part of the input signal, that part will be multiplied to a lesser extent, resulting in unfaithful amplification. So if V_{CE} is greater than V_{KNEE} the collector-base junction is properly reverse biased and the value of β remains constant, resulting in faithful amplification.

Biasing is the process of providing DC voltage which helps in the functioning of the circuit. A transistor is biased in order to make the emitter base junction forward biased and collector base junction reverse biased, so that it maintains in active region, to work as an amplifier.

In the previous chapter, we explained how a transistor acts as a good amplifier, if both the input and output sections are biased.

Transistor Biasing

The proper flow of zero signal collector current and the maintenance of proper collector-emitter voltage during the passage of signal is known as **Transistor Biasing**. The circuit which provides transistor biasing is called as **Biasing Circuit**.

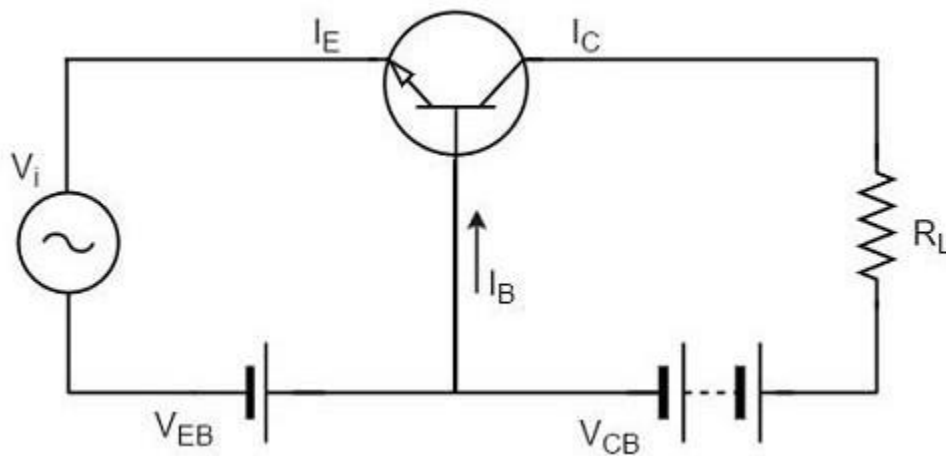
Need for DC biasing

If a signal of very small voltage is given to the input of BJT, it cannot be amplified. Because, for a BJT, to amplify a signal, two conditions have to be met.

- The input voltage should exceed **cut-in voltage** for the transistor to be **ON**.
- The BJT should be in the **active region**, to be operated as an **amplifier**.

If appropriate DC voltages and currents are given through BJT by external sources, so that BJT operates in active region and superimpose the AC signals to be amplified, then this problem can be avoided. The given DC voltage and currents are so chosen that the transistor remains in active region for entire input AC cycle. Hence DC biasing is needed.

The below figure shows a transistor amplifier that is provided with DC biasing on both input and output circuits.



For a transistor to be operated as a faithful amplifier, the operating point should be stabilized. Let us have a look at the factors that affect the stabilization of operating point.

Factors affecting the operating point

The main factor that affects the operating point is the temperature. The operating point shifts due to change in temperature.

As temperature increases, the values of I_{CE} , β , V_{BE} gets affected.

- I_{CBO} gets doubled (for every 10° rise)
- V_{BE} decreases by 2.5mV (for every 1° rise)

So the main problem which affects the operating point is temperature. Hence operating point should be made independent of the temperature so as to achieve stability. To achieve this, biasing circuits are introduced.

STABILIZATION

The process of making the operating point independent of temperature changes or variations in transistor parameters is known as **Stabilization**.

Once the stabilization is achieved, the values of I_C and V_{CE} become independent of temperature variations or replacement of transistor. A good biasing circuit helps in the stabilization of operating point.

Need for Stabilization

Stabilization of the operating point has to be achieved due to the following reasons.

- Temperature dependence of I_C
- Individual variations
- Thermal runaway

Let us understand these concepts in detail.

Temperature Dependence of I_C

As the expression for collector current I_C is

$$I_C = \beta I_B + I_{CEO} \\ = \beta I_B + (\beta + 1) I_{CBO}$$

The collector leakage current I_{CBO} is greatly influenced by temperature variations. To come out of this, the biasing conditions are set so that zero signal collector current $I_C = 1 \text{ mA}$. Therefore, the operating point needs to be stabilized i.e. it is necessary to keep I_C constant.

Individual Variations

As the value of β and the value of V_{BE} are not same for every transistor, whenever a transistor is replaced, the operating point tends to change. Hence it is necessary to stabilize the operating point.

Thermal Runaway

As the expression for collector current I_C is

$$I_C = \beta I_B + I_{CEO} \\ = \beta I_B + (\beta + 1) I_{CBO}$$

The flow of collector current and also the collector leakage current causes heat dissipation. If the operating point is not stabilized, there occurs a cumulative effect which increases this heat dissipation.

The self-destruction of such an unstabilized transistor is known as **Thermal run away**.

In order to avoid **thermal runaway** and the destruction of transistor, it is necessary to stabilize the operating point, i.e., to keep I_C constant.

STABILITY FACTOR

It is understood that I_C should be kept constant in spite of variations of I_{CBO} or I_{CO} . The extent to which a biasing circuit is successful in maintaining this is measured by **Stability factor**. It denoted by S .

By definition, the rate of change of collector current I_C with respect to the collector leakage current I_{CO} at constant β and I_B is called **Stability factor**.

$$S = dI_C / dI_{CO} \text{ at constant } I_B \text{ and } \beta$$

Hence we can understand that any change in collector leakage current changes the collector current to a great extent. The stability factor should be as low as possible so that the collector current doesn't get affected. $S=1$ is the ideal value.

The general expression of stability factor for a CE configuration can be defined as under.

$$I_C = \beta I_B + (\beta + 1) I_{CO}$$

Differentiating above expression with respect to I_C , we get

$$1 = \beta (dI_B / dI_C) + (\beta + 1) (dI_{CO} / dI_C)$$

Or

$$1 = [\beta (dI_B / dI_C)] + [(\beta + 1) / S]$$

$$\text{Since } (dI_{CO} / dI_C) = (1/S)$$

Or

$$S = (\beta + 1) / [1 - \beta (dI_B / dI_C)]$$

Hence the stability factor S depends on β , I_B and I_C .

METHODS OF TRANSISTOR BIASING:

The biasing in transistor circuits is done by using two DC sources V_{BB} and V_{CC} . It is economical to minimize the DC source to one supply instead of two which also makes the circuit simple.

The commonly used methods of transistor biasing are

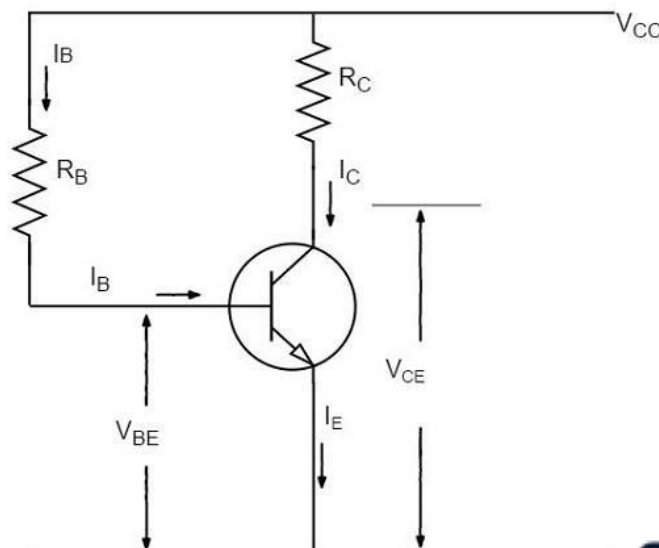
- Base Resistor method
- Collector to Base bias
- Biasing with Collector feedback resistor
- Voltage-divider bias

All of these methods have the same basic principle of obtaining the required value of I_B and I_C from V_{CC} in the zero signal conditions.

Base Resistor Method

In this method, a resistor R_B of high resistance is connected in base, as the name implies. The required zero signal base current is provided by V_{CC} which flows through R_B . The base emitter junction is forward biased, as base is positive with respect to emitter.

The required value of zero signal base current and hence the collector current (as $I_C = \beta I_B$) can be made to flow by selecting the proper value of base resistor R_B . Hence the value of R_B is to be known. The figure below shows how a base resistor method of biasing circuit looks like.



Let I_C be the required zero signal collector current. Therefore,

$$I_B = I_C / \beta$$

Considering the closed circuit from V_{CC} , base, emitter and ground, while applying the Kirchhoff's voltage law, we get,

$$V_{CC} = I_B R_B + V_{BE}$$

Or

$$I_B R_B = V_{CC} - V_{BE}$$

Therefore

$$R_B = (V_{CC} - V_{BE}) / I_B$$

Since V_{BE} is generally quite small as compared to V_{CC} , the former can be neglected with little error. Then,

$$R_B = V_{CC} / I_B$$

We know that V_{CC} is a fixed known quantity and I_B is chosen at some suitable value. As R_B can be found directly, this method is called as **fixed bias method**.

Stability factor

$$S = (\beta + 1) [1 - \beta (dI_B / dI_C)]$$

In fixed-bias method of biasing, I_B is independent of I_C so that,

$$dI_B / dI_C = 0$$

Substituting the above value in the previous equation,

$$\text{Stability factor, } S = \beta + 1$$

Thus the stability factor in a fixed bias is $(\beta + 1)$ which means that I_C changes $(\beta + 1)$ times as much as any change in I_{CO} .

Advantages

- The circuit is simple.
- Only one resistor R_E is required.
- Biasing conditions are set easily.
- No loading effect as no resistor is present at base-emitter junction.

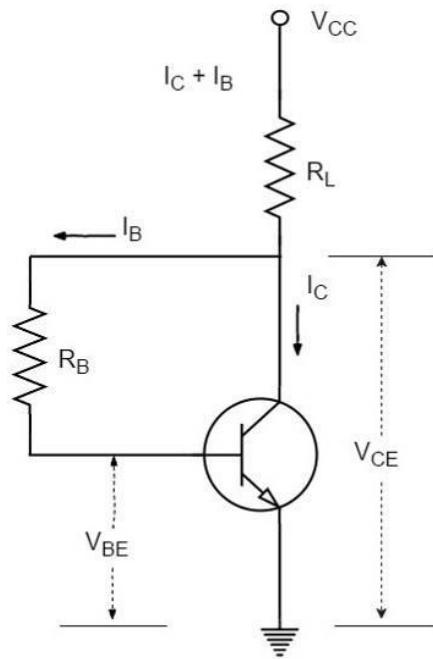
Disadvantages

- The stabilization is poor as heat development can't be stopped.
- The stability factor is very high. So, there are strong chances of thermal run away.

Hence, this method is rarely employed.

Collector to Base Bias

The collector to base bias circuit is same as base bias circuit except that the base resistor R_B is returned to collector, rather than to V_{CC} supply as shown in the figure below.



This circuit helps in improving the stability considerably. If the value of I_C increases, the voltage across R_L increases and hence the V_{CE} also increases. This in turn reduces the base current I_B . This action somewhat compensates the original increase.

The required value of R_B needed to give the zero signal collector current I_C can be calculated as follows.

Voltage drop across R_L will be

$$R_L = (I_C + I_B)R_L \cong I_C R_L$$

From the figure,

$$I_C R_L + I_B R_B + V_{BE} = V_{CC}$$

Or

$$I_B R_B = V_{CC} - V_{BE} - I_C R_L$$

Therefore

$$R_B = (V_{CC} - V_{BE} - I_C R_L) / I_B$$

Or

$$R_B = [(V_{CC} - V_{BE} - I_C R_L) \beta] / I_C$$

Applying KVL we have

$$(I_B + I_C)R_L + I_B R_B + V_{BE} = V_{CC}$$

Or

$$I_B (R_L + R_B) + I_C R_L + V_{BE} = V_{CC}$$

Therefore

$$I_B = (V_{CC} - V_{BE} - I_C R_L) / (R_L + R_B)$$

Since V_{BE} is almost independent of collector current, we get

$$dI_B / dI_C = -(R_L) / (R_L + R_B)$$

We know that

$$S = (1 + \beta) / [1 - \beta(dI_B/dI_C)]$$

Therefore

$$S = (1 + \beta) / [1 + \beta(R_L / (R_L + R_B))]$$

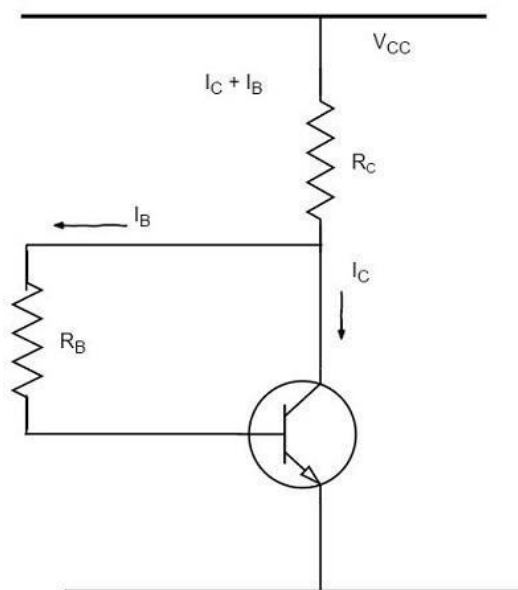
This value is smaller than $(1 + \beta)$ which is obtained for fixed bias circuit. Thus there is an improvement in the stability.

This circuit provides a negative feedback which reduces the gain of the amplifier. So the increased stability of the collector to base bias circuit is obtained at the cost of AC voltage gain.

Biasing with Collector Feedback resistor

In this method, the base resistor R_B has its one end connected to base and the other to the collector as its name implies. In this circuit, the zero signal base current is determined by V_{CB} but not by V_{CC} .

It is clear that V_{CB} forward biases the base-emitter junction and hence base current I_B flows through R_B . This causes the zero signal collector current to flow in the circuit. The below figure shows the biasing with collector feedback resistor circuit.



The required value of R_B needed to give the zero signal current I_C can be determined as follows.

$$V_{CC} = I_C R_C + I_B R_B + V_{BE}$$

Or

$$\begin{aligned} R_B &= (V_{CC} - V_{BE} - I_C R_C) / I_B \\ &= (V_{CC} - V_{BE} - \beta I_B R_C) / I_B \\ &\text{Since } I_C = \beta I_B \end{aligned}$$

Alternatively,

$$V_{CE} = V_{BE} + V_{CB}$$

Or

$$V_{CB} = V_{CE} - V_{BE}$$

Since

$$R_B = V_{CB} / I_B = (V_{CE} - V_{BE}) / I_B$$

Where

$$I_B = I_C / \beta$$

Mathematically,

$$\text{Stability factor, } S < (\beta + 1)$$

Therefore, this method provides better thermal stability than the fixed bias.

The Q-point values for the circuit are shown as

$$I_C = (V_{CC} - V_{BE}) / [(R_B / \beta) + R_C]$$

$$V_{CE} = V_{CC} - I_C R_C$$

Advantages

- The circuit is simple as it needs only one resistor.
- This circuit provides some stabilization, for lesser changes.

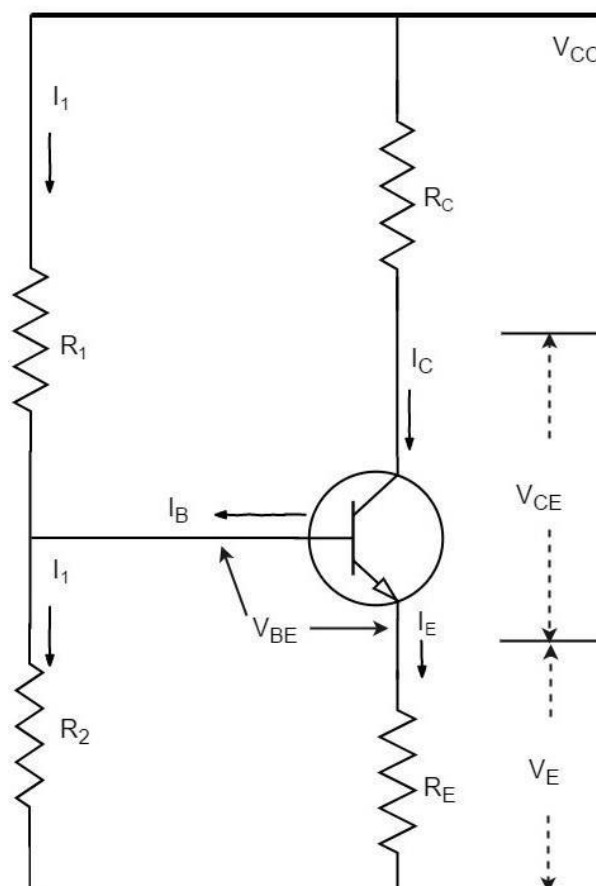
Disadvantages

- The circuit doesn't provide good stabilization.
- The circuit provides negative feedback.

Voltage Divider Bias Method

Among all the methods of providing biasing and stabilization, the **voltage divider bias method** is the most prominent one. Here, two resistors R_1 and R_2 are employed, which are connected to V_{CC} and provide biasing. The resistor R_E employed in the emitter provides stabilization.

The name voltage divider comes from the voltage divider formed by R_1 and R_2 . The voltage drop across R_2 forward biases the base-emitter junction. This causes the base current and hence collector current flow in the zero signal conditions. The figure below shows the circuit of voltage divider bias method.



Suppose that the current flowing through resistance R_1 is I_1 . As base current I_B is very small, therefore, it can be assumed with reasonable accuracy that current flowing through R_2 is also I_1 .

Now let us try to derive the expressions for collector current and collector voltage.

Collector Current, I_C

From the circuit, it is evident that,

$$I_1 = V_{CC}/(R_1 + R_2)$$

Therefore, the voltage across resistance R_2 is

$$V_2 = [V_{CC}/(R_1 + R_2)]R_2$$

Applying Kirchhoff's voltage law to the base circuit,

$$V_2 = V_{BE} + V_E$$

$$V_2 = V_{BE} + I_E R_E$$

$$I_E = (V_2 - V_{BE})/R_E$$

Since $I_E \approx I_C$,

$$I_C = (V_2 - V_{BE})/R_E$$

From the above expression, it is evident that I_C doesn't depend upon β . V_{BE} is very small that I_C doesn't get affected by V_{BE} at all. Thus I_C in this circuit is almost independent of transistor parameters and hence good stabilization is achieved.

Collector-Emitter Voltage, V_{CE}

Applying Kirchhoff's voltage law to the collector side,

$$V_{CC} = I_C R_C + V_{CE} + I_E R_E$$

Since $I_E \cong I_C$

$$= I_C R_C + V_{CE} + I_C R_E$$

$$= I_C (R_C + R_E) + V_{CE}$$

Therefore,

$$V_{CE} = V_{CC} - I_C (R_C + R_E)$$

R_E provides excellent stabilization in this circuit.

$$V_2 = V_{BE} + I_C R_E$$

Suppose there is a rise in temperature, then the collector current I_C decreases, which causes the voltage drop across R_E to increase. As the voltage drop across R_2 is V_2 , which is independent of I_C , the value of V_{BE} decreases. The reduced value of I_B tends to restore I_C to the original value.

Stability Factor

The equation for **Stability factor** of this circuit is obtained as

$$\begin{aligned} \text{Stability Factor} = S &= [(\beta + 1)(R_0 + R_3)]/[R_0 + R_E + \beta R_E] \\ &= [(\beta + 1)] \times [(1 + R_0/R_E)/[\beta + 1 + (R_0/R_E)]] \end{aligned}$$

Where

$$R_0 = [(R_1 R_2)/(R_1 + R_2)]$$

If the ratio R_0/R_E is very small, then R_0/R_E can be neglected as compared to 1 and the stability factor becomes

$$\text{Stability Factor} = S = (\beta + 1) \times [1 / (\beta + 1)] = 1$$

This is the smallest possible value of S and leads to the maximum possible thermal stability.

BIAS COMPENSATION USING DIODES

So far we have seen different stabilization techniques. The stabilization occurs due to negative feedback action. The negative feedback, although improves the stability of operating point, it reduces the gain of the amplifier.

As the gain of the amplifier is a very important consideration, some compensation techniques are used to maintain excellent bias and thermal stabilization. Let us now go through such bias compensation techniques.

Diode Compensation for Instability

These are the circuits that implement compensation techniques using diodes to deal with biasing instability. The stabilization techniques refer to the use of resistive biasing circuits which permit I_B to vary so as to keep I_C relatively constant.

There are two types of diode compensation methods. They are –

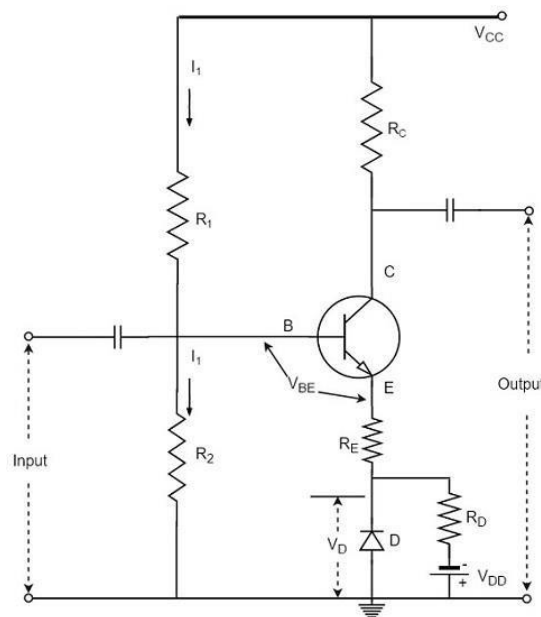
- Diode compensation for instability due to V_{BE} variation
- Diode compensation for instability due to I_{CO} variation

Let us understand these two compensation methods in detail.

Diode Compensation for Instability due to V_{BE} Variation

In a Silicon transistor, the changes in the value of V_{BE} results in the changes in I_C . A diode can be employed in the emitter circuit in order to compensate the variations in V_{BE} or I_{CO} . As the diode and transistor used are of same material, the voltage V_D across the diode has same temperature coefficient as V_{BE} of the transistor.

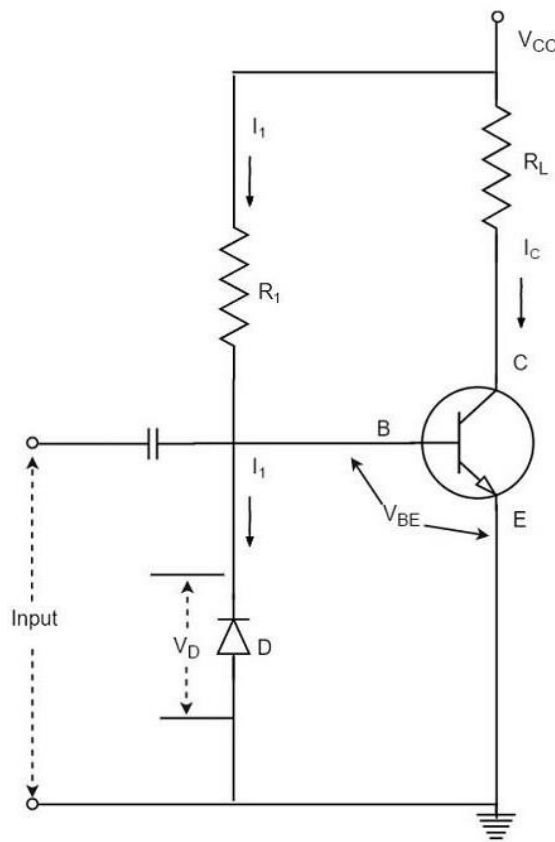
The following figure shows self-bias with stabilization and compensation.



The diode D is forward biased by the source V_{DD} and the resistor R_D . The variation in V_{BE} with temperature is same as the variation in V_D with temperature, hence the quantity $(V_{BE} - V_D)$ remains constant. So the current I_C remains constant in spite of the variation in V_{BE} .

Diode Compensation for Instability due to I_{CO} Variation

The following figure shows the circuit diagram of a transistor amplifier with diode D used for compensation of variation in I_{CO} .



So, the reverse saturation current I_o of the diode will increase with temperature at the same rate as the transistor collector saturation current I_{CO} .

$$I = (V_{CC} - V_{BE}) / R \cong V_{CC} / R = \text{Constant}$$

The diode D is reverse biased by V_{BE} and the current through it is the reverse saturation current I_o .

Now the base current is,

$$I_B = I - I_o$$

Substituting the above value in the expression for collector current.

$$I_C = \beta(I - I_o) + (1 + \beta)I_{CO}$$

If $\beta \gg 1$,

$$I_C = \beta I - \beta I_o + \beta I_{CO}$$

I is almost constant and if I_o of diode and I_{CO} of transistor track each other over the operating temperature range, then I_C remains constant.

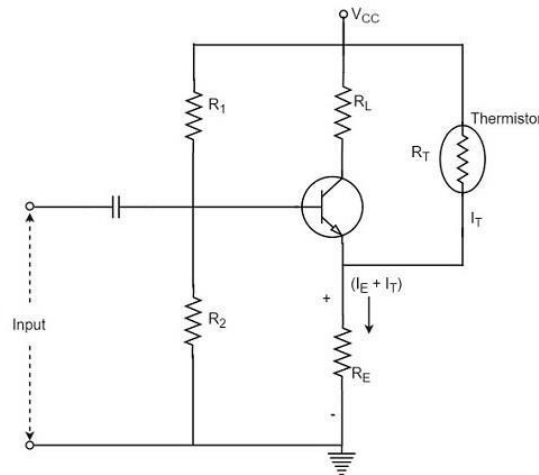
Other Compensations

There are other compensation techniques which refer to the use of temperature sensitive devices such as diodes, transistors, thermistors, Sensistors, etc. to compensate for the variation in currents.

There are two popular types of circuits in this method, one using a thermistor and the other using a Sensistor. Let us have a look at them.

Thermistor Compensation

Thermistor is a temperature sensitive device. It has negative temperature coefficient. The resistance of a thermistor increases when the temperature decreases and it decreases when the temperature increases. The below figure shows a self-bias amplifier with thermistor compensation.

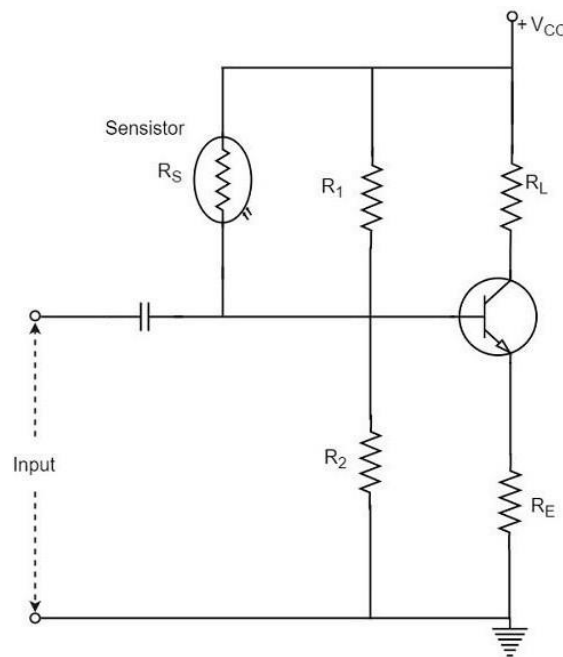


In an amplifier circuit, the changes that occur in I_{CO} , V_{BE} and β with temperature, increases the collector current. Thermistor is employed to minimize the increase in collector current. As the temperature increases, the resistance R_T of thermistor decreases, which increases the current through it and the resistor R_E . Now, the voltage developed across R_E increases, which reverse biases the emitter junction. This reverse bias is so high that the effect of resistors R_1 and R_2 providing forward bias also gets reduced. This action reduces the rise in collector current.

Thus the temperature sensitivity of thermistor compensates the increase in collector current, occurred due to temperature.

Sensistor Compensation

A Sensistor is a heavily doped semiconductor that has positive temperature coefficient. The resistance of a Sensistor increases with the increase in temperature and decreases with the decrease in temperature. The figure below shows a self-bias amplifier with Sensistor compensation.



In the above figure, the Sensistor may be placed in parallel with R_1 or in parallel with R_E . As the temperature increases, the resistance of the parallel combination, thermistor and R_1 increases and

their voltage drop also increases. This decreases the voltage drop across R_2 . Due to the decrease of this voltage, the net forward emitter bias decreases. As a result of this, I_C decreases.

Hence by employing the Sensistor, the rise in the collector current which is caused by the increase of I_{CO} , V_{BE} and β due to temperature, gets controlled.

THERMAL RESISTANCE

The transistor is a temperature dependent device. When the transistor is operated, the collector junction gets heavy flow of electrons and hence has much heat generated. This heat if increased further beyond the permissible limit, damages the junction and thus the transistor.

In order to protect itself from damage, the transistor dissipates heat from the junction to the transistor case and from there to the open air surrounding it.

Let, the ambient temperature or the temperature of surrounding air = $T_A^\circ\text{C}$

And, the temperature of collector-base junction of the transistor = $T_J^\circ\text{C}$

As $T_J > T_A$, the difference $T_J - T_A$ is greater than the power dissipated in the transistor P_D will be greater. Thus,

$$\begin{aligned} T_J - T_A &\propto P_D \\ T_J - T_A &= H P_D \end{aligned}$$

Where H is the constant of proportionality, and is called as **Thermal resistance**.

Thermal resistance is the resistance to heat flow from junction to surrounding air. It is denoted by H .

$$H = \frac{T_J - T_A}{P_D}$$

The unit of H is $^\circ\text{C}/\text{watt}$.

If the thermal resistance is low, the transfer of heat from the transistor into the air, will be easy. If the transistor case is larger, the heat dissipation will be better. This is achieved by the use of Heat sink.

Heat Sink

The transistor that handle larger powers, dissipates more heat during operation. This heat if not dissipated properly, could damage the transistor. Hence the power transistors are generally mounted on large metal cases to provide a larger area to get the heat radiated that is generated during its operation.



The metal sheet that helps to dissipate the additional heat from the transistor is known as the **heat sink**. The ability of a heat sink depends upon its material, volume, area, shape, contact between case and sink, and the movement of air around the sink.

The heat sink is selected after considering all these factors. The image shows a power transistor with a heat sink.

A tiny transistor in the above image is fixed to a larger metal sheet in order to dissipate its heat, so that the transistor doesn't get damaged.

Thermal Runaway

The use of heat sink avoids the problem of **Thermal Runaway**. It is a situation where an increase in temperature leads to the condition that further increase in temperature, leads to the destruction of the device itself. This is a kind of uncontrollable positive feedback.

Heat sink is not the only consideration; other factors such as operating point, ambient temperature, and the type of transistor used can also cause thermal runaway.

CONDITION FOR THERMAL STABILITY:

For preventing thermal runaway, the required condition is the rate at which the heat is released at the collector junction should not exceed the rate at which the heat can be dissipated under steady state condition.

Hence the condition to be satisfied to avoid thermal runaway is given by

$$\frac{\partial P_C}{\partial T_j} < \frac{1}{\theta}$$

If the circuit is properly designed, then the transistor cannot runaway below a specified ambient temperature or even under any conditions.

In the self biased circuit the transistor is biased in the active region. The power generated at the junction without any signal is

$$P_C = I_C V_{CE} \approx I_C V_{CC}$$

Let us assume that the quiescent collector and the emitter currents are equal. Then

$$P_C = I_C V_{CC} - I_C^2 (R_E + R_C) \quad P_C = I_C V_{CC} - I_C^2 (R_E + R_C) \dots \dots \dots (1)$$

The condition to prevent thermal runaway can be written as

$$\frac{\partial P_C}{\partial I_C} \frac{\partial I_C}{\partial T_j} < \frac{1}{\theta} \dots \dots \dots (2)$$

As θ and $\frac{\partial I_C}{\partial T_j}$ are positive, $\frac{\partial P_C}{\partial I_C}$ should be negative in order to satisfy the above condition.

Differentiating equation (1) w.r.t I_C we get

$$\frac{\partial P_C}{\partial I_C} = V_{CC} - 2I_C (R_E + R_C) \dots \dots \dots (3)$$

Hence to avoid thermal runaway it is necessary that

$$I_C > \frac{V_{CC}}{2(R_E + R_C)} \dots \dots \dots (4)$$

Since $V_{CE} = V_{CC} - I_C(R_E + R_C)$ then eq(4) implies that $V_{CE} < V_{CC}/2$. IF the inequality of eq(4) is not satisfied and $V_{CE} < V_{CC}/2$, then from eq(3), $\frac{\partial P_C}{\partial I_C} \frac{\partial P_C}{\partial I_C}$ is positive., and the corresponding eq(2) should be satisfied. Other wise thermal runaway will occur.

UNIT III

JUNCTION FIELD EFFECT TRANSISTOR (FET)

In the *Bipolar Junction Transistor* tutorials, we saw that the output Collector current of the transistor is proportional to input current flowing into the Base terminal of the device, thereby making the bipolar transistor a “CURRENT” operated device (Beta model) as a smaller current can be used to switch a larger load current.

The **Field Effect Transistor**, or simply **FET** however, uses the voltage that is applied to their input terminal, called the Gate to control the current flowing through them resulting in the output current being proportional to the input voltage. As their operation relies on an electric field (hence the name field effect) generated by the input Gate voltage, this then makes the **Field Effect Transistor** a “VOLTAGE” operated device.



Typical Field Effect Transistor

The **Field Effect Transistor** is a three terminal unipolar semiconductor device that has very similar characteristics to those of their *Bipolar Transistor* counterparts. For example, high efficiency, instant operation, robust and cheap and can be used in most electronic circuit applications to replace their equivalent bipolar junction transistors (BJT) cousins.

Field effect transistors can be made much smaller than an equivalent BJT transistor and along with their low power consumption and power dissipation makes them ideal for use in integrated circuits such as the CMOS range of digital logic chips.

We remember from the previous tutorials that there are two basic types of bipolar transistor construction, NPN and PNP, which basically describes the physical arrangement of the P-type and N-type semiconductor materials from which they are made. This is also true of FET's as there are also two basic classifications of Field Effect Transistor, called the N-channel FET and the P-channel FET.

The field effect transistor is a three terminal device that is constructed with no PN-junctions within the main current carrying path between the Drain and the Source terminals. These terminals correspond in function to the Collector and the Emitter respectively of the bipolar transistor. The current path between these two terminals is called the “channel” which may be made of either a P-type or an N-type semiconductor material.

The control of current flowing in this channel is achieved by varying the voltage applied to the Gate. As their name implies, Bipolar Transistors are “Bipolar” devices because they operate with both types of charge carriers, Holes and Electrons. The Field Effect Transistor on the other hand is a “Unipolar” device that depends only on the conduction of electrons (N-channel) or holes (P-channel).

The **Field Effect Transistor** has one major advantage over its standard bipolar transistor cousins, in that their input impedance, (R_{in}) is very high, (thousands of Ohms), while the BJT is comparatively low. This very high input impedance makes them very sensitive to input voltage signals, but the price of this high sensitivity also means that they can be easily damaged by static electricity.

There are two main types of field effect transistor, the **Junction Field Effect Transistor** or **JFET** and the **Insulated-gate Field Effect Transistor** or **IGFET**), which is more commonly known as the standard **Metal Oxide Semiconductor Field Effect Transistor** or **MOSFET** for short.

The Junction Field Effect Transistor

We saw previously that a bipolar junction transistor is constructed using two PN-junctions in the main current carrying path between the Emitter and the Collector terminals. The **Junction Field Effect Transistor** (JUGFET or JFET) has no PN-junctions but instead has a narrow piece of high resistivity semiconductor material forming a “Channel” of either N-type or P-type silicon for the majority carriers to flow through with two ohmic electrical connections at either end commonly called the Drain and the Source respectively.

There are two basic configurations of junction field effect transistor, the N-channel JFET and the P-channel JFET. The N-channel JFET’s channel is doped with donor impurities meaning that the flow of current through the channel is negative (hence the term N-channel) in the form of electrons.

Likewise, the P-channel JFET’s channel is doped with acceptor impurities meaning that the flow of current through the channel is positive (hence the term P-channel) in the form of holes. N-channel JFET’s have a greater channel conductivity (lower resistance) than their equivalent P-channel types, since electrons have a higher mobility through a conductor compared to holes. This makes the N-channel JFET’s a more efficient conductor compared to their P-channel counterparts.

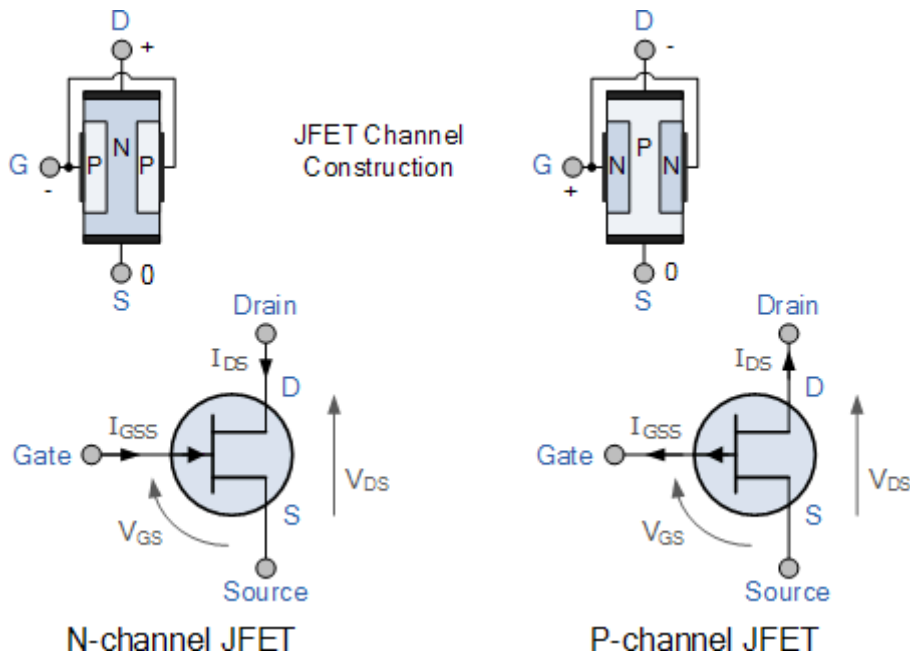
We have said previously that there are two ohmic electrical connections at either end of the channel called the Drain and the Source. But within this channel there is a third electrical connection which is called the Gate terminal and this can also be a P-type or N-type material forming a PN-junction with the main channel.

The relationship between the connections of a junction field effect transistor and a bipolar junction transistor are compared below.

Comparison of Connections between a JFET and a BJT

Bipolar Transistor (BJT)		Field Effect Transistor (FET)
Emitter – (E)	>>	Source – (S)
Base – (B)	>>	Gate – (G)
Collector – (C)	>>	Drain – (D)

The symbols and basic construction for both configurations of JFETs are shown below.



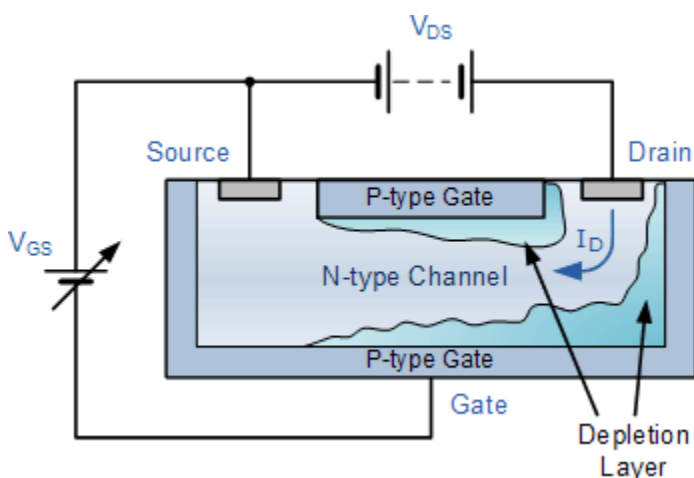
The semiconductor “channel” of the **Junction Field Effect Transistor** is a resistive path through which a voltage V_{DS} causes a current I_D to flow and as such the junction field effect transistor can conduct current equally well in either direction. As the channel is resistive in nature, a voltage gradient is thus formed down the length of the channel with this voltage becoming less positive as we go from the Drain terminal to the Source terminal.

The result is that the PN-junction therefore has a high reverse bias at the Drain terminal and a lower reverse bias at the Source terminal. This bias causes a “depletion layer” to be formed within the channel and whose width increases with the bias.

The magnitude of the current flowing through the channel between the Drain and the Source terminals is controlled by a voltage applied to the Gate terminal, which is a reverse-biased. In an N-channel JFET this Gate voltage is negative while for a P-channel JFET the Gate voltage is positive.

The main difference between the JFET and a BJT device is that when the JFET junction is reverse-biased the Gate current is practically zero, whereas the Base current of the BJT is always some value greater than zero.

Biasing of an N-channel JFET



The cross sectional diagram above shows an N-type semiconductor channel with a P-type region called the Gate diffused into the N-type channel forming a reverse biased PN-junction and it is this junction which forms the *depletion region* around the Gate area when no external voltages are applied. JFETs are therefore known as depletion mode devices.

This depletion region produces a potential gradient which is of varying thickness around the PN-junction and restrict the current flow through the channel by reducing its effective width and thus increasing the overall resistance of the channel itself.

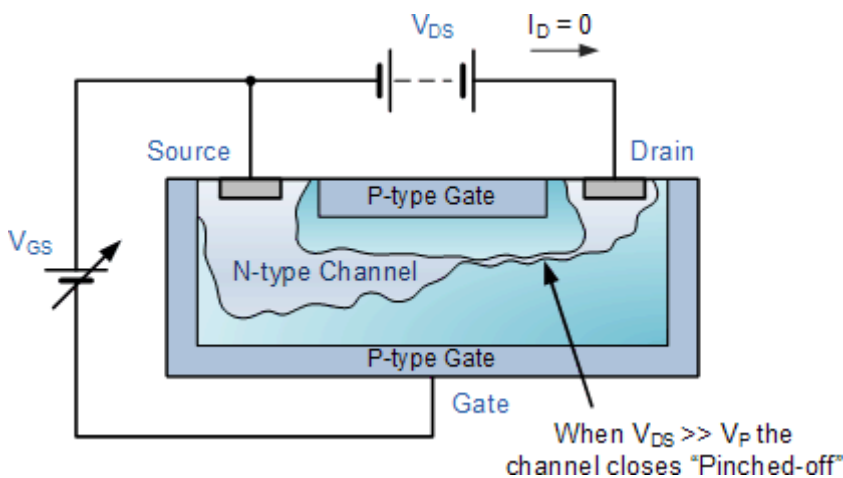
Then we can see that the most-depleted portion of the depletion region is in between the Gate and the Drain, while the least-depleted area is between the Gate and the Source. Then the JFET's channel conducts with zero bias voltage applied (ie, the depletion region has near zero width).

With no external Gate voltage ($V_G = 0$), and a small voltage (V_{DS}) applied between the Drain and the Source, maximum saturation current (I_{DSS}) will flow through the channel from the Drain to the Source restricted only by the small depletion region around the junctions.

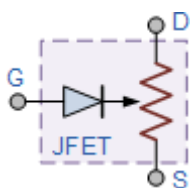
If a small negative voltage ($-V_{GS}$) is now applied to the Gate the size of the depletion region begins to increase reducing the overall effective area of the channel and thus reducing the current flowing through it, a sort of “squeezing” effect takes place. So by applying a reverse bias voltage increases the width of the depletion region which in turn reduces the conduction of the channel.

Since the PN-junction is reverse biased, little current will flow into the gate connection. As the Gate voltage ($-V_{GS}$) is made more negative, the width of the channel decreases until no more current flows between the Drain and the Source and the FET is said to be “pinched-off” (similar to the cut-off region for a BJT). The voltage at which the channel closes is called the “pinch-off voltage”, (V_P).

JFET Channel Pinched-off



In this pinch-off region the Gate voltage, V_{GS} controls the channel current and V_{DS} has little or no effect.



JFET Model

The result is that the FET acts more like a voltage controlled resistor which has zero resistance when $V_{GS} = 0$ and maximum “ON” resistance (R_{DS}) when the Gate voltage is very negative. Under normal operating conditions, the JFET gate is always negatively biased relative to the source.

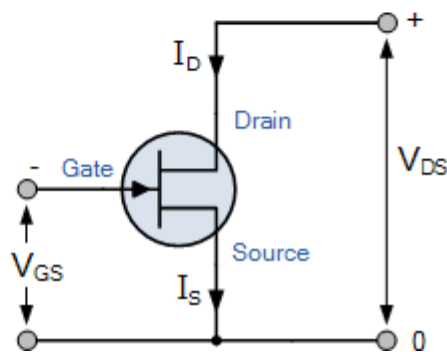
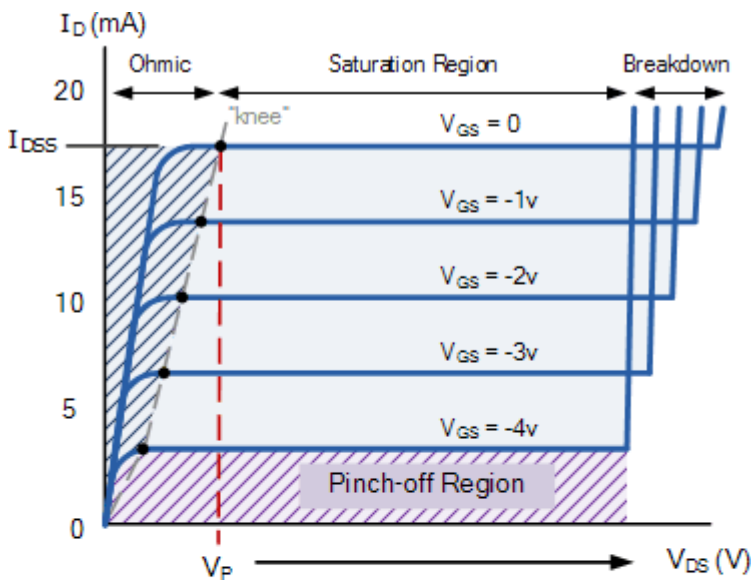
It is essential that the Gate voltage is never positive since if it is all the channel current will flow to the Gate and not to the Source, the result is damage to the JFET. Then to close the channel:

- No Gate Voltage (V_{GS}) and V_{DS} is increased from zero.
- No V_{DS} and Gate control is decreased negatively from zero.
- V_{DS} and V_{GS} varying.

The P-channel **Junction Field Effect Transistor** operates exactly the same as the N-channel above, with the following exceptions: 1). Channel current is positive due to holes, 2). The polarity of the biasing voltage needs to be reversed.

The output characteristics of an N-channel JFET with the gate short-circuited to the source is given as:

Output characteristic V-I curves of a typical junction FET



The voltage V_{GS} applied to the Gate controls the current flowing between the Drain and the Source terminals. V_{GS} refers to the voltage applied between the Gate and the Source while V_{DS} refers to the voltage applied between the Drain and the Source.

Because a **Junction Field Effect Transistor** is a voltage controlled device, **“NO current flows into the gate!”** then the Source current (I_S) flowing out of the device equals the Drain current flowing into it and therefore ($I_D = I_S$).

The characteristics curves example shown above, shows the four different regions of operation for a JFET and these are given as:

- Ohmic Region – When $V_{GS} = 0$ the depletion layer of the channel is very small and the JFET acts like a voltage controlled resistor.
- Cut-off Region – This is also known as the pinch-off region where the Gate voltage, V_{GS} is sufficient to cause the JFET to act as an open circuit as the channel resistance is at maximum.
- Saturation or Active Region – The JFET becomes a good conductor and is controlled by the Gate-Source voltage, (V_{GS}) while the Drain-Source voltage, (V_{DS}) has little or no effect.
- Breakdown Region – The voltage between the Drain and the Source, (V_{DS}) is high enough to causes the JFET's resistive channel to break down and pass uncontrolled maximum current.

The characteristics curves for a P-channel junction field effect transistor are the same as those above, except that the Drain current I_D decreases with an increasing positive Gate-Source voltage, V_{GS} .

The Drain current is zero when $V_{GS} = V_P$. For normal operation, V_{GS} is biased to be somewhere between V_P and 0. Then we can calculate the Drain current, I_D for any given bias point in the saturation or active region as follows:

Drain current in the active region.

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

Note that the value of the Drain current will be between zero (pinch-off) and I_{DSS} (maximum current). By knowing the Drain current I_D and the Drain-Source voltage V_{DS} the resistance of the channel (R_{DS}) is given as:

Drain-Source Channel Resistance.

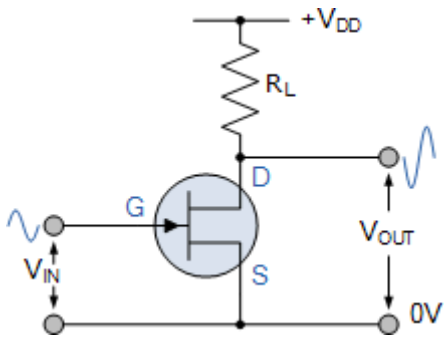
$$R_{DS} = \frac{\Delta V_{DS}}{\Delta I_D} = \frac{1}{g_m}$$

Where: g_m is the “transconductance gain” since the JFET is a voltage controlled device and which represents the rate of change of the Drain current with respect to the change in Gate-Source voltage.

Modes of FET's

Like the bipolar junction transistor, the field effect transistor being a three terminal device is capable of three distinct modes of operation and can therefore be connected within a circuit in one of the following configurations.

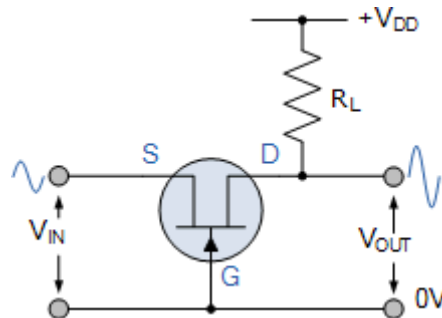
Common Source (CS) Configuration



In the **Common Source** configuration (similar to common emitter), the input is applied to the Gate and its output is taken from the Drain as shown. This is the most common mode of operation of the FET due to its high input impedance and good voltage amplification and as such Common Source amplifiers are widely used.

The common source mode of FET connection is generally used audio frequency amplifiers and in high input impedance pre-amps and stages. Being an amplifying circuit, the output signal is 180° “out-of-phase” with the input.

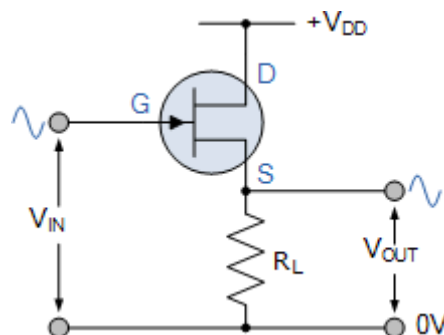
Common Gate (CG) Configuration



In the **Common Gate** configuration (similar to common base), the input is applied to the Source and its output is taken from the Drain with the Gate connected directly to ground (0V) as shown. The high input impedance feature of the previous connection is lost in this configuration as the common gate has a low input impedance, but a high output impedance.

This type of FET configuration can be used in high frequency circuits or in impedance matching circuits where a low input impedance needs to be matched to a high output impedance. The output is “in-phase” with the input.

Common Drain (CD) Configuration



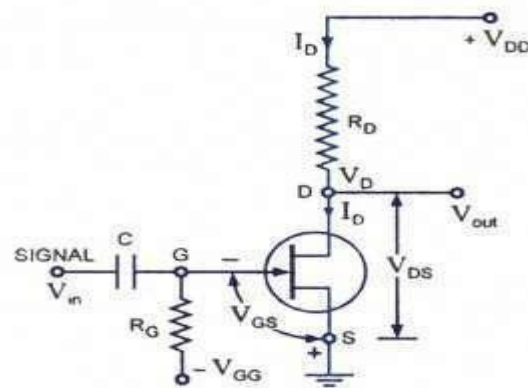
In the **Common Drain** configuration (similar to common collector), the input is applied to the Gate and its output is taken from the Source. The common drain or “source follower” configuration has a high input impedance and a low output impedance and near-unity voltage gain so is therefore used in buffer amplifiers. The voltage gain of the source follower configuration is less than unity, and the output signal is “in-phase”, 0° with the input signal.

This type of configuration is referred to as “Common Drain” because there is no signal available at the drain connection, the voltage present, $+V_{DD}$ just provides a bias. The output is in-phase with the input.

FET BIASING METHODS:

Unlike BJTs, thermal runaway does not occur with FETs. However, the wide differences in maximum and minimum **transfer characteristics** make I_D levels unpredictable with simple fixed- gate bias voltage. To obtain reasonable limits on quiescent drain currents I_D and drain-source voltage V_{DS} , source resistor and potential divider bias techniques must be used. With few exceptions, MOSFET bias circuits are similar to those used for **JFETs**. Various FET biasing circuits are discussed below

Fixed Bias:



Fixed Biasing Circuit For JFET

DC bias of a FET device needs setting of gate-source voltage V_{GS} to give desired drain current I_D

. For a JFET drain current is limited by the saturation current I_{DS} .

FET has such a high input impedance that no gate current flows and the dc voltage of the gate set by a voltage divider or a fixed battery voltage is not affected or loaded by the FET.

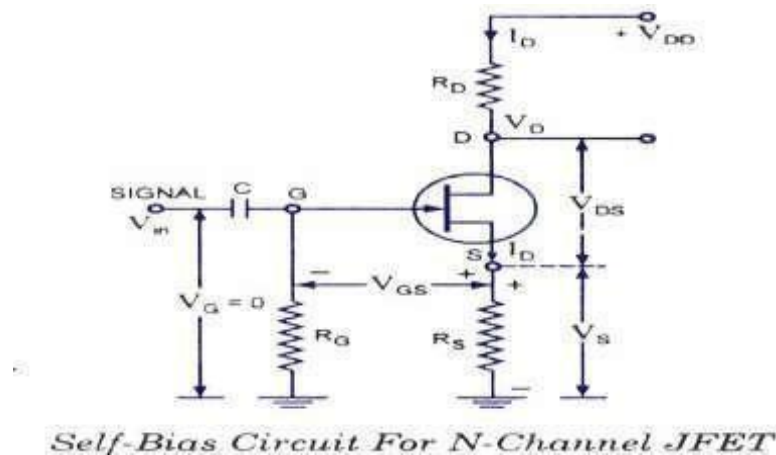
Fixed dc bias is obtained using a battery V_{GG} . This battery ensures that the gate is always negative with respect to source and no current flows through resistor R_G and gate terminal that is $I_G = 0$. The battery provides a voltage V_{GS} to bias the N-channel JFET, but no resulting current is drawn from the battery V_{GG} . Resistor R_G is included to allow any ac signal applied through capacitor C to develop across R_G . While any ac signal will develop across R_G , the dc voltage drop across R_G is equal to $I_G R_G$ i.e. 0 volt.

The gate-source voltage V_{GS} is then

$V_{GS} = -V_G - V_S = -V_{GG} - 0 = -V_{GG}$ The drain -source current I_D is then fixed by the gate-source voltage as determined by equation.

This current then causes a voltage drop across the drain resistor R_D and is given as $V_{RD} = I_D R_D$ and output voltage , $V_{out} = V_{DD} - I_D R_D$

Self-Bias:



This is the most common method for biasing a JFET. Self-bias circuit for N-channel JFET is shown in figure.

Since no gate current flows through the reverse-biased gate-source, the gate current $I_G = 0$ and, therefore, $V_G = I_G R_G = 0$. With a drain current I_D the voltage at the S is $V_S = I_D R_S$.

The gate-source voltage is then $V_{GS} = V_G - V_S = 0 - I_D R_S = -I_D R_S$.

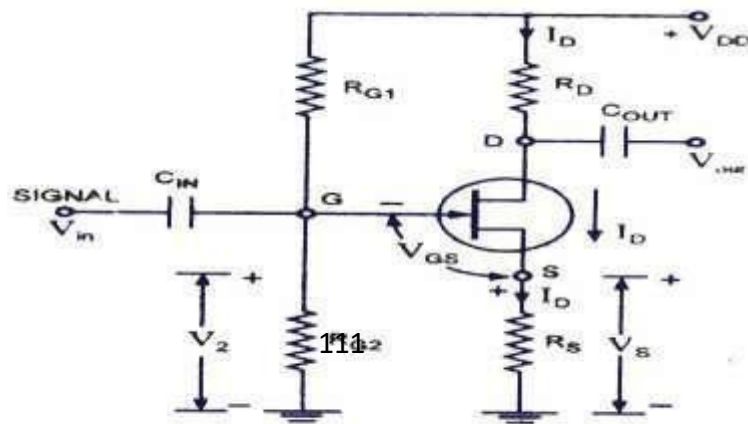
So voltage drop across resistance R_S provides the biasing voltage V_{GS} and no external source is required for biasing and this is the reason that it is called self-biasing.

The operating point (i.e zero signal I_D and V_{DS}) can easily be determined from equation and equation given below

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

Thus dc conditions of JFET amplifier are fully specified. Self biasing of a JFET stabilizes its quiescent operating point against any change in its parameters like trans conductance. Let the given JFET be replaced by another JFET having the double conductance then drain current will also try to be double but since any increase in voltage drop across R_S , therefore, gate-source voltage, V_{GS} becomes more negative and thus increase in drain current is reduced.

Potential-Divider Biasing.



*Potential-Divider Bias Circuit
For N-Channel JFET*

A slightly modified form of dc bias is provided by the circuit shown in figure. The resistors R_{G1} and R_{G2} form a potential divider across drain supply V_{DD} . The voltage V_2 across R_{G2} provides the necessary bias. The additional gate resistor R_{G1} from gate to supply voltage facilitates in larger adjustment of the dc bias point and permits use of larger valued R_S .

The gate is reverse biased so that $I_G = 0$ and gate voltage

$$V_G = V_2 = (V_{DD}/R_{G1} + R_{G2}) * R_{G2} \text{ And } V_{GS} = V_G - V_S = V_G - I_D R_S$$

The circuit is so designed that $I_D R_S$ is greater than V_2 so that V_{GS} is negative. This provides correct bias voltage.

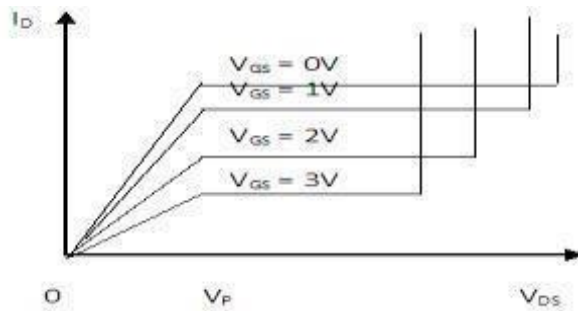
The operating point can be determined as $I_D = (V_2 - V_{GS})/R_S$ And $V_{DS} = V_{DD} - I_D (R_D + R_S)$

FET AS A VOLTAGE –VARIABLE RESISTOR (VVR):

FET is operated in the constant current portion of its output characteristics for the linear applications. In the region before pinch off, where V_{ds} is small the drain to source resistance r_d can be controlled by the bias voltage V_{gs} . The FET is useful as a voltage variable resistor (VVR) or Voltage Dependent resistor.

In JFET the drain source conductance $g_d = I_d/V_{ds}$ for small values of V_{ds} which may be expressed as $g_d = g_{do} [1 - (V_{gs}/V_p)^2]$ where g_{do} is the value of drain conductance when the bias voltage V_{gs} is zero. Small signal FET drain resistance r_d varies with applied gate voltage V_{gs} and FET act like a VARIABLE PASSIVE RESISTOR.

When $V_{DS} < V_p$, I_D V_{DS} , when V_{GS} is constant. i.e., FET acts as a resistor. In this region FET is used as a Voltage controlled resistor. Or Voltage variable resistor. Or Voltage dependant resistor.



Differences between BJT and FET

- FET is unipolar device – current I_D is due to majority (Where as BJT is Bipolar) charge carries only.
- FET is less noisy as there are no junctions(in conduction channel) FET
- FET Input impedance is very high (100 M) (due to reverse bias)
- FET is voltage controlled device, BJT is current controlled device
- FETs are easy to fabricate
- FET performance does not change much with temperature. FET has $-Ve$ temp. Coefficient, BJT has $+Ve$ temp. coefficient.
- FET has higher switching speeds
- FET is useful for small signal operation only
- BJT is cheaper than FET

SPECIAL PURPOSE DEVICES:

ZENER DIODE

A normal p-n junction diode allows electric current only in forward biased condition. When forward biased voltage is applied to the p-n junction diode, it allows large amount of electric current and blocks only a small amount of electric current. Hence, a forward biased p-n junction diode offer only a small resistance to the electric current.

When reverse biased voltage is applied to the p-n junction diode, it blocks large amount of electric current and allows only a small amount of electric current. Hence, a reverse biased p-n junction diode offer large resistance to the electric current.

If reverse biased voltage applied to the p-n junction diode is highly increased, a sudden rise in current occurs. At this point, a small increase in voltage will rapidly increases the electric current. This sudden rise in electric current causes a junction breakdown called zener or avalanche breakdown. The voltage at which zener breakdown occurs is called zener voltage and the sudden increase in current is called zener current.

A normal p-n junction diode does not operate in breakdown region because the excess current permanently damages the diode. Normal p-n junction diodes are not designed to operate in reverse breakdown region. Therefore, a normal p-n junction diode does not operate in reverse breakdown region.

What is zener diode?

A zener diode is a special type of device designed to operate in the zener breakdown region. Zener diodes act like normal p-n junction diodes under forward biased condition. When forward biased voltage is applied to the zener diode it allows large amount of electric current and blocks only a small amount of electric current.

Zener diode is heavily doped than the normal p-n junction diode. Hence, it has very thin depletion region. Therefore, zener diodes allow more electric current than the normal p-n junction diodes.

Zener diode allows electric current in forward direction like a normal diode but also allows electric current in the reverse direction if the applied reverse voltage is greater than the zener voltage. Zener diode is always connected in reverse direction because it is specifically designed to work in reverse direction.

Zener diode definition

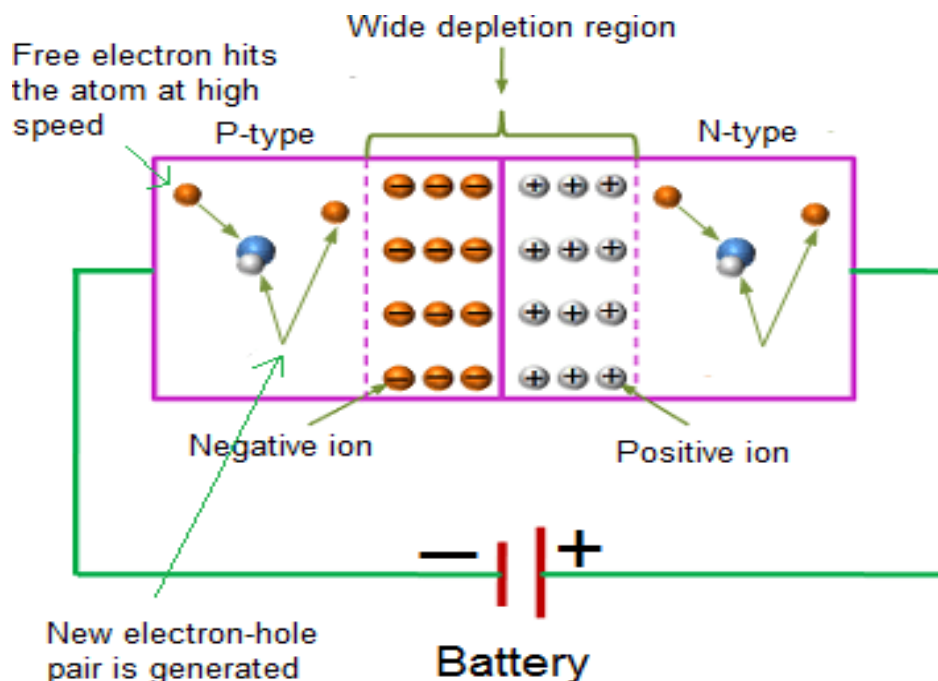
A zener diode is a p-n junction semiconductor device designed to operate in the reverse breakdown region. The breakdown voltage of a zener diode is carefully set by controlling the doping level during manufacture.

The name zener diode was named after the American physicist Clarence Melvin Zener who discovered the zener effect. Zener diodes are the basic building blocks of electronic circuits. They are widely used in all kinds of electronic equipments. Zener diodes are mainly used to protect electronic circuits from over voltage.

Breakdown in zener diode

There are two types of reverse breakdown regions in a zener diode: avalanche breakdown and zener breakdown.

Avalanche breakdown The avalanche breakdown occurs in both normal diodes and zener diodes at high reverse voltage. When high reverse voltage is applied to the p-n junction diode, the free electrons (minority carriers) gain large amount of energy and accelerated to greater velocities.

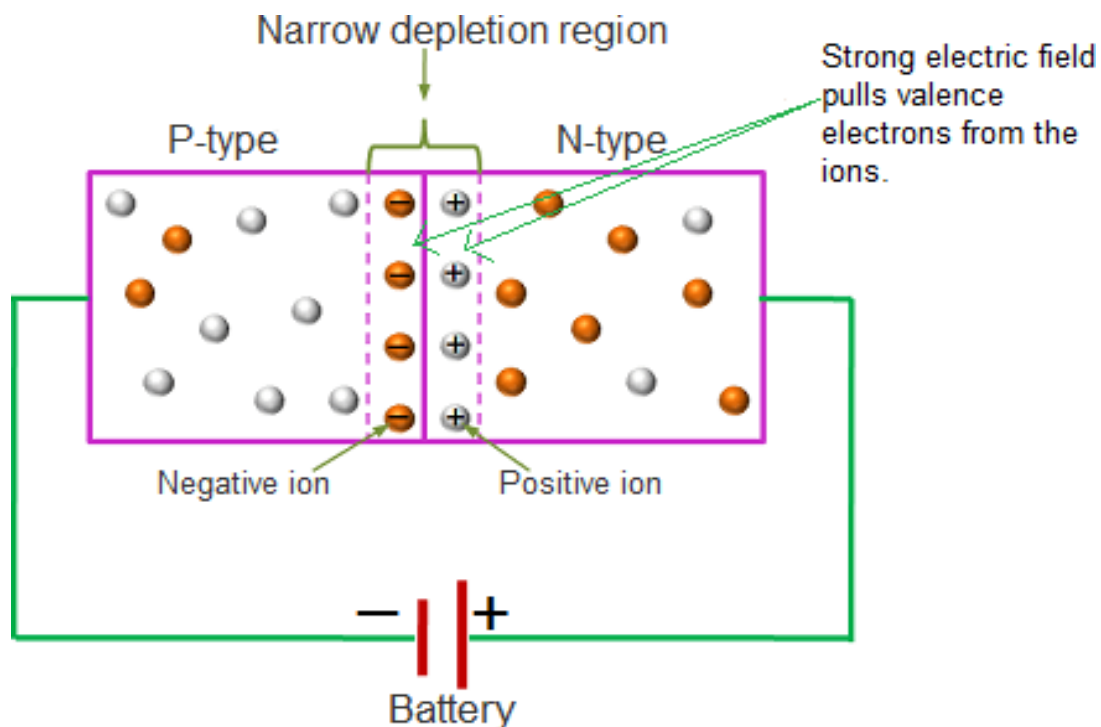


The free electrons moving at high speed will collide with the atoms and knock off more

electrons. These electrons are again accelerated and collide with other atoms. Because of this continuous collision with the atoms, a large number of free electrons are generated. As a result, electric current in the diode increases rapidly. This sudden increase in electric current may permanently destroys the normal diode. However, avalanche diodes may not be destroyed because they are carefully designed to operate in avalanche breakdown region. Avalanche breakdown occurs in zener diodes with zener voltage (V_z) greater than 6V.

Zener breakdown :

The zener breakdown occurs in heavily doped p-n junction diodes because of their narrow depletion region. When reverse biased voltage applied to the diode is increased, the narrow depletion region generates strong electric field.



When reverse biased voltage applied to the diode reaches close to zener voltage, the electric field in the depletion region is strong enough to pull electrons from their valence band. The valence electrons which gains sufficient energy from the strong electric field of depletion region will breaks bonding with the parent atom. The valance electrons which break bonding with parent atom will become free electrons. This free electrons carry electric current from one place to another place. At zener breakdown region, a small increase in voltage will rapidly increases the electric current.

Zener breakdown occurs at low reverse voltage whereas avalanche breakdown occurs at high reverse voltage.

Zener breakdown occurs in zener diodes because they have very thin depletion region.

Breakdown region is the normal operating region for a zener diode.

Zener breakdown occurs in zener diodes with zener voltage (V_z) less than 6V.

Symbol of Zener diode

The symbol of zener diode is shown in below figure. Zener diode consists of two terminals: cathode and anode.

Zener diode symbol



In zener diode, electric current flows from both anode to cathode and cathode to anode.

The symbol of zener diode is similar to the normal p-n junction diode, but with bend edges on the vertical bar.

VI characteristics of zener diode

The VI characteristics of a zener diode is shown in the below figure. When forward biased voltage is applied to the zener diode, it works like a normal diode. However, when reverse biased voltage is applied to the zener diode, it works in different manner.

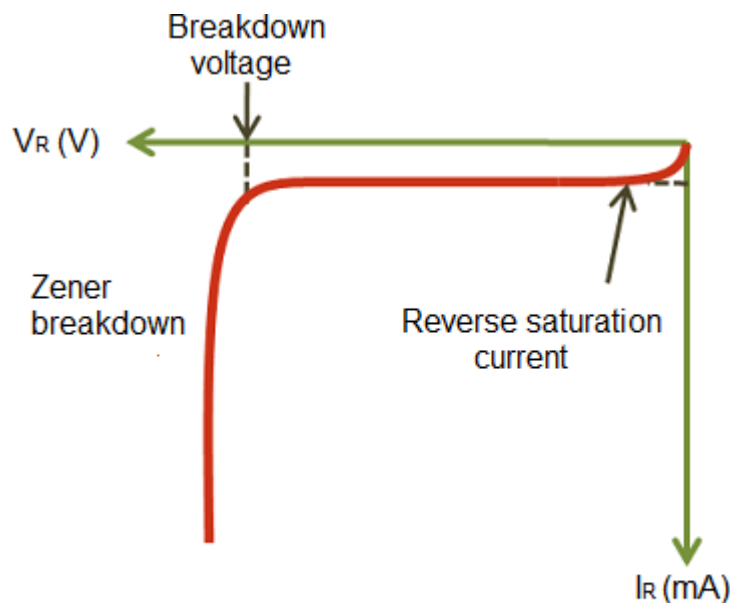


Fig: Zener breakdown

Copyright © Physics and Radio-Electronics, All rights reserved

When reverse biased voltage is applied to a zener diode, it allows only a small amount of leakage current until the voltage is less than zener voltage. When reverse biased voltage applied to the zener diode reaches zener voltage, it starts allowing large amount of electric current. At this point, a small increase in reverse voltage will rapidly increases the electric current. Because of this sudden rise in electric current, breakdown occurs called zener breakdown. However, zener diode exhibits a controlled breakdown that does damage the device.

The zener breakdown voltage of the zener diode is depends on the amount of doping applied. If the diode is heavily doped, zener breakdown occurs at low reverse volt. On the other hand,

if the diode is lightly doped, the zener breakdown occurs at high reverse voltages. Zener diodes are available with zener voltages in the range of 1.8V to 400V.

Advantages of zener diode

Power dissipation capacity is very high

High accuracy

Small size

Low cost

Applications of zener diode

- It is normally used as voltage reference
- Zener diodes are used in voltage stabilizers or shunt regulators.
- Zener diodes are used in switching operations
- Zener diodes are used in clipping and clamping circuits.
- Zener diodes are used in various protection circuits

Zener diode
as
voltage regulator

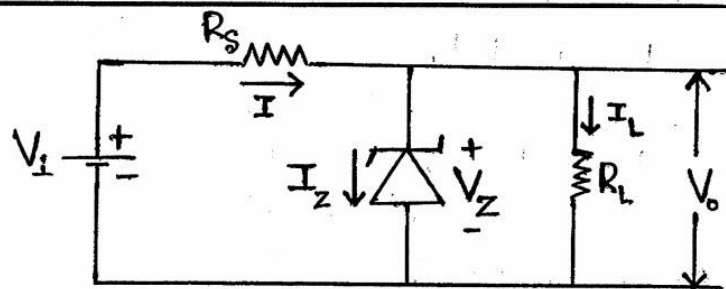


Fig ① : **Zener diode voltage regulator**

- * The Zener diode is called as voltage-regulated diode because it maintains a constant o/p voltage even though the current through it changes.
- * The i/p voltage ' V_i ' is the unregulated (varying) dc voltage. The Zener diode is used in reverse biased region. Under this condition, the current through the diode is very small i.e. in terms of μA .
- * When sufficient reverse voltage is applied, the Zener breakdown occurs (Zener diode conducts) & the large current flows through the diode.
- * The voltage at which Zener diode conducts is called Zener voltage ' V_Z '. Under this condition, the voltage across the Zener is constant & equal to V_Z . As V_Z is connected across the load & hence the load voltage V_o is equal to the Zener voltage V_Z .

Thus Zener diode acts as a voltage regulator.

$$\therefore V_o = V_Z$$

* From Fig ①, the current

$$I = I_Z + I_L$$

$$I_Z = I - I_L \rightarrow \textcircled{1}$$

* The current 'I' is given by

$$I = \frac{V_i - V_o}{R} \rightarrow \textcircled{2}$$

Substituting eq ② in eq ①, we get

$$I_Z = \frac{V_i - V_o}{R} - I_L$$

$$I_Z + I_L = \frac{V_i - V_o}{R}$$

$$R = \frac{V_i - V_o}{I_Z + I_L}$$

* The load resistance

$$R_L = \frac{V_o}{I_L}$$

* The power

$$P_{Z_{max}} = I_{Z_{max}} V_Z$$

❖ What is a voltage regulator? Why is it necessary?

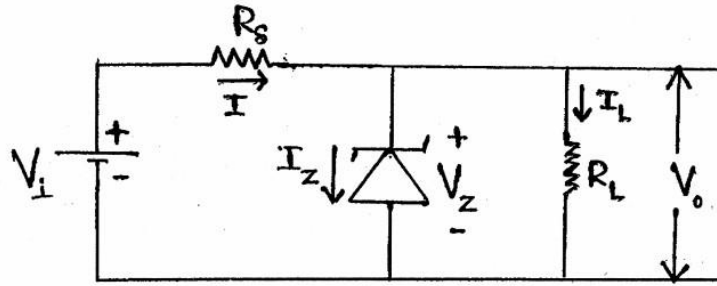


Fig ①: Zener diode voltage Regulator

- * A voltage regulator is a circuit which accepts unregulated dc as I/p & provides a constant dc o/p voltage irrespective of changes in the line voltage & the load current.
- * Most of the electronic circuits require a stable dc voltage for their proper operation. Hence it is necessary to regulate the voltage before giving to electronic circuits. Fig ① Shows the example of regulated dc power supply.

FORMULAE

1> o/p voltage $V_o = \text{Zener voltage } V_Z$ i.e. $V_o = V_Z$

2> The Current in the CKT is given by

$$I = I_Z + I_L$$

$$I_{Z\max} = I_Z$$

3> Zener Current $I_Z = I - I_L$

$$4> V_i = IR + V_o$$

$$5> I = \frac{V_i - V_o}{R}$$

$$6> R = \frac{V_i - V_o}{I_Z + I_L}$$

$$7> i) R_{\min} = \frac{V_{i\max} - V_o}{I_{Z\max} + I_{L\min}}$$

$$ii) R_{\max} = \frac{V_{i\min} - V_o}{I_{Z\min} + I_{L\max}}$$

$$8> R_L = \frac{V_o}{I_L}$$

$$9> P_{Z\max} = I_{Z\max} V_Z$$

$$P_o = P_{Z\max}$$

$$10> I_{Z\max} = \frac{P_{Z\max}}{V_Z}$$

$$11> \text{F81 } V_{i(\min)} : I_Z = I_{Z(\min)}$$

$$I = I_{Z\min} + I_L$$

$$V_{i(\min)} = IR + V_o$$

$$12> \text{F81 } V_{i(\max)} :$$

$$I_Z = I_{Z(\max)}$$

$$I = I_{Z(\max)} + I_L$$

$$V_{i(\max)} = IR + V_o$$

Silicon Controlled Rectifier

We know that the diode allows electric current in one direction and blocks electric current in another direction. In other words, the diode converts the AC current into DC current. This unique behavior of the diodes makes it possible to build different types of rectifiers such as half wave, full wave and bridge rectifiers. These rectifiers convert the Alternating Current into Direct Current.

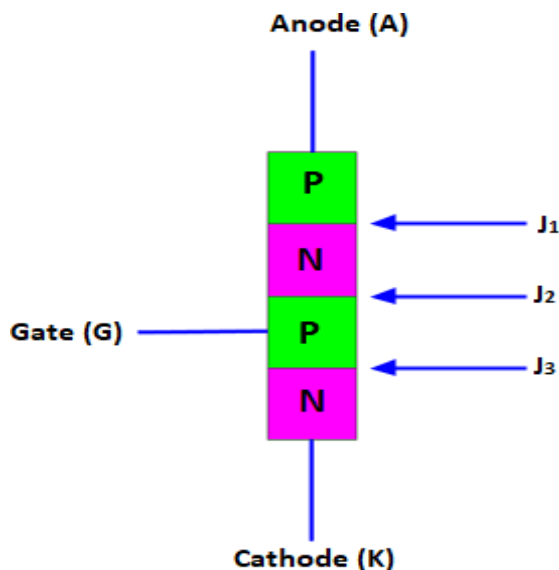
The half wave, full wave, and bridge rectifiers use normal p-n junction diodes (two layer diodes). So if the voltage applied to these diodes is high enough, then the diodes may get destroyed. So the rectifiers cannot operate at high voltages.

To overcome these drawbacks, scientists have developed a special type of rectifier known as Silicon Controlled Rectifier. These rectifiers can withstand high voltages.

Silicon Controlled Rectifier Definition

A Silicon Controlled Rectifier is a 3 terminal and 4 layer semiconductor current controlling device. It is mainly used in the devices for the control of high power. Silicon controlled rectifier is also sometimes referred to as SCR diode, 4-layer diode, 4-layer device, or Thyristor. It is made up of a silicon material which controls high power and converts high AC current into DC current (rectification). Hence, it is named as silicon controlled rectifier.

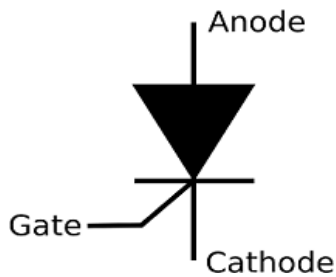
What is Silicon Controlled Rectifier?



Silicon controlled rectifier is a unidirectional current controlling device. Just like a normal p-n junction diode, it allows electric current in only one direction and blocks electric current in another direction. A normal p-n junction diode is made of two semiconductor layers namely P-type and N-type. However, a SCR diode is made of 4 semiconductor layers of alternating P and N type materials.

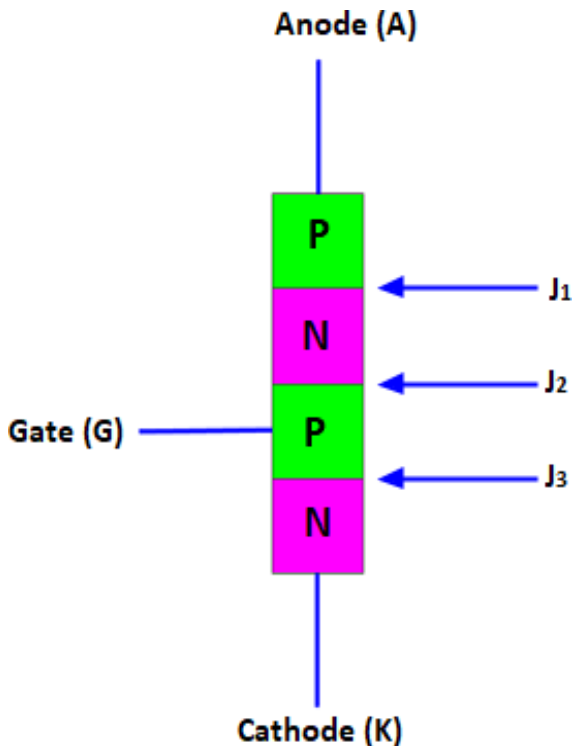
Silicon controlled rectifiers are used in power control applications such as power delivered to electric motors, relay controls or induction heating elements where the power delivered has to be controlled.

Silicon Controlled Rectifier Symbol



The schematic symbol of a silicon controlled rectifier is shown in the below figure. A SCR diode consists of three terminals namely anode (A), cathode (K), Gate (G). The diode arrow represents the direction of conventional current.

Construction of Silicon Controlled Rectifier



A silicon controlled rectifier is made up of 4 semiconductor layers of alternating P and N type materials, which forms NPNP or PNP structures. It has three P-N junctions namely J_1 , J_2 , J_3 with three terminals attached to the semiconductors materials namely anode (A), cathode (K), and gate (G). Anode is a positively charged electrode through which the conventional current enters into an electrical device, cathode is a negatively charged electrode through which the conventional current leaves an electrical device, gate is a terminal that controls the flow of current between anode and cathode. The gate terminal is also sometimes referred to as control terminal.

The anode terminal of SCR diode is connected to the first p-type material of a PNP structure, cathode terminal is connected to the last n-type material, and gate terminal is connected to the second p-type material of a PNP structure which is nearest to the cathode.

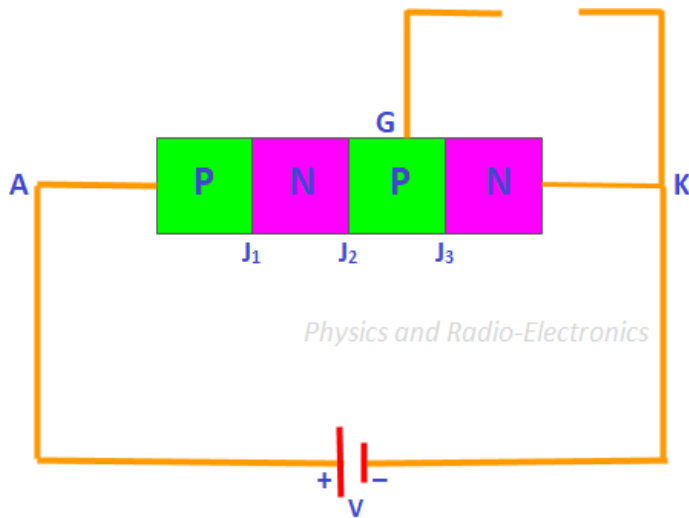
In silicon controlled rectifier, silicon is used as an intrinsic semiconductor. When pentavalent impurities are added to this intrinsic semiconductor, an N-type semiconductor is formed. When trivalent impurities are added to an intrinsic semiconductor, a p-type semiconductor is formed.

When 4 semiconductor layers of alternating P and N type materials are placed one over another (hypothetically), three junctions are formed in PNP structure. In a PNP structure, the junction J_1 is formed between the first P-N layer, the junction J_2 is formed between the N-P layer and the junction J_3 is formed between the last P-N layer. The doping of PNP structure is depends on the application of SCR diode

Modes of Operation in SCR

There are three modes of operation for a Silicon Controlled Rectifier (SCR), depending upon the biasing given to it.

1) Forward Blocking Mode (Off State)



Forward Blocking Mode of SCR

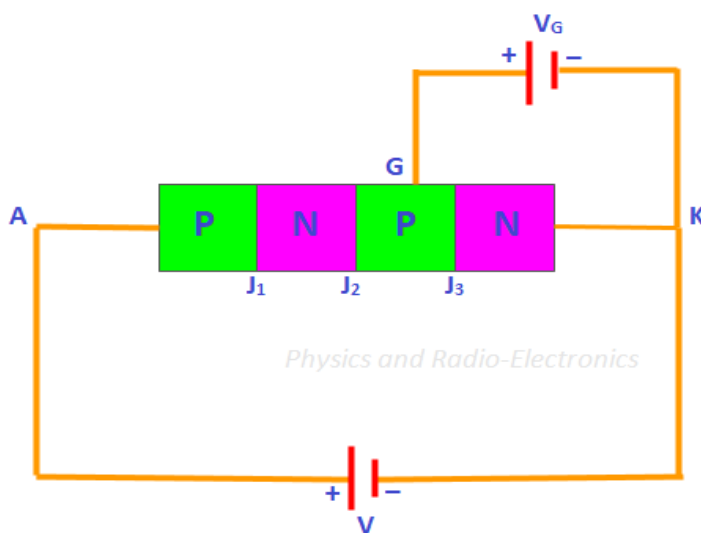
current flows between junction J_1 and junction J_3 .

In this mode of operation, the positive voltage (+) is given to anode A (+), negative voltage (-) is given to cathode K (-), and gate G is open circuited as shown in the below figure. In this case, the junction J_1 and junction J_3 are forward biased whereas the junction J_2 becomes reverse biased. Due to the reverse bias voltage, the width of depletion region increases at junction J_2 . This depletion region at junction J_2 acts as a wall or obstacle between the junction J_1 and junction J_3 . It blocks the current flowing between junction J_1 and junction J_3 . Therefore, the majority of the current does not flow between junction J_1 and junction J_3 . However, a small amount of leakage

When the voltage applied to the SCR reaches a breakdown value, the high energy minority carriers causes avalanche breakdown. At this breakdown voltage, current starts flowing through the SCR. But below this breakdown voltage, the SCR offers very high resistance to the current and so it will be in off state.

In this mode of operation, SCR is forward biased but still current does flows through it. Hence, it is named as Forward Blocking Mode.

2) Forward Conducting Mode (On State)



Forward Conducting Mode of SCR

The Silicon Controlled Rectifier can be made to conduct in two ways:

- By increasing the forward bias voltage applied between anode and cathode beyond the breakdown voltage
- By applying positive voltage at gate terminal.

In the first case, the forward bias voltage applied between anode and cathode is increased beyond the breakdown voltage, the minority

carriers (free electrons in anode and holes in cathode) gains large amount of energy and accelerated to greater velocities. This high speed minority carriers collides with other atoms and generates more charge carriers. Likewise, many collisions happens with other atoms. Due to this, millions of charge carriers are generated. As a result depletion region breakdown occurs at junction J_2 and current starts flowing through the SCR. So the SCR is in On state. The

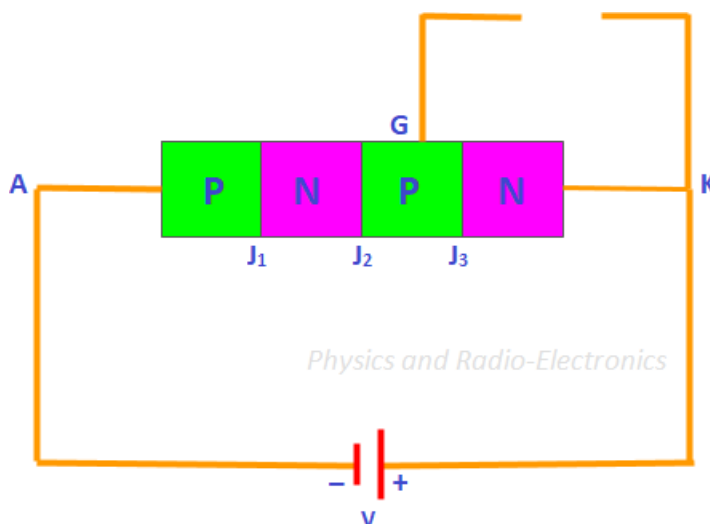
current flow in the SCR increases rapidly after junction breakdown occurs.

In the second case, a small positive voltage V_G is applied to the gate terminal. As we know that, in forward blocking mode, current does not flow through the circuit because of the wide depletion region present at the junction J_2 . This depletion region was formed because of the reverse biased gate terminal. So this problem can be easily solved by applying a small positive voltage at the Gate terminal. When a small positive voltage is applied to the gate terminal, it will become forward biased. So the depletion region width at junction J_2 becomes very narrow. Under this condition, applying a small forward bias voltage between anode and cathode is enough for electric current to penetrate through this narrow depletion region. Therefore, electric current starts flowing through the SCR circuit.

In second case, we no need to apply large voltage between anode and cathode. A small voltage between anode and cathode, and positive voltage to gate terminal is enough to brought SCR from blocking mode to conducting mode.

In this mode of operation, SCR is forward biased and current flows through it. Hence, it is named as Forward Conducting Mode.

3) Reverse Blocking Mode (On State)



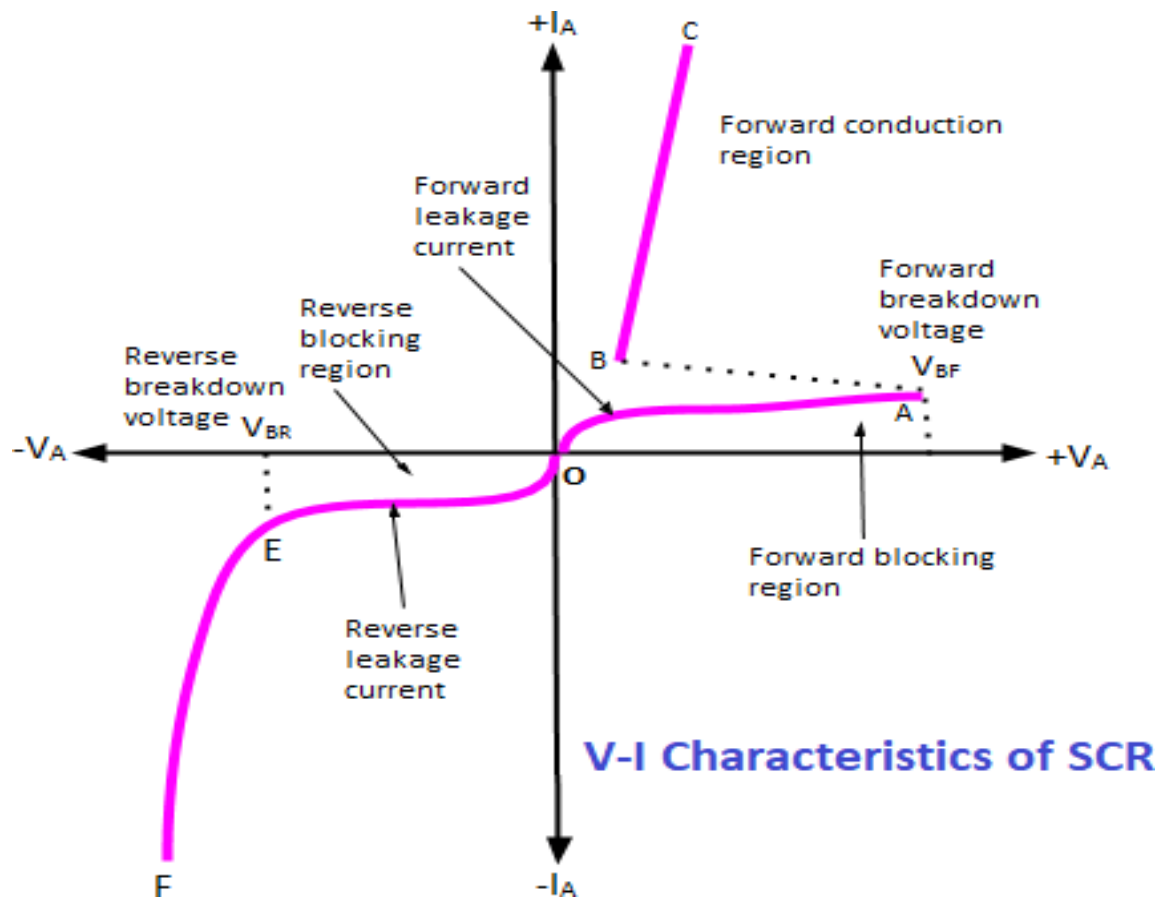
Reverse Blocking Mode of SCR

junction J_2 . This small leakage current is not enough to turn on the SCR. So the SCR will be in Off state.

In this mode of operation, the negative voltage (-) is given to anode (+), positive voltage (+) is given to cathode (-), and gate is open circuited as shown in the below figure. In this case, the junction J_1 and junction J_3 are reverse biased whereas the junction J_2 becomes forward biased.

As the junctions J_1 and junction J_3 are reverse biased, no current flows through the SCR circuit. But a small leakage current flows due to drift of charge carriers in the forward biased

V-I Characteristics of SCR



V_A = Anode voltage, I_A = Anode current, $+V_A$ = Forward anode voltage, $+I_A$ = Forward anode current, $-V_A$ = Reverse anode voltage, $-I_A$ = Reverse anode current

The V-I characteristics of SCR is divided into three regions:

Forward blocking region

In this region, the positive voltage (+) is given to anode (+), negative voltage (-) is given to cathode (-), and gate is open circuited. Due to this the junction J_1 and J_3 become forward biased while J_2 become reverse biased. Therefore, a small leakage current flows from anode to cathode terminals of the SCR. This small leakage current is known as forward leakage current.

The region OA of V-I characteristics is known as forward blocking region in which the SCR does not conduct electric current.

Forward Conduction region

If the forward bias voltage applied between anode and cathode is increased beyond the breakdown voltage, the minority carriers (free electrons in anode and holes in cathode) gain a large amount of energy and are accelerated to greater velocities. This high speed minority carrier collides with other atoms and generates more charge carriers. Likewise, many collisions happen with atoms. Due to this, millions of charge carriers are generated. As a result depletion region breakdown occurs at junction J_2 and current starts flowing through the SCR. So the SCR will be in On state. The current flow in the SCR increases rapidly after junction breakdown occurs.

The voltage at which the junction J_2 gets broken when the gate is open is called forward breakdown voltage (V_{BF}).

The region BC of the V-I characteristics is called conduction region. In this region, the current flowing from anode to cathode increases rapidly. The region AB indicates that as soon as the device becomes on, the voltage across the SCR drops to some volts.

Reverse Blocking Region

In this region, the negative voltage (-) is given to anode (+), positive voltage (+) is given to cathode (-), and gate is open circuited. In this case, the junction J_1 and junction J_3 are reverse biased whereas the junction J_2 becomes forward biased.

As the junctions J_1 and junction J_3 are reverse biased, no current flows through the SCR circuit. But a small leakage current flows due to drift of charge carriers in the forward biased junction J_2 . This small leakage current is called reverse leakage current. This small leakage current is not sufficient to turn on the SCR.

If the reverse bias voltage applied between anode and cathode is increased beyond the reverse breakdown voltage (V_{BR}), an avalanche breakdown occurs. As a result, the current increases rapidly. The region EF is called reverse avalanche region. This rapid increase in current may damage the SCR device.

Tunnel diode

Definition: 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



Tunnel diode symbol

What is a tunnel diode?

Tunnel diode was invented in 1958 by Leo Esaki.

Leo Esaki observed that if a semiconductor diode is heavily doped with impurities, it will exhibit negative resistance. Negative resistance means the current across the tunnel diode decreases when the voltage increases.

In 1973 Leo Esaki received the Nobel Prize in physics for discovering the electron tunneling effect used in these diodes.

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

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

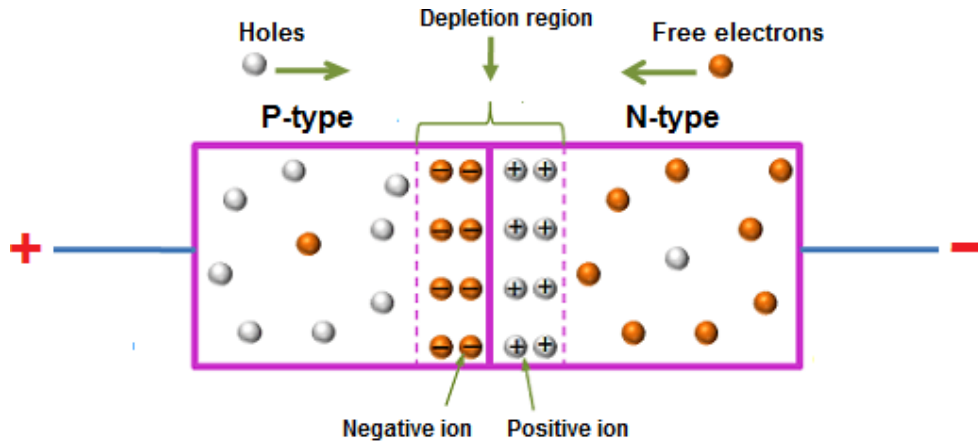
The depletion region or depletion layer in a p-n junction diode is made up of positive ions and negative ions. Because of these positive and negative ions, there exists a built-in-potential or electric field in the depletion region. This electric field in the depletion region exerts electric force in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

Another thing we need to remember is that 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.

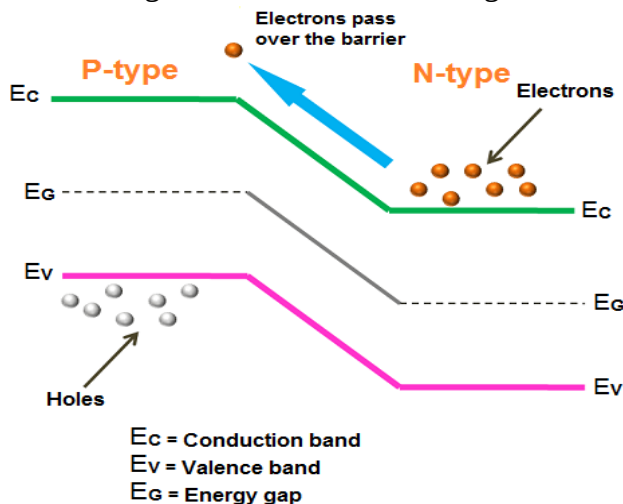
Electric current in ordinary p-n junction diode

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.

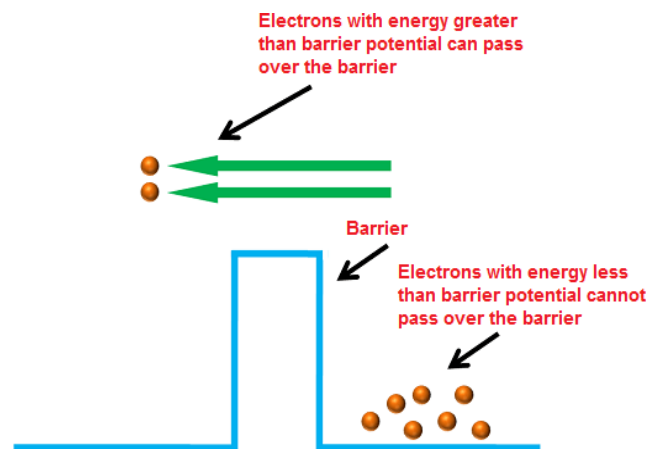


Forward bias p-n junction diode

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



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.

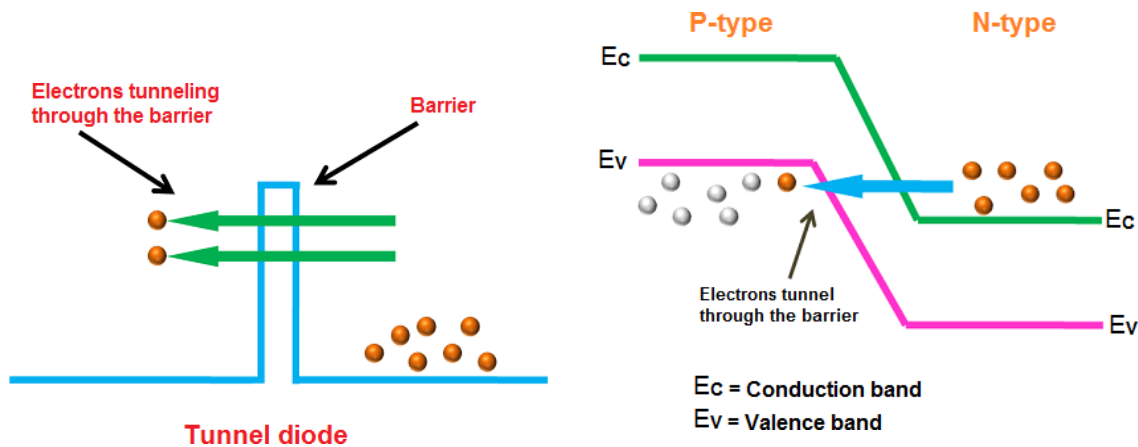
Quantum mechanics says that the electrons will directly penetrate through the depletion layer or barrier if the depletion width is very small.

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.

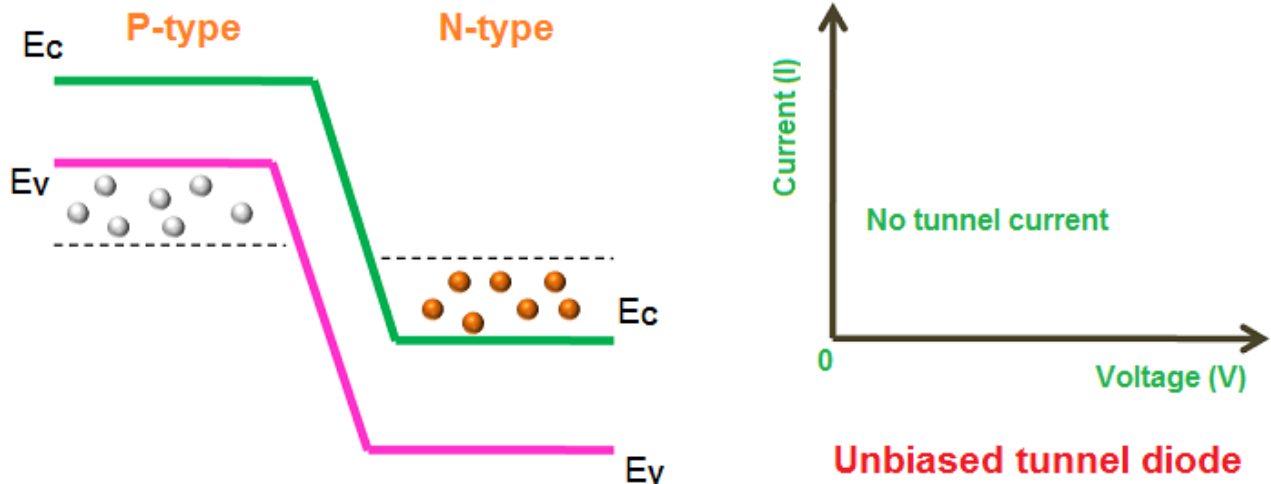
In tunnel diodes, the electrons need not overcome the opposing force from the depletion layer to produce electric current. The electrons can directly tunnel from the conduction band of n-region into the valence band of p-region. Thus, electric current is produced in tunnel diode.



How tunnel diode works?

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.

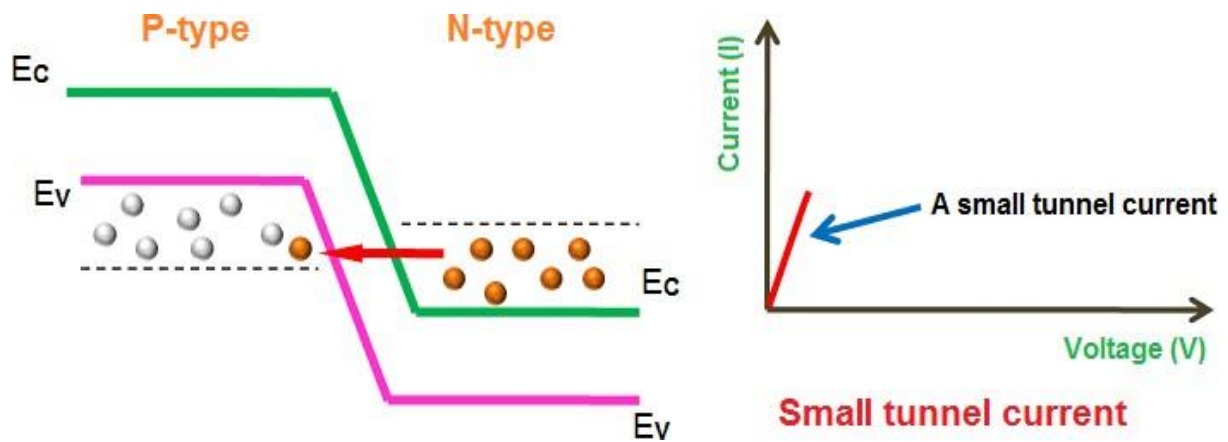
However, the net current flow will be zero because an equal number of charge carriers (free electrons and holes) flow in opposite directions.

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.

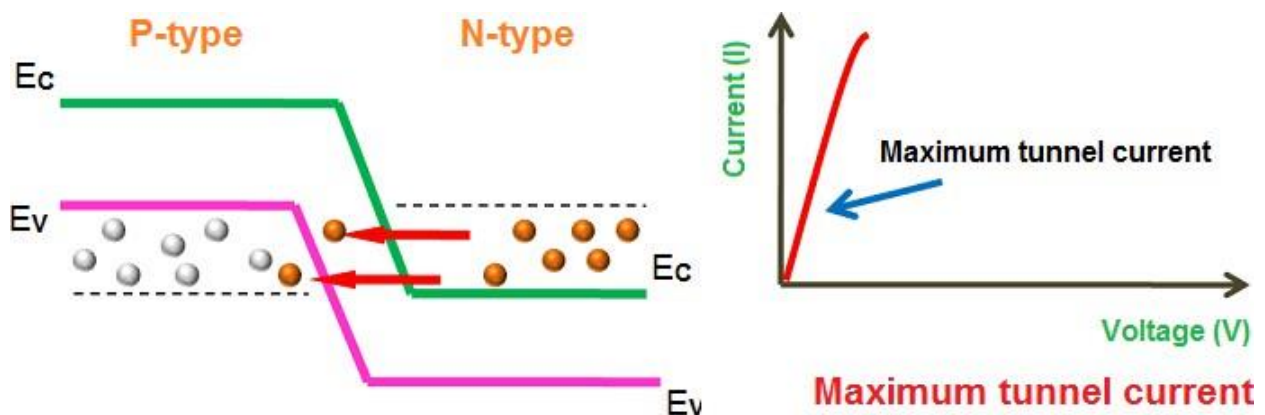
However, a small number of electrons in the conduction band of the n-

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

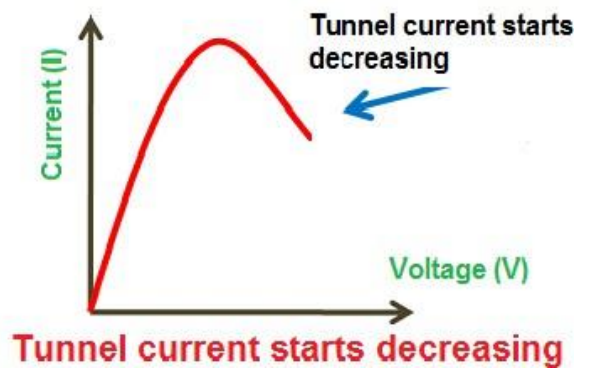
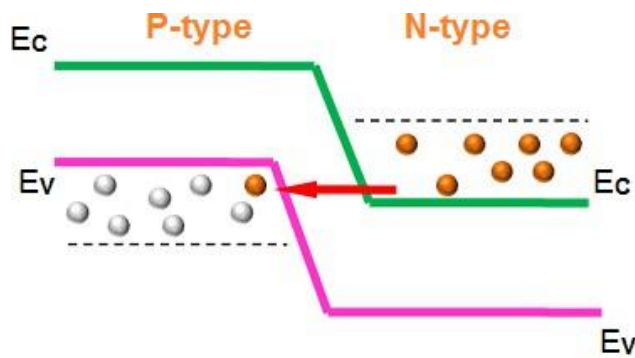
When the voltage applied to the tunnel diode is slightly increased, a large number of free electrons at n-side and holes at p-side are generated. Because of the increase in voltage, the overlapping of the conduction band and valence band is increased.



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.

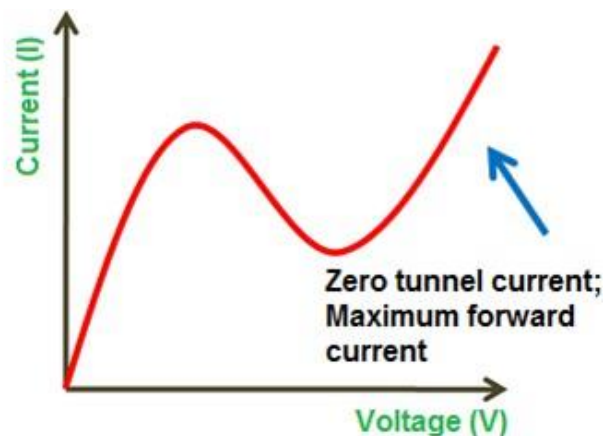
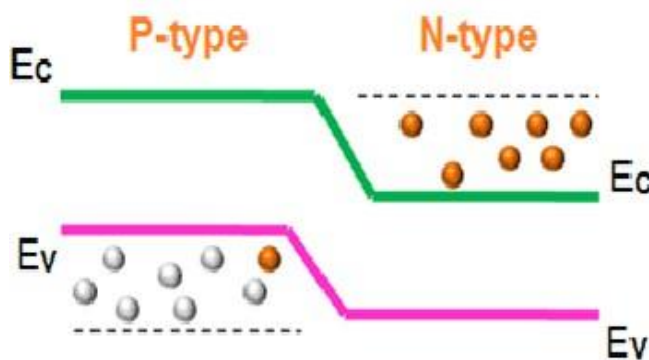
Step 4: Applied voltage is further increased

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 n-region to the valence band of 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.



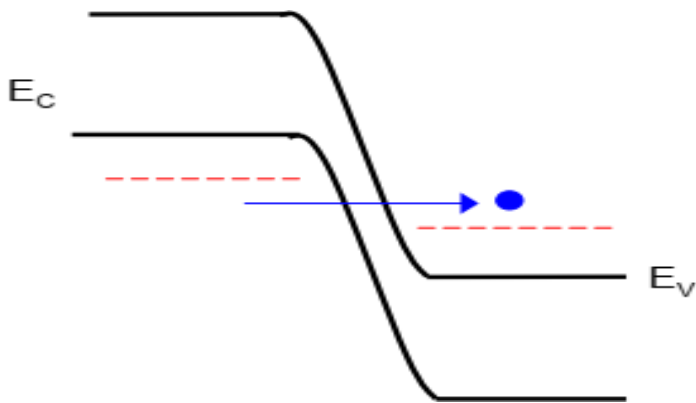
If this applied voltage is greater than the built-in potential of the depletion layer, the regular forward current starts flowing through the tunnel diode.

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.

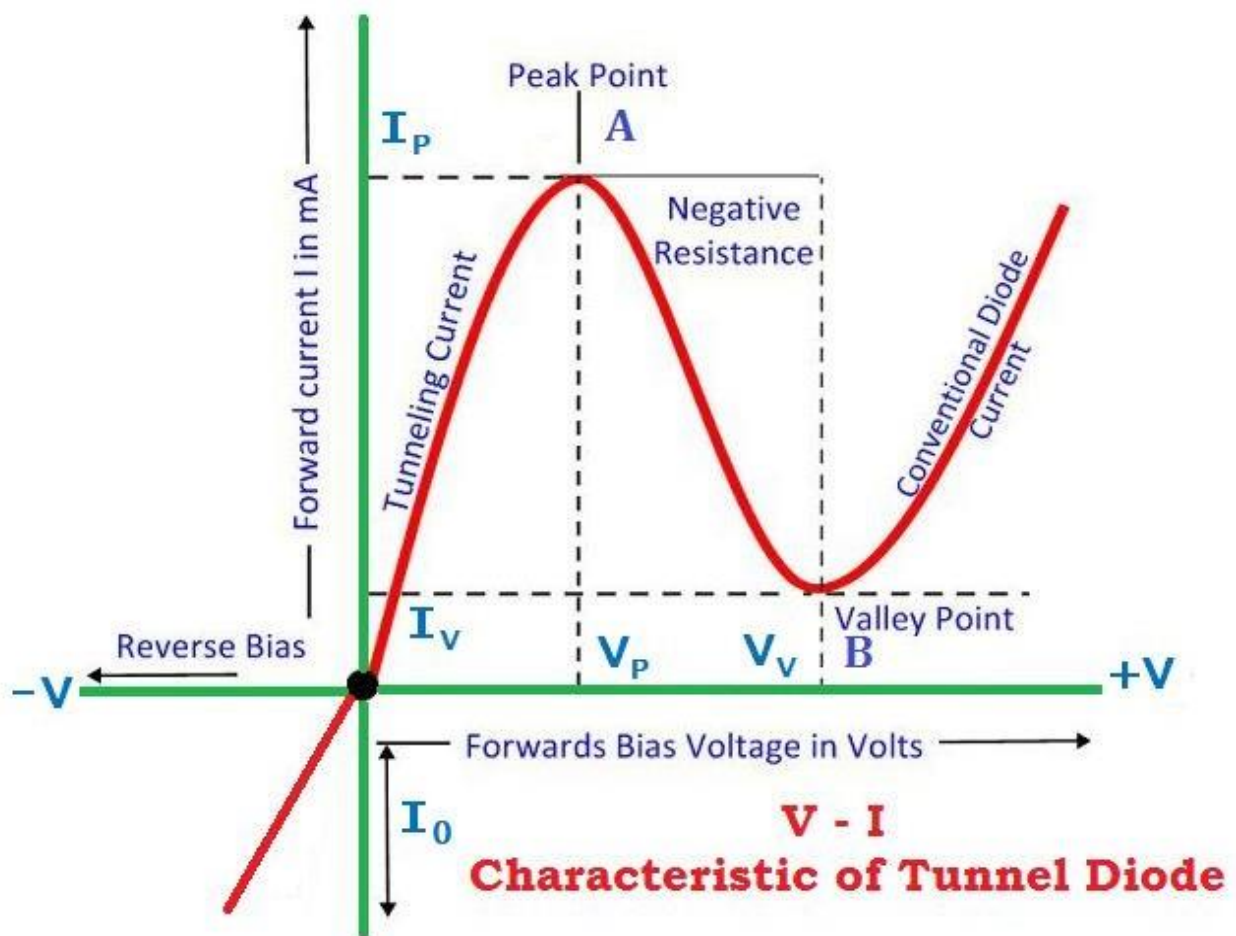
A tunnel diode operating in the negative resistance region can be used as an amplifier or an oscillator.

Under Reverse Bias

In this case the, electrons in the valence band of the p side tunnel directly towards the empty states present in the conduction band of the n side creating large tunneling current which increases with the application of reverse voltage.



V-I Characteristics of Tunnel Diode



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.

UNIUNCTION TRANSISTOR(UJT)

The **Unijunction Transistor** or **UJT** for short, is another solid state three terminal device that can be used in gate pulse, timing circuits and trigger generator applications to switch and control either thyristors and triac's for AC power control type applications.

Like diodes, unijunction transistors are constructed from separate P-type and N-type semiconductor materials forming a single (hence its name Uni-Junction) PN-junction within the main conducting N-type channel of the device.

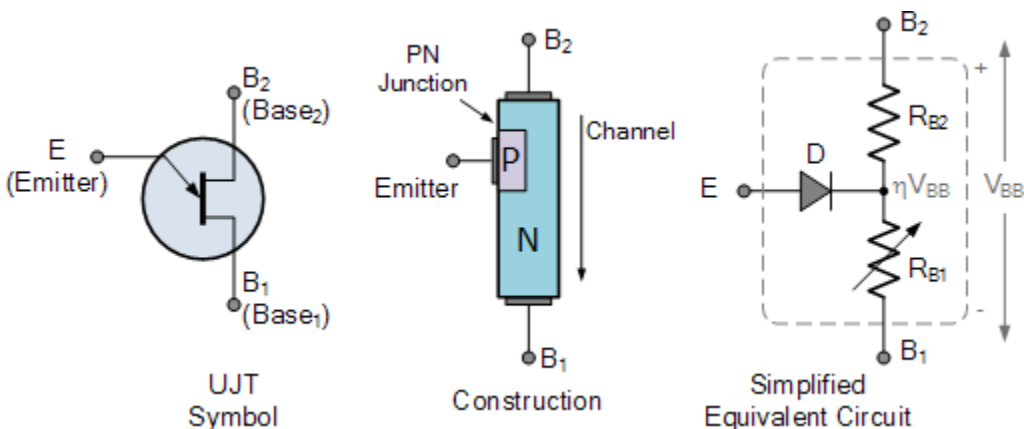
Although the *Unijunction Transistor* has the name of a transistor, its switching characteristics are very different from those of a conventional bipolar or field effect transistor as it can not be used to amplify a signal but instead is used as a ON-OFF switching transistor. UJT's have unidirectional conductivity and negative impedance characteristics acting more like a variable voltage divider during breakdown.

Like N-channel FET's, the UJT consists of a single solid piece of N-type semiconductor material forming the main current carrying channel with its two outer connections marked as *Base 2* (B_2) and *Base 1* (B_1). The third connection, confusingly marked as the *Emitter* (E) is located along the channel. The emitter terminal is represented by an arrow pointing from the P-type emitter to the N-type base.

The Emitter rectifying p-n junction of the unijunction transistor is formed by fusing the P-type material into the N-type silicon channel. However, P-channel UJT's with an N-type Emitter terminal are also available but these are little used.

The Emitter junction is positioned along the channel so that it is closer to terminal B_2 than B_1 . An arrow is used in the UJT symbol which points towards the base indicating that the Emitter terminal is positive and the silicon bar is negative material. Below shows the symbol, construction, and equivalent circuit of the UJT.

Unijunction Transistor Symbol and Construction



Notice that the symbol for the unijunction transistor looks very similar to that of the junction field effect transistor or JFET, except that it has a bent arrow representing the Emitter (E) input. While similar in respect of their ohmic channels, JFET's and UJT's operate very differently and should not be confused.

So how does it work? We can see from the equivalent circuit above, that the N-type channel basically consists of two resistors R_{B2} and R_{B1} in series with an equivalent (ideal) diode, D representing the p-n junction connected to their center point. The Emitter p-n junction

is fixed in position along the ohmic channel during manufacture and can therefore not be changed.

Resistance R_{B1} is given between the Emitter, E and terminal B_1 , while resistance R_{B2} is given between the Emitter, E and terminal B_2 . As the physical position of the p-n junction is closer to terminal B_2 than B_1 the resistive value of R_{B2} will be less than R_{B1} .

The total resistance of the silicon bar (its Ohmic resistance) will be dependent upon the semiconductors actual doping level as well as the physical dimensions of the N-type silicon channel but can be represented by R_{BB} . If measured with an ohmmeter, this static resistance would typically measure somewhere between about $4k\Omega$ and $10k\Omega$'s for most common UJT's such as the 2N1671, 2N2646 or the 2N2647.

These two series resistances produce a voltage divider network between the two base terminals of the unijunction transistor and since this channel stretches from B_2 to B_1 , when a voltage is applied across the device, the potential at any point along the channel will be in proportion to its position between terminals B_2 and B_1 . The level of the voltage gradient therefore depends upon the amount of supply voltage.

When used in a circuit, terminal B_1 is connected to ground and the Emitter serves as the input to the device. Suppose a voltage V_{BB} is applied across the UJT between B_2 and B_1 so that B_2 is biased positive relative to B_1 . With zero Emitter input applied, the voltage developed across R_{B1} (the lower resistance) of the resistive voltage divider can be calculated as:

Unijunction Transistor R_{B1} Voltage

$$V_{RB1} = \frac{R_{B1}}{R_{B1} + R_{B2}} \times V_{BB}$$

For a unijunction transistor, the resistive ratio of R_{B1} to R_{BB} shown above is called the **intrinsic stand-off ratio** and is given the Greek symbol: η (eta). Typical standard values of η range from 0.5 to 0.8 for most common UJT's.

If a small positive input voltage which is less than the voltage developed across resistance, R_{B1} (ηV_{BB}) is now applied to the Emitter input terminal, the diode p-n junction is reverse biased, thus offering a very high impedance and the device does not conduct. The UJT is switched "OFF" and zero current flows.

However, when the Emitter input voltage is increased and becomes greater than V_{RB1} (or $\eta V_{BB} + 0.7V$, where $0.7V$ equals the p-n junction diode volt drop) the p-n junction becomes forward biased and the unijunction transistor begins to conduct. The result is that Emitter current, ηI_E now flows from the Emitter into the Base region.

The effect of the additional Emitter current flowing into the Base reduces the resistive portion of the channel between the Emitter junction and the B_1 terminal. This reduction in the value of R_{B1} resistance to a very low value means that the Emitter junction becomes even more forward biased resulting in a larger current flow. The effect of this results in a negative resistance at the Emitter terminal.

Likewise, if the input voltage applied between the Emitter and B_1 terminal decreases to a value below breakdown, the resistive value of R_{B1} increases to a high value. Then the **Unijunction Transistor** can be thought of as a voltage breakdown device.

So we can see that the resistance presented by R_{B1} is variable and is dependant on the value of Emitter current, I_E . Then forward biasing the Emitter junction with respect to B_1 causes more current to flow which reduces the resistance between the Emitter, E and

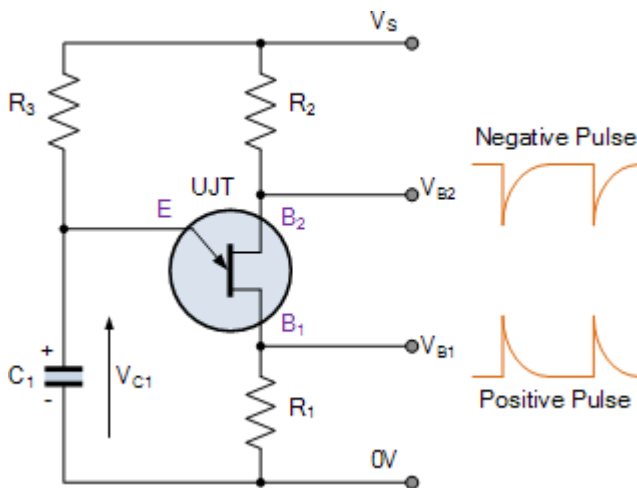
In other words, the flow of current into the UJT's Emitter causes the resistive value of R_{B1} to decrease and the voltage drop across it, V_{RB1} must also decrease, allowing more current to flow producing a negative resistance condition.

Unijunction Transistor Applications

Now that we know how a *unijunction transistor* works, what can they be used for. The most common application of a unijunction transistor is as a triggering device for *SCR's* and *Triacs* but other UJT applications include sawtoothed generators, simple oscillators, phase control, and timing circuits. The simplest of all UJT circuits is the Relaxation Oscillator producing non-sinusoidal waveforms.

In a basic and typical UJT relaxation oscillator circuit, the Emitter terminal of the unijunction transistor is connected to the junction of a series connected resistor and capacitor, RC circuit as shown below.

Unijunction Transistor Relaxation Oscillator

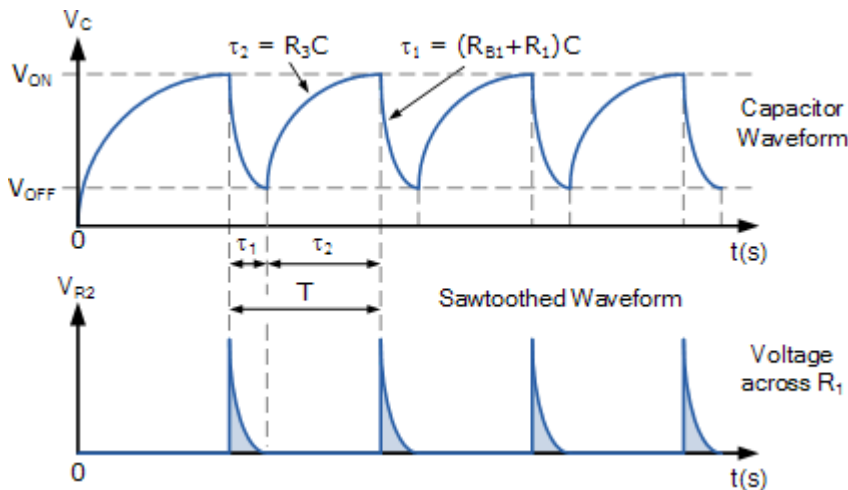


When a voltage (V_S) is firstly applied, the unijunction transistor is “OFF” and the capacitor C_1 is fully discharged but begins to charge up exponentially through resistor R_3 . As the Emitter of the UJT is connected to the capacitor, when the charging voltage V_C across the capacitor becomes greater than the diode volt drop value, the p-n junction behaves as a normal diode and becomes forward biased triggering the UJT into conduction. The unijunction transistor is “ON”. At this point the Emitter to B_1 impedance collapses as the Emitter goes into a low impedance saturated state with the flow of Emitter current through R_1 taking place.

As the ohmic value of resistor R_1 is very low, the capacitor discharges rapidly through the UJT and a fast rising voltage pulse appears across R_1 . Also, because the capacitor discharges more quickly through the UJT than it does charging up through resistor R_3 , the discharging time is a lot less than the charging time as the capacitor discharges through the low resistance UJT.

When the voltage across the capacitor decreases below the holding point of the p-n junction (V_{OFF}), the UJT turns “OFF” and no current flows into the Emitter junction so once again the capacitor charges up through resistor R_3 and this charging and discharging process between V_{ON} and V_{OFF} is constantly repeated while there is a supply voltage, V_S applied.

UJT Oscillator Waveforms



Then we can see that the unijunction oscillator continually switches “ON” and “OFF” without any feedback. The frequency of operation of the oscillator is directly affected by the value of the charging resistance R_3 , in series with the capacitor C_1 and the value of η . The output pulse shape generated from the Base1 (B1) terminal is that of a sawtooth waveform and to regulate the time period, you only have to change the ohmic value of resistance, R_3 since it sets the RC time constant for charging the capacitor.

The time period, T of the sawtoothed waveform will be given as the charging time plus the discharging time of the capacitor. As the discharge time, τ_1 is generally very short in comparison to the larger RC charging time, τ_2 the time period of oscillation is more or less equivalent to $T \cong \tau_2$. The frequency of oscillation is therefore given by $f = 1/T$.

VARACTOR DIODE

Introduction

A Diode is an electronic component that has two terminals and allows current to flow only in one direction. Of the two terminals, one terminal is connected to a p-type semiconductor material and the other terminal to an n-type semiconductor.

Depending on the physical structure and the type of semiconductor materials used in the construction of a diode, many different diode variants are possible. They range from diodes with optical properties like photo-diodes, LEDs, laser diodes etc., to the ones like rectifier, Zener and Tunnel diodes.

In this article we will look at one such variant of the diode – the Varactor Diode.

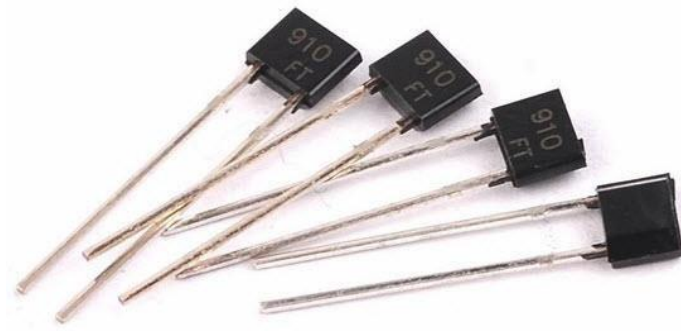
For comprehensive details of different types of diodes, visit this page: [Different Types of Diodes](#)

What is a Varactor Diode?

A Varactor Diode (also known with the names Varicap Diode, Varactor Diode, Tuning Diode) is a p-n junction diode which acts as a variable capacitor under varying reverse bias voltage across its terminals.

In other words, it is a specially designed semiconductor diode whose capacitance at the p-n semiconductor junction changes with the change in voltage applied across its terminals. And

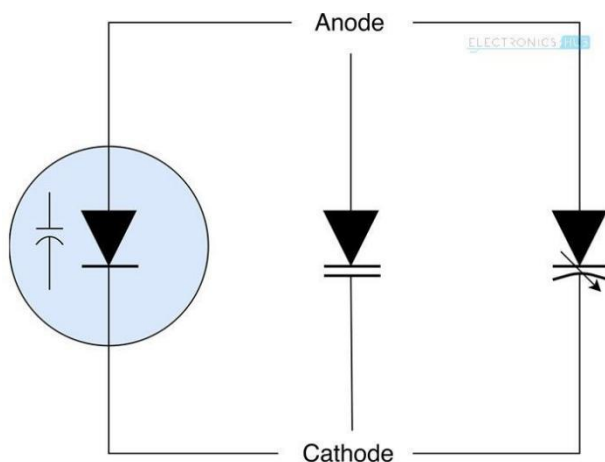
because it is a diode that can behave as a variable capacitor, it is named as a Varactor Diode in short.



It is mainly used to replace variable capacitors that need to be operated mechanically for changing the value of capacitance. One advantage is that the capacitance of a varactor diode can be changed just by changing the voltage across its terminals. We will learn more about its operation in the following sections.

Symbolic Representation

A number of schematic symbols are used to represent a Varactor Diode in circuit diagrams. Of them, the three most popular representations are shown below:

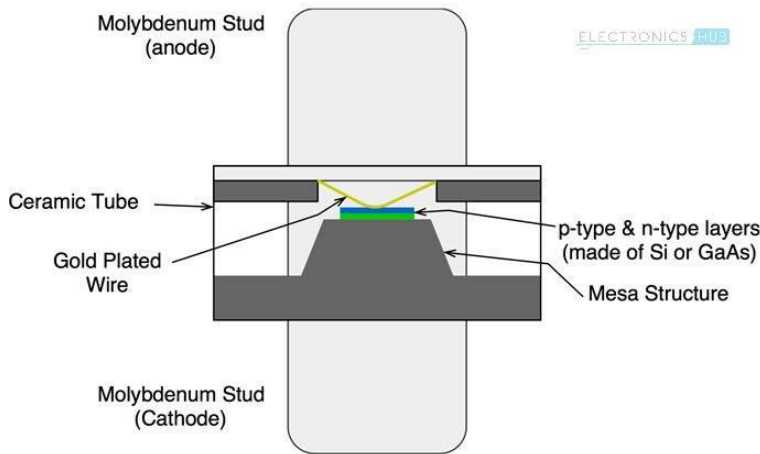


Among these three symbols of a Varactor Diode, the second symbol is more frequently used. It is a combination of symbol of a p-n junction diode and a capacitor. In the symbol, the triangle indicates the presence of a diode component, and the two lines at the tip of triangle indicate parallel plates of a capacitor.

Construction

A Varactor Diode consists of p-type and n-type semiconductor layers sandwiched together, with the n-type layer attached to a mesa (table-shaped) structure. A gold plated molybdenum stud is connected to n-type layer via the mesa structure and it acts as cathode terminal.

The p-type layer is connected to another gold plated molybdenum stud (which acts as anode) via a gold wire. Except for some portion of the molybdenum studs, the entire arrangement is enclosed in a ceramic layer.



The p-type and n-type layers of the varactor diode are made up of silicon or gallium arsenide depending on the type of application for which it is used. For low frequency applications, silicon is used, and for high-frequency applications gallium arsenide is used.

For conventional diodes, the p-type and n-type semiconductor layers are uniformly doped with impurities to improve conductivity. But in the case of varactor diodes, the concentration of impurities near the pn junction is very less and it gradually increases as we move towards the layer's other surface.

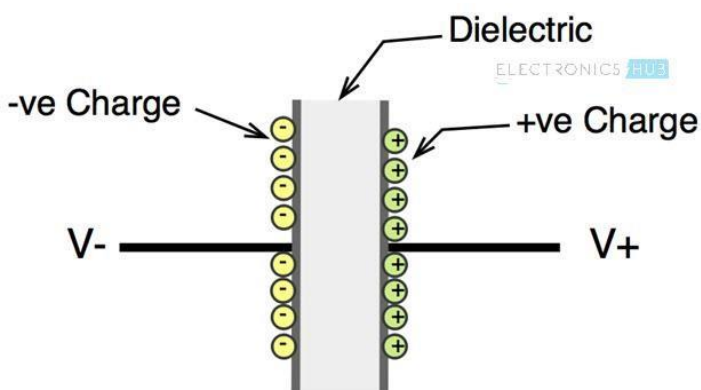
Working

To learn how a varactor diode works, you need to first understand the working principle of a variable capacitor:

A capacitor consists of two conducting surfaces separated by a non-conducting dielectric medium (see figure below). When one of surfaces is connected to a positive voltage and the other to negative voltage, because of attraction between positive and negative carriers, positive charge accumulates on one surface and negative charge on the other.

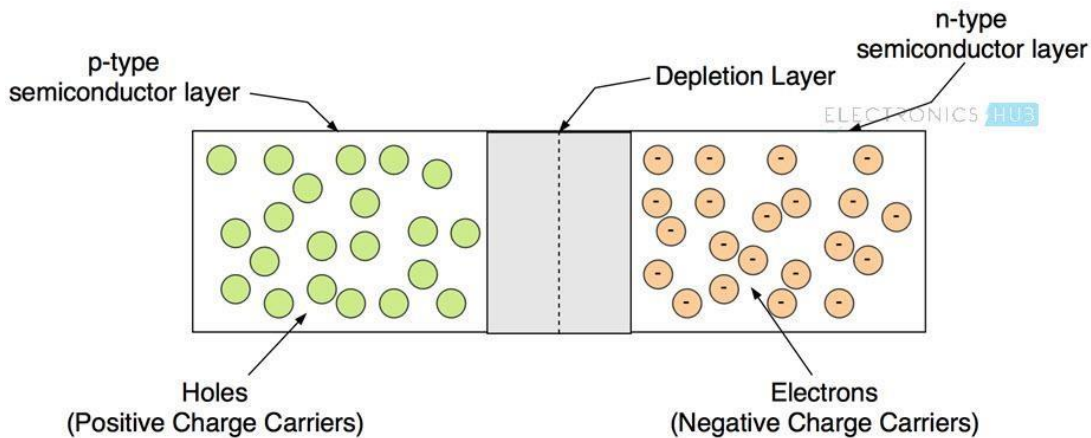
The amount of charge that accumulates is termed as capacitance. If we reduce the gap between the two surfaces, the force of attraction between positive and negative charge carriers increases and so more charge accumulates on the surface i.e. the capacitance increases.

The opposite happens when the surfaces are moved away from each other i.e. the capacitance decreases. A variable capacitor has a mechanical arrangement which allows us to change the gap between surfaces, which effectively changes the capacitance.



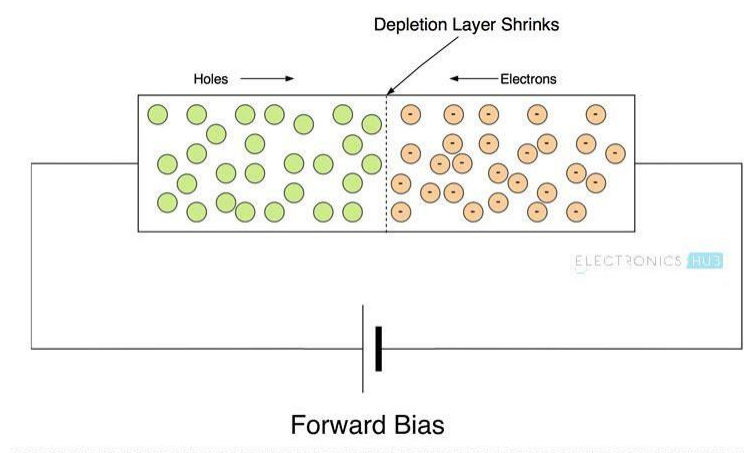
Let us now come back to working of a p-n junction diode. The p-type layer of the diode is filled with positive charge carriers and the n-type layer with negative charge carriers. Near the surface of contact between the two layers, positive and negative charges diffuse and neutralize each other. This region is known as depletion region.

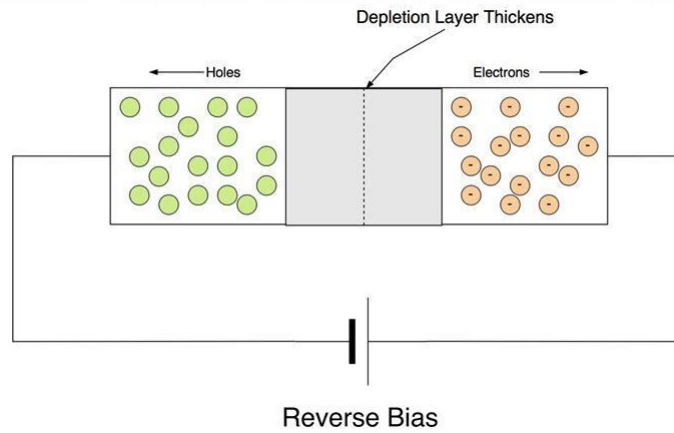
No further diffusion of charge carriers through the depletion layer is possible unless an external voltage is applied. So effectively the depletion layer acts as an insulator.



The width of depletion layer depends on the voltage applied across the p and n-type layers. If a forward bias voltage is applied, i.e. positive voltage is applied to the p-type layer and negative voltage is applied to the n-type layer, the width of depletion decreases, and above a certain voltage it completely disappears.

If a reverse bias is applied, i.e., positive voltage to n-type layer and negative voltage to p-type layer, the width of depletion layer increases. The below diagram illustrates the two scenarios:

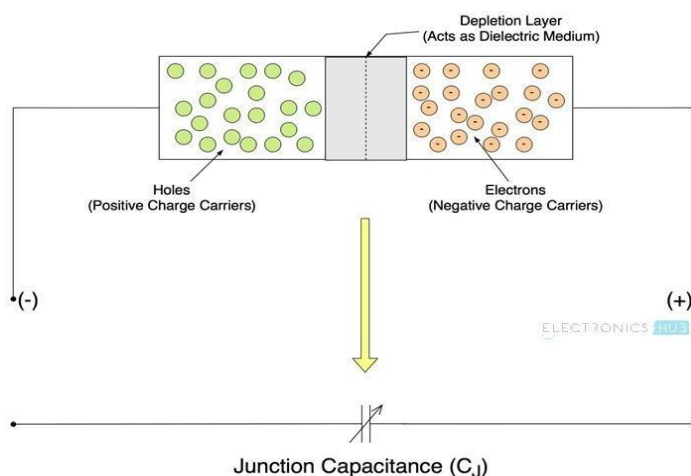




Simply put, the width of the depletion region can be changed to a required value just by adjusting the voltage across the p & n-type semiconductor layers. So, by now you must have observed the similarities between capacitor, and the diode in reverse bias. The depletion layer in the diode is akin to the dielectric medium in a capacitor, which acts as an insulator and prevents the flow of charge carriers from one side to the other.

So, when reverse bias voltage is applied across the diode, respective charge carriers accumulate on either side of the depletion layer. This makes the diode acquire some capacitance and it is termed as junction capacitance.

A Varactor Diode is specially designed to enhance this ability to store charge carriers when reverse bias is applied, thus allowing it to act as a capacitor.

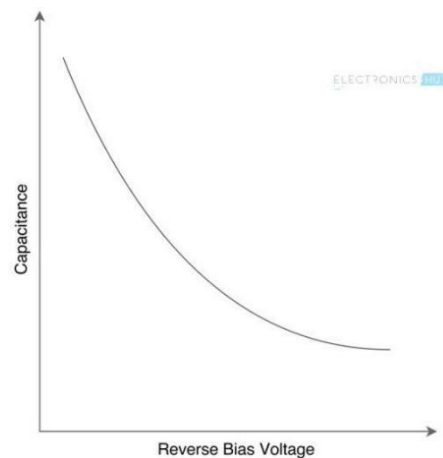


The junction capacitance is inversely proportional to the width of depletion layer i.e. if the width of depletion layer is less, the capacitance will be more, and vice versa. So if we need to increase the capacitance of a varactor diode, the reverse bias voltage should be decreased. It causes the width of depletion layer to decrease, resulting in higher capacitance. Similarly increasing the reverse bias voltage should decrease the capacitance.

This ability to get different values of capacitances just by changing the voltage applied is the biggest advantage of a varactor diode when compared with a normal variable capacitor.

Characteristics

The below graph illustrates the relation between the magnitude of reverse bias voltage p-type & n-type layers of a varactor diode, and the magnitude of junction capacitance.



You can observe that the junction capacitance of a varactor diode is inversely proportional to reverse bias voltage. Also because of difference in the way impurities are added to the p-type and n-type layers, the capacitance of a varactor diode is always higher than that of a conventional diode.

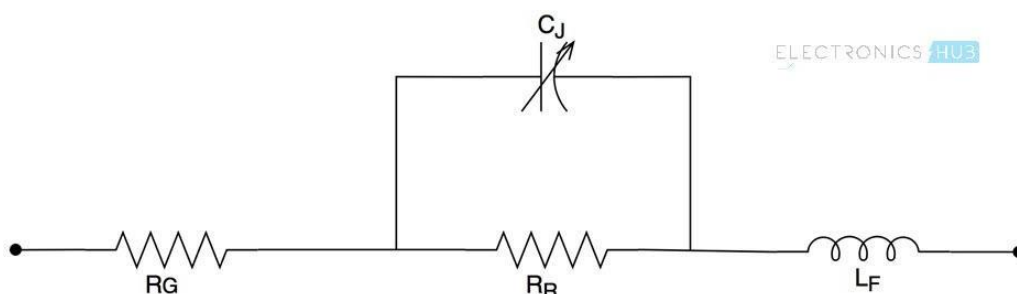
The exact magnitude of junction capacitance of a varactor diode is calculated with respect to the capacitance at zero-bias condition (C_J), forward bias voltage (Voltage required to remove depletion layer completely, V_B) and the actual reverse bias voltage (V_R) applied across the junction. It is given by the formula:

$$C_J = C_0(1 + V_R/V_B)^{-n}$$

'n' is a constant and varies depending on the way p and n-type layers are doped. Its value lies between (0.5 – 0.33). The value of reverse bias voltage (V_R) should be below the breakdown voltage, above which the value of capacitance C_J equals zero because of breakdown of depletion region and the free flow of holes and electrons.

Equivalent Circuit

For designing circuits that use the properties of a varactor diode, we first need to break down its electrical properties and visualize an equivalent circuit with basic components like resistors, capacitors and inductors. The below image shows such approximate equivalent circuit of a varactor diode when operated under reverse bias condition.



R_R is the reverse bias resistance and R_G is the geometric resistance of the varactor diode. C_J is the junction capacitance which is dependent on the reverse bias voltage. L_F is the effective inductance and it depends on the frequency limits under which a varactor diode should be operated.

Applications

Owing to the special property of varying capacitance with varying voltage, varactor diodes are mostly used in frequency modulation or tuning circuits where the value of capacitance determines the output modulation frequency. Some of the other applications include:

- Automatic Frequency Controllers (AFCs)
- Ultra High Frequency Television sets
- High frequency Radios
- Frequency Multipliers
- Band Pass Filters
- Harmonic Generators

UNIT – IV

ANALYSIS AND DESIGN OF SMALL SIGNAL LOW

FREQUENCY BJT AMPLIFIERS

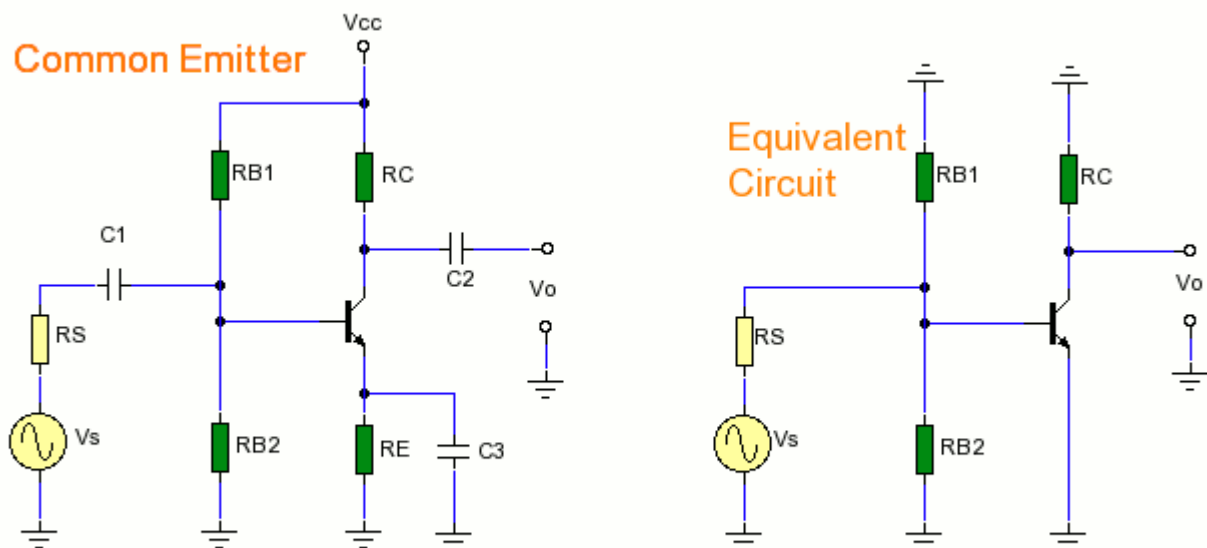
INTRODUCTION

The primary function of a "model" is to predict the behaviour of a device in a particular operating region. At dc the bipolar junction transistor (BJT) and some of its biasing techniques have already been described, see these articles: Biasing, Transistor as switch

The behaviour of the BJT in the sinusoidal ac domain is quite different from its dc domain. At dc the BJT usually works at in either saturation or cutoff regions. In the ac domain the transistor works in the linear region and effects of capacitance between terminals, input impedance, output conductance, etc all have to be accounted for. The small-signal ac response can be described by two common models: the *hybrid model* and *r_e model*. The models are equivalent circuits (or combination of circuit elements) that allow methods of circuit analysis to predict performance.

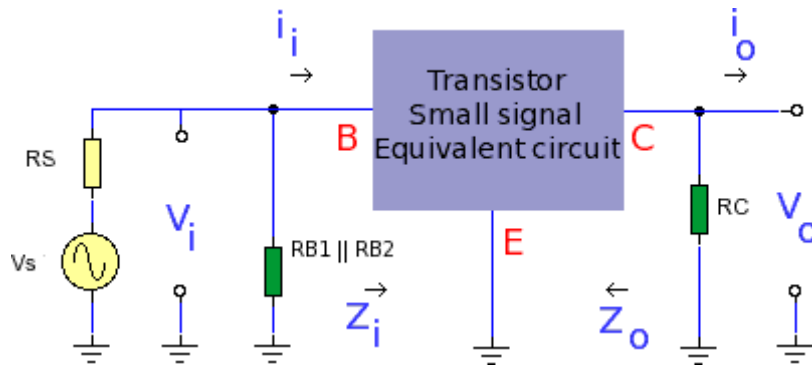
TRANSISTOR HYBRID MODEL

To demonstrate the Hybrid transistor model an ac equivalent circuit must be produced. The left hand diagram below is a single common emitter stage for analysis.



At ac the reactance of coupling capacitors C1 and C2 is so low that they are virtual short circuits, as does the bypass capacitor C3. The power supply (which will have filter capacitors) is also a short circuit as far as ac signals are concerned. The equivalent circuit is shown above on the right hand diagram. The input signal generator is shown as V_s and the generators source impedance as R_s .

As RB1 and RB2 are now in parallel the input impedance will be $RB1 \parallel RB2$. The collector resistor RC also appears from collector to emitter (as emitter is bypassed). See below :

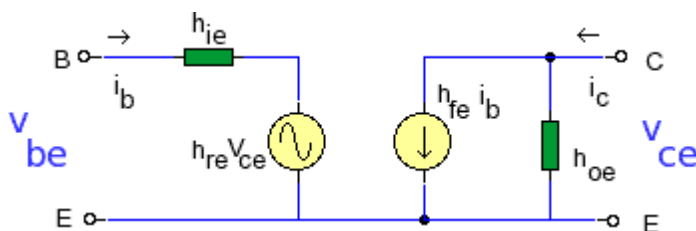


The blue rectangle now represents the small signal ac equivalent circuit and can now start work on the hybrid equivalent circuit.

The hybrid model has four h-parameters. The "h" stands for hybrid because the parameters are a mix of impedance, admittance and dimensionless units. In common emitter the parameters are:

hie	input impedance (Ω)
hre	reverse voltage ratio (dimensionless)
hfe	forward current transfer ratio (dimensionless)
hoe	output admittance (Siemen)

Note that lower case suffixes indicate small signal values and the last suffix indicates the mode so hie is input impedance in common emitter, hfb would be forward current transfer ration in common base mode, etc. The hybrid model for the BJT in common emitter mode is shown below:



The hybrid model is suitable for small signals at mid band and describes the action of the transistor. Two equations can be derived from the diagram, one for input voltage vbe and one for the output ic:

$$v_{be} = h_{ie} i_b + h_{re} v_{ce}$$

$$i_c = h_{fe} i_b + h_{oe} v_{ce}$$

If i_b is held constant ($i_b=0$) then hre and hoe can be solved:

$$h_{re} = v_{be} / v_{ce} | i_b = 0$$

$$h_{oe} = i_c / v_{ce} | i_b = 0$$

Also if v_{ce} is held constant ($v_{ce}=0$) then hie and hfe can be solved:

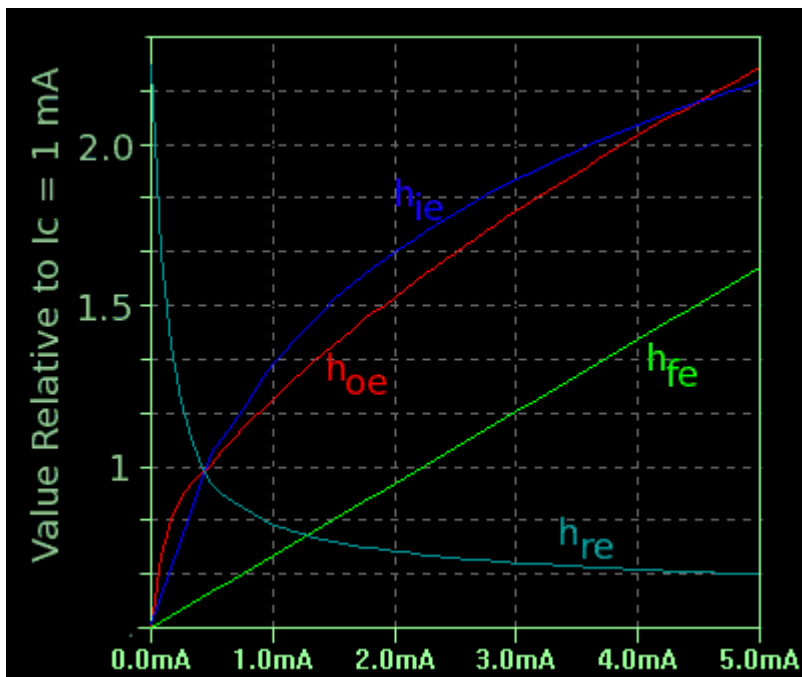
$$h_{ie} = v_{be} / i_b \mid v_{ce} = 0$$

$$h_{fe} = i_c / i_b \mid v_{ce} = 0$$

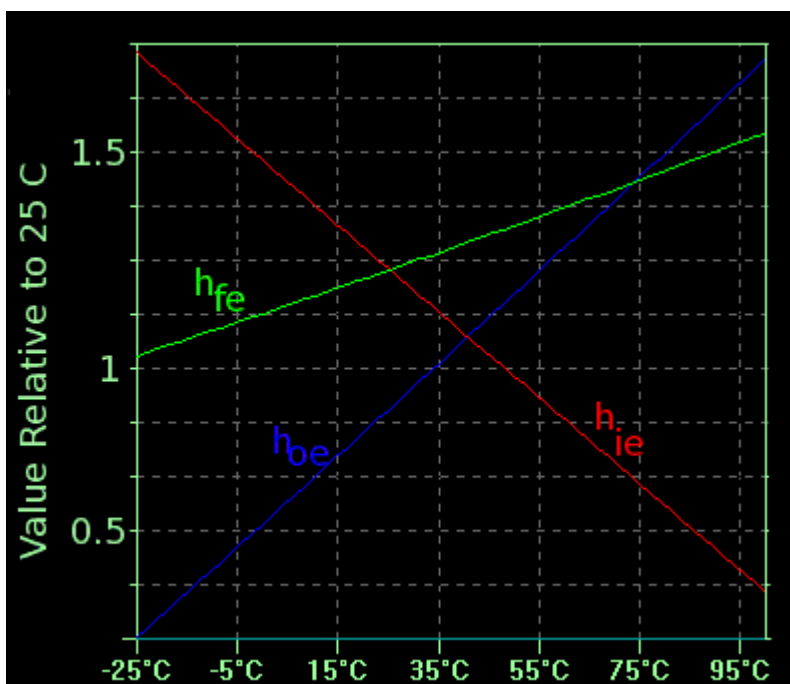
These are the four basic parameters for a BJT in common emitter. Typical values are $h_{re} = 1 \times 10^{-4}$, h_{oe} typical value 20uS, h_{ie} typically 1k to 20k and h_{fe} can be 50 - 750. The H-parameters can often be found on the transistor datasheets. The table below lists the four h-parameters for the BJT in common base and common collector (emitter follower) mode.

Common Base	Common Emitter	Common Collector	Definitions
$h_{ib} = \frac{v_{eb}}{i_e}$	$h_{ie} = \frac{v_{be}}{i_b}$	$h_{ic} = \frac{v_{bc}}{i_b}$	Input Impedance with Output Short Circuit
$h_{rb} = \frac{v_{eb}}{v_{cb}}$	$h_{re} = \frac{v_{be}}{v_{ce}}$	$h_{rc} = \frac{v_{bc}}{v_{ec}}$	Reverse Voltage Ratio Input Open Circuit
$h_{fb} = \frac{i_c}{i_e}$	$h_{fe} = \frac{i_c}{i_b}$	$h_{fc} = \frac{i_e}{i_b}$	Forward Current Gain Output Short Circuit
$h_{ob} = \frac{i_c}{v_{cb}}$	$h_{oe} = \frac{i_c}{v_{ce}}$	$h_{oc} = \frac{i_e}{v_{ec}}$	Output Admittance Input Open Circuit

H-parameters are not constant and vary with temperature, collector current and collector emitter voltage. For this reason when designing a circuit the hybrid parameters should be measured under the same conditions as the actual circuit. Below are graphs of the variation of h-parameters with temperature and collector current.

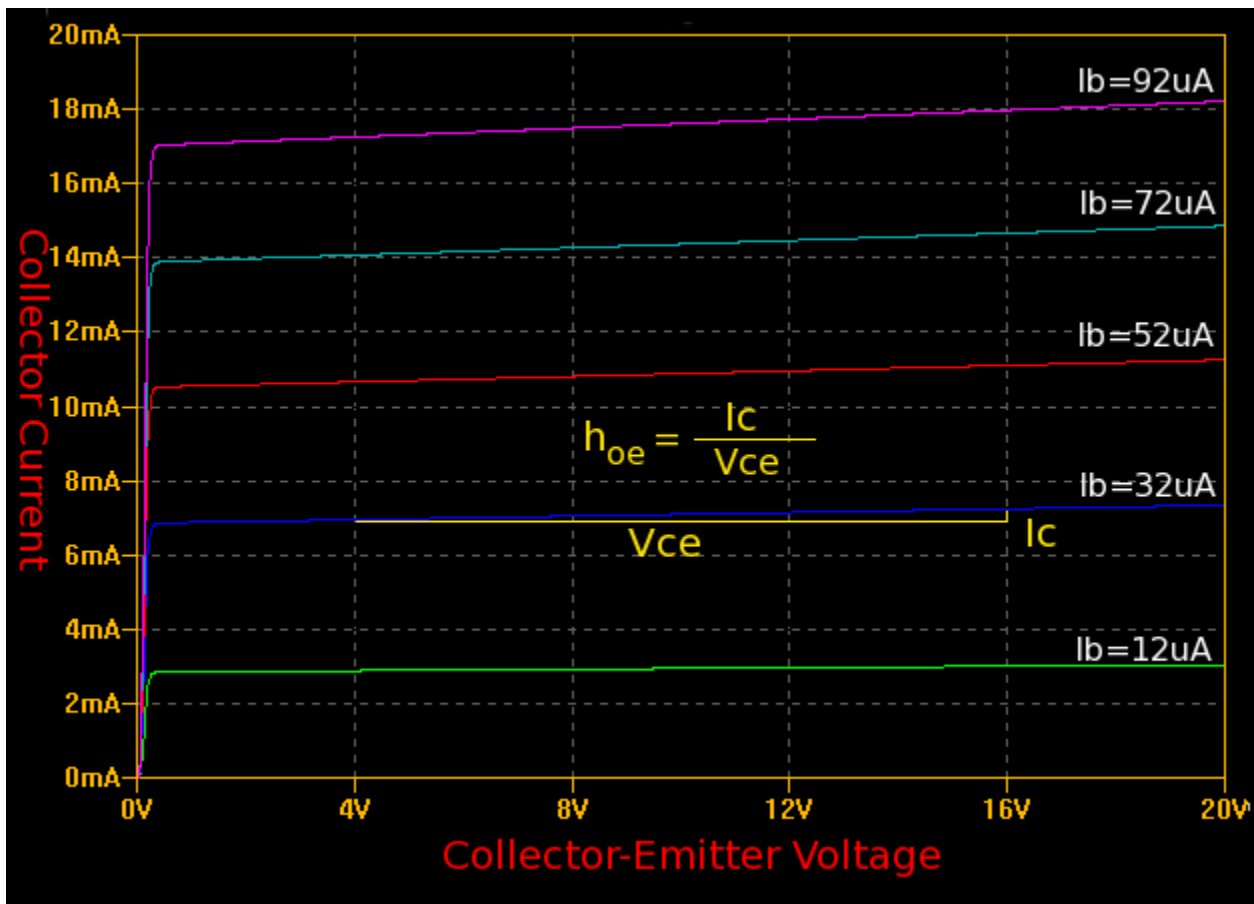


Variation of h-parameters with Temperature



OUTPUT CHARACTERISTIC CURVES

The graph below, shows the output characteristic curves for a 2N3904 transistor in common emitter mode. The output curves are quite useful as they show the change in collector current for a range of collector emitter voltages.



In addition, because the base currents are also known, then two small signal parameters, h_{fe} and h_{oe} can be determined straight from the graph. The almost flat portion of the curves, shows that the transistor behaves as a constant current generator. However, in saturation the steepness of the curves (between 0 and $0.4 V_{ce}$) show a sharp drop in h_{fe} . This is an important fact to consider, if using the transistor as a switch.

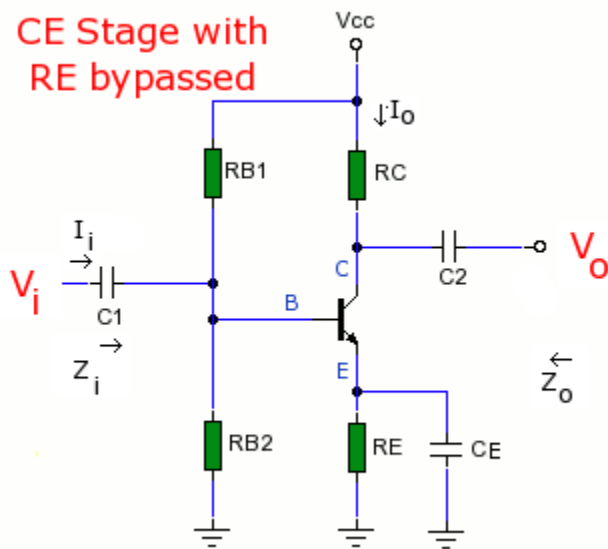
TYPICAL H-PARAMETER VALUES

h-parameters are not constant and vary with both temperature and collector current. Typical values for 1mA collector currents are:

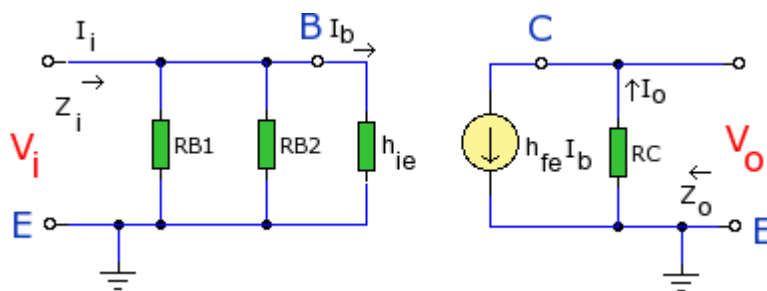
$$h_{ie} = 1000 \Omega \quad h_{re} = 3 \times 10^{-4} \quad h_{oe} = 3 \times 10^{-6} S \quad h_{fe} = 250$$

Examples

The h-parameter model will be applied to a single common emitter (CE) stage with the emitter resistor (R_E) bypassed. The model will be used to build equations for voltage gain, current gain, input and output impedance. The circuit is shown below:



The small signal parameter $h_{re}V_{ce}$ is often too small to be considered so the input resistance is just h_{ie} . Often the output resistance h_{oe} is often large compared with the collector resistor R_C and its effects can be ignored. The h-parameter equivalent model is now simplified and drawn below:



The input impedance is the parallel combination of bias resistors R_{B1} and R_{B2} . As the power supply is considered short circuit at small signal levels then R_{B1} and R_{B2} are in parallel. R_{BB} will represent the parallel combination:

$$R_{BB} = R_{B1} \parallel R_{B2} = \frac{R_{B1} R_{B2}}{R_{B1} + R_{B2}}$$

As R_{BB} is in parallel with h_{ie} then:

$$Z_i = R_{BB} \parallel h_{ie}$$

As $h_{fe}I_b$ is an ideal current generator with infinite output impedance, then output impedance looking into the circuit is:

$$Z_o = R_C$$

Note the $-$ sign in the equation, this indicates phase inversion of the output waveform.

$$V_o = -I_o RC = -h_{fe} I_b RC$$

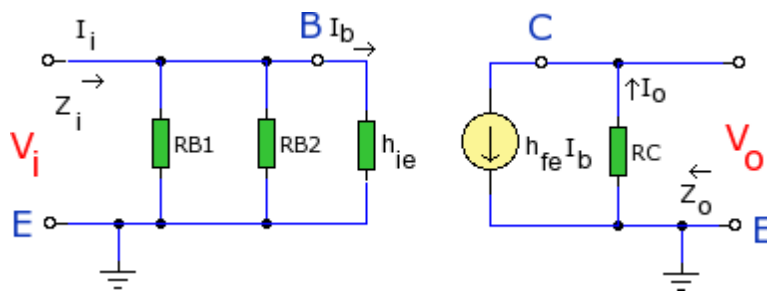
as $I_b = V_i / h_{ie}$ then:

$$= -h_{fe} \frac{V_i}{h_{ie}} RC$$

$$= \frac{-h_{fe}}{h_{ie}} RC V_i$$

$$A_v = \frac{V_o}{V_i} = \frac{-h_{fe} RC}{h_{ie}}$$

The current gain is the ratio I_o / I_i . At the input the current is split between the parallel branch R_{BB} and h_{ie} . So looking at the equivalent h-parameter model again (shown below):



The current divider rule can be used for I_b :

$$I_b = \frac{R_{BB} I_i}{R_{BB} + h_{ie}}$$

$$\frac{I_b}{I_i} = \frac{R_{BB}}{R_{BB} + h_{ie}}$$

At the output side, $I_o = h_{fe} I_b$

re-arranging $I_o / I_b = h_{fe}$

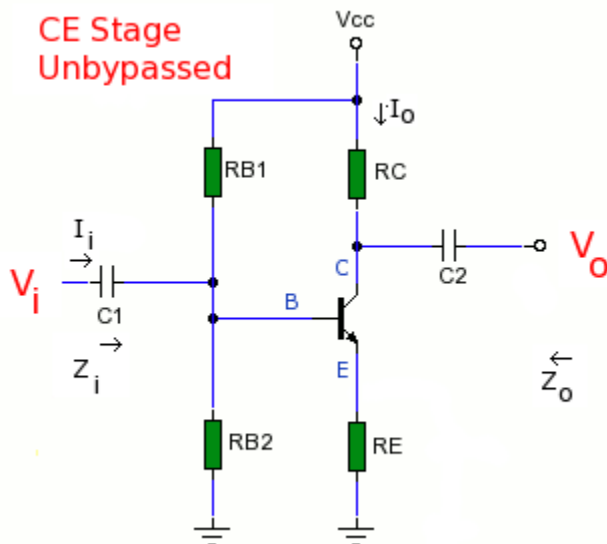
$$A_i = \frac{I_o}{I_i} = \frac{I_o I_b}{I_b I_i} = h_{fe} \frac{R_{BB}}{R_{BB} + h_{ie}}$$

$$A_i = \frac{R_{BB} h_{fe}}{R_{BB} + h_{ie}}$$

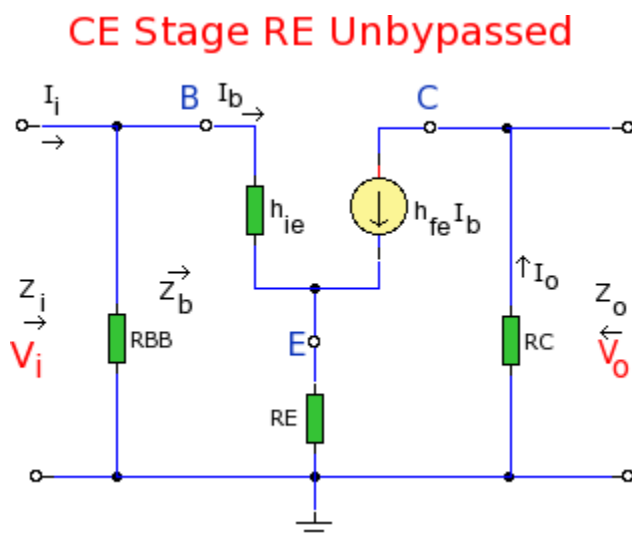
If $R_{BB} \gg h_{ie}$ then,

$$A_i \approx \frac{R_{BB} h_{fe}}{R_{BB}} = h_{fe}$$

The h-parameter model of a common emitter stage with the emitter resistor unbypassed is now shown. The model will be used to build equations for voltage gain, current gain, input and output impedance. The circuit is shown below:



As in the previous example, R_{B1} and R_{B2} are in parallel, the bias resistors are replaced by resistance R_{BB} , but as R_E is now unbypassed this resistor appears in series with the emitter terminal. The hybrid small signal model is shown below, once again effects of small signal parameters $h_{re}V_{ce}$ and h_{oe} have been omitted.



The input impedance Z_i is the bias resistors R_{BB} in parallel with the impedance of the base, Z_b .

$$Z_b = h_{ie} + (1 + h_{fe}) R_E$$

Since h_{fe} is normally much larger than 1, the equation can be reduced to:

$$Z_b = h_{ie} + h_{fe} R_E$$

$$Z_i = R_{BB} \parallel (h_{ie} + h_{fe} R_E)$$

With V_i set to zero, then $I_b = 0$ and $h_{fe}I_b$ can be replaced by an open-circuit. The output impedance is:

$$Z_o = RC$$

Note the $-$ sign in the equation, this indicates phase inversion of the output waveform.

$$I_b = \frac{V_i}{Z_b}$$

$$V_o = -I_o RC = -h_{fe} I_b RC$$

$$= -h_{fe} \frac{V_i}{Z_b} RC$$

$$A_v = \frac{V_o}{V_i} = \frac{-h_{fe} RC}{Z_b}$$

As $Z_b = h_{ie} + h_{fe} RE$ often the product $h_{fe}RE$ is much larger than h_{ie} , so Z_b can be reduced to the approximation:

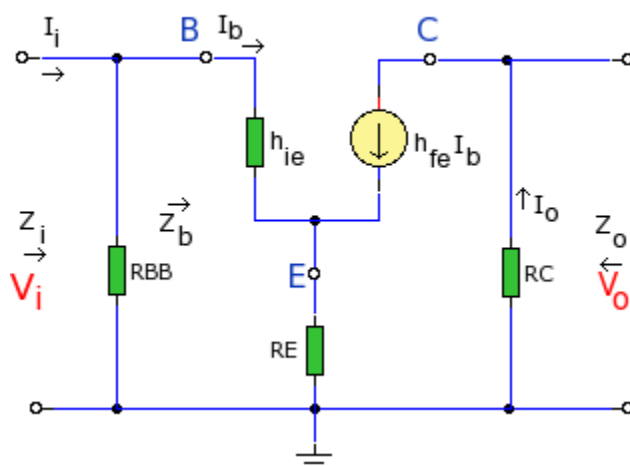
$$Z_b \approx h_{fe}RE$$

$$\therefore A_v = \frac{-h_{fe}RC}{h_{fe}RE}$$

$$A_v = \frac{V_o}{V_i} = - \frac{RC}{RE}$$

The current gain is the ratio I_o / I_i . At the input the current is split between the parallel branch R_{BB} and Z_b . So looking at the equivalent h-parameter model again (shown below):

CE Stage RE Unbypassed



The current divider rule can be used for I_b :

$$I_b = \frac{R_{BB} I_i}{R_{BB} + Z_b}$$

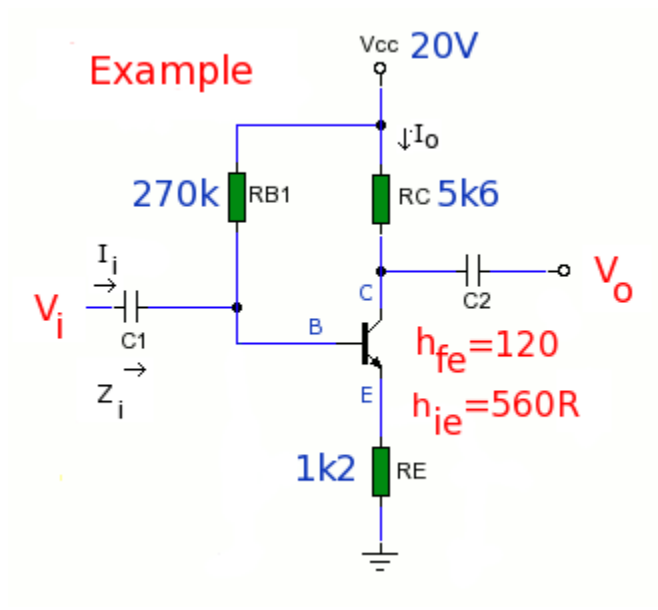
$$\frac{I_b}{I_i} = \frac{R_{BB}}{R_{BB} + Z_b}$$

At the output side, $I_o = h_{fe} I_b$

re-arranging $I_o / I_b = h_{fe}$

$$A_i = \frac{I_o}{I_i} = \frac{I_o I_b}{I_b I_i} = h_{fe} \frac{R_{BB}}{R_{BB} + Z_b}$$

$$A_i = \frac{R_{BB} h_{fe}}{R_{BB} + Z_b}$$



The hybrid parameters must be known to use the hybrid model, either from the datasheet or measured. In the above circuit, Z_i , Z_o , A_v , and A_i will now be calculated. Note that this CE stage uses a single bias resistor R_{B1} which is the value R_{BB} .

$$Z_b = h_{ie} + (1 + h_{fe}) R_E$$

$$= 0.56k + (1 + 120) 1.2k = 145.76k$$

$$Z_i = R_B \parallel Z_b$$

$$Z_i = 270k \parallel 145.76k = 94.66k$$

$$Z_o \approx 5.6k$$

$$A_v = - \frac{h_{fe} R_C}{Z_b}$$

$$= - \frac{120 \times 5.6k}{145.76k}$$

$$A_v = -4.61$$

$$A_i = \frac{R_{BB} h_{fe}}{R_{BB} + Z_b}$$

$$= \frac{270k \times 120}{270k + 145.76k}$$

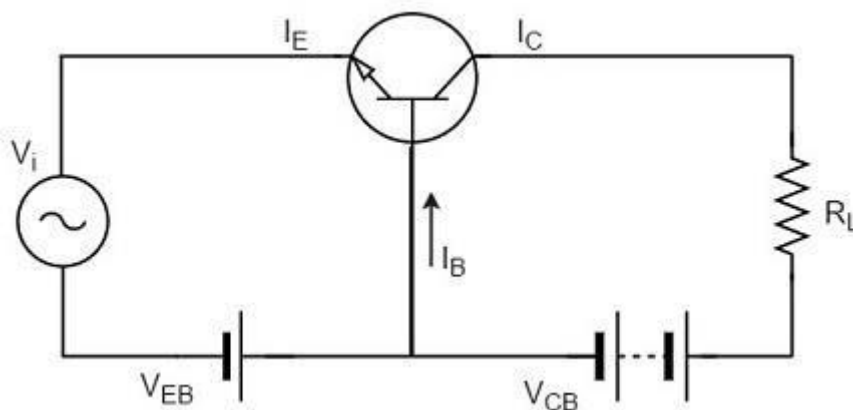
$$A_i = 77.93$$

Summary

The hybrid model is limited to a particular set of operating conditions for accuracy. If the device is operated at a different collector current, temperature or V_{ce} level from the manufacturers datasheet then the h parameters will have to be measured for these new conditions. The hybrid model has parameters for output impedance and reverse voltage ratio which can be important in some circuits.

TRANSISTOR AMPLIFIER

A transistor acts as an amplifier by raising the strength of a weak signal. The DC bias voltage applied to the emitter base junction, makes it remain in forward biased condition. This forward bias is maintained regardless of the polarity of the signal. The below figure shows how a transistor looks like when connected as an amplifier.



The low resistance in input circuit, lets any small change in input signal to result in an appreciable change in the output. The emitter current caused by the input signal contributes the collector current, which when flows through the load resistor R_L , results in a large voltage drop across it. Thus a small input voltage results in a large output voltage, which shows that the transistor works as an amplifier.

Example

Let there be a change of 0.1v in the input voltage being applied, which further produces a change of 1mA in the emitter current. This emitter current will obviously produce a change in collector current, which would also be 1mA.

A load resistance of 5kΩ placed in the collector would produce a voltage of

$$5 \text{ k}\Omega \times 1 \text{ mA} = 5\text{V}$$

Hence it is observed that a change of 0.1v in the input gives a change of 5v in the output, which means the voltage level of the signal is amplified.

Performance of Amplifier

As the common emitter mode of connection is mostly adopted, let us first understand a few important terms with reference to this mode of connection.

Input Resistance

As the input circuit is forward biased, the input resistance will be low. The input resistance is the opposition offered by the base-emitter junction to the signal flow.

By definition, it is the ratio of small change in base-emitter voltage (ΔV_{BE}) to the resulting change in base current (ΔI_B) at constant collector-emitter voltage.

$$\text{Input resistance, } R_i = \Delta V_{BE} / \Delta I_B$$

Where R_i = input resistance, V_{BE} = base-emitter voltage, and I_B = base current.

Output Resistance

The output resistance of a transistor amplifier is very high. The collector current changes very slightly with the change in collector-emitter voltage.

By definition, it is the ratio of change in collector-emitter voltage (ΔV_{CE}) to the resulting change in collector current (ΔI_C) at constant base current.

$$\text{Output resistance} = R_o = \Delta V_{CE} / \Delta I_C$$

Where R_o = Output resistance, V_{CE} = Collector-emitter voltage, and I_C = Collector-emitter voltage.

Effective Collector Load

The load is connected at the collector of a transistor and for a single-stage amplifier, the output voltage is taken from the collector of the transistor and for a multi-stage amplifier, the same is collected from a cascaded stages of transistor circuit.

By definition, it is the total load as seen by the a.c. collector current. In case of single stage amplifiers, the effective collector load is a parallel combination of R_C and R_o .

$$\begin{aligned} \text{Effective Collector Load, } R_{AC} &= R_C // R_o \\ &= \frac{R_C \times R_o}{R_C + R_o} \end{aligned}$$

Hence for a single stage amplifier, effective load is equal to collector load R_C .

In a multi-stage amplifier (i.e. having more than one amplification stage), the input resistance R_i of the next stage also comes into picture.

Effective collector load becomes parallel combination of R_C , R_o and R_i i.e.,

$$\begin{aligned} \text{Effective Collector Load, } R_{AC} &= R_C // R_o // R_i \\ &= \frac{R_C // R_o \times R_i}{R_C // R_o + R_i} \end{aligned}$$

As input resistance R_i is quite small, therefore effective load is reduced.

Current Gain

The gain in terms of current when the changes in input and output currents are observed, is called as **Current gain**. By definition, it is the ratio of change in collector current (ΔI_C) to the change in base current (ΔI_B).

$$\text{Current gain, } \beta = \Delta I_C / \Delta I_B$$

The value of β ranges from 20 to 500. The current gain indicates that input current becomes β times in the collector current.

Voltage Gain

The gain in terms of voltage when the changes in input and output currents are observed, is called as **Voltage gain**. By definition, it is the ratio of change in output voltage (ΔV_{CE}) to the change in input voltage (ΔV_{BE}).

$$\begin{aligned}\text{Voltage gain, } A_V &= \frac{\Delta V_{CE}}{\Delta V_{BE}} \\ &= \frac{\text{Change in output voltage}}{\text{Change in input voltage}} \\ &= \frac{\Delta I_C \times R_{AC}}{\Delta I_B \times R_i} = \frac{\Delta I_C}{\Delta I_B} \times \frac{R_{AC}}{R_i} = \beta \times \frac{R_{AC}}{R_i} \\ &= \beta \times \frac{R_{AC}}{R_i}\end{aligned}$$

For a single stage, $R_{AC} = R_C$.

However, for Multistage,

$$R_{AC} = R_C \times \frac{R_i}{R_C + R_i}$$

Where R_i is the input resistance of the next stage.

Power Gain

The gain in terms of power when the changes in input and output currents are observed, is called as **Power gain**.

By definition, it is the ratio of output signal power to the input signal power.

$$\begin{aligned}\text{Power gain, } A_P &= \frac{(\Delta I_C)^2 \times R_{AC}}{(\Delta I_B)^2 \times R_i} \\ &= \left(\frac{\Delta I_C}{\Delta I_B} \right)^2 \times \frac{R_{AC}}{R_i} \\ &= \left(\frac{\Delta I_C}{\Delta I_B} \right) \times \left(\frac{\Delta I_C}{\Delta I_B} \right) \times \frac{R_{AC}}{R_i} \\ &= \text{Current gain} \times \text{Voltage gain}\end{aligned}$$

Hence these are all the important terms which refer the performance of amplifiers.

ANALYSIS OF CE, CC, CB AMPLIFIERS AND CE AMPLIFIER WITH EMITTER RESISTANCE

Any transistor amplifier, uses a transistor to amplify the signals which is connected in one of the three configurations. For an amplifier it is a better state to have a high input impedance, in order to avoid loading effect in Multi-stage circuits and lower output impedance, in order to deliver maximum output to the load. The voltage gain and power gain should also be high to produce a better output.

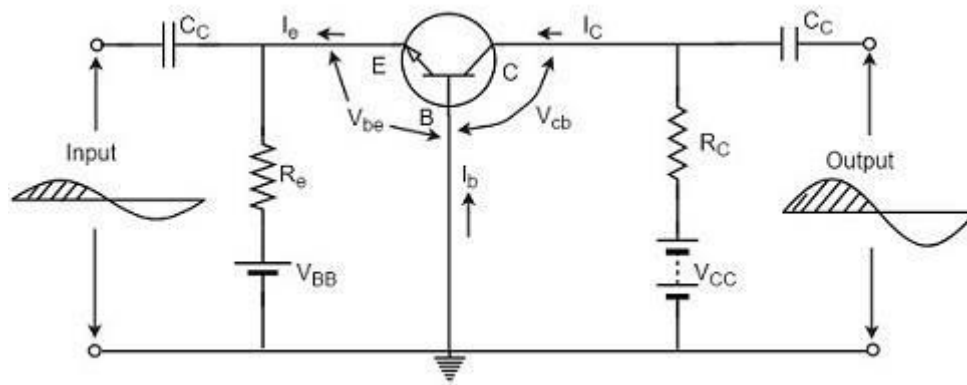
Let us now study different configurations to understand which configuration suits better for a transistor to work as an amplifier.

CB AMPLIFIER

The amplifier circuit that is formed using a CB configured transistor combination is called as CB amplifier.

Construction

The common base amplifier circuit using NPN transistor is as shown below, the input signal being applied at emitter base junction and the output signal being taken from collector base junction.

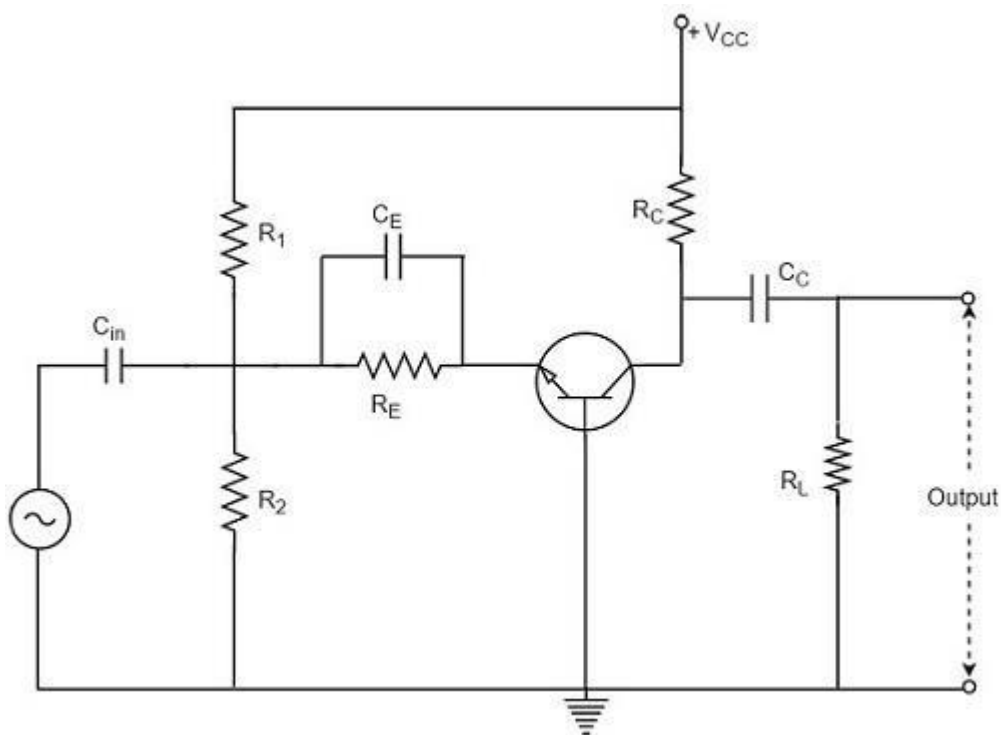


The emitter base junction is forward biased by V_{EE} and collector base junction is reverse biased by V_{CC} . The operating point is adjusted with the help of resistors R_E and R_C . Thus the values of I_C , I_B and I_{CB} are decided by V_{CC} , V_{EE} , R_E and R_C .

Operation

When no input is applied, the quiescent conditions are formed and no output is present. As V_{be} is at negative with respect to ground, the forward bias is decreased, for the positive half of the input signal. As a result of this, the base current I_B also gets decreased.

The below figure shows the CB amplifier with self-bias circuit.



As we know that,

$$I_C \cong I_E \cong \beta I_B$$

Both the collector current and emitter current get decreased.

The voltage drop across R_C is

$$V_C = I_C R_C$$

This V_C also gets decreased.

As $I_C R_C$ decreases, V_{CB} increases. It is because,

$$V_{CB} = V_{CC} - I_C R_C$$

Thus, a positive half cycle output is produced.

In CB configuration, a positive input produces a positive output and hence input and output are in phase. So, there is no phase reversal between input and output in a CB amplifier.

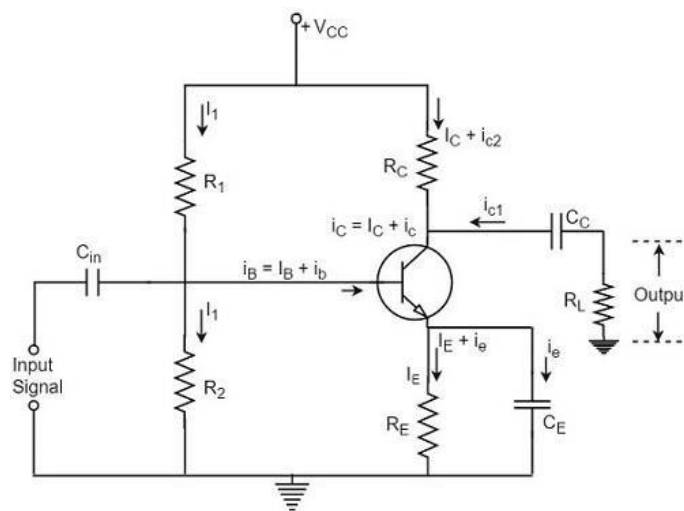
If CB configuration is considered for amplification, it has low input impedance and high output impedance. The voltage gain is also low compared to CE configuration. Hence CB configured amplifiers are used at high frequency applications.

CE AMPLIFIER

The amplifier circuit that is formed using a CE configured transistor combination is called as CE amplifier.

Construction

The common emitter amplifier circuit using NPN transistor is as shown below, the input signal being applied at emitter base junction and the output signal being taken from collector base junction.



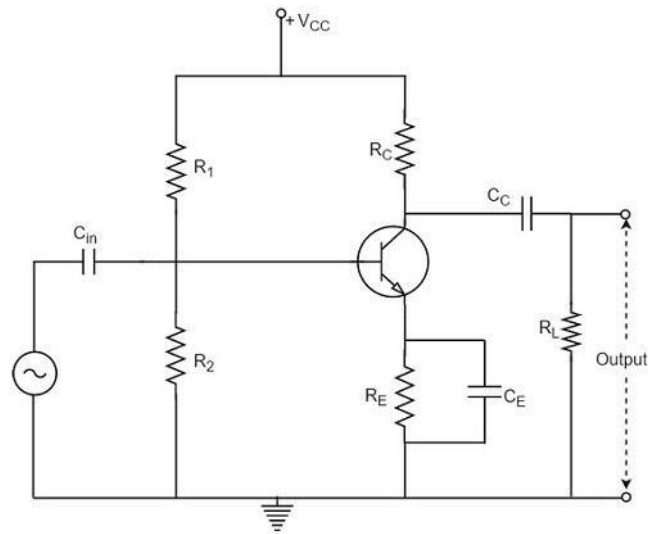
The emitter base junction is forward biased by V_{EE} and collector base junction is reverse biased by V_{CC} . The operating point is adjusted with the help of resistors R_e and R_c . Thus the values of I_c , I_b and I_{cb} are decided by V_{CC} , V_{EE} , R_e and R_c .

Operation

When no input is applied, the quiescent conditions are formed and no output is present. When positive half of the signal is being applied, the voltage between base and emitter V_{be} is increased because it is already positive with respect to ground.

As forward bias increases, the base current too increases accordingly. Since $I_c = \beta I_b$, the collector current increases as well.

The following circuit diagram shows a CE amplifier with self-bias circuit.



The collector current when flows through R_C , the voltage drop increases.

$$V_C = I_C R_C$$

As a consequence of this, the voltage between collector and emitter decreases. Because,

$$V_{CB} = V_{CC} - I_C R_C$$

Thus, the amplified voltage appears across R_C .

Therefore, in a CE amplifier, as the positive going signal appears as a negative going signal, it is understood that there is a phase shift of 180° between input and output.

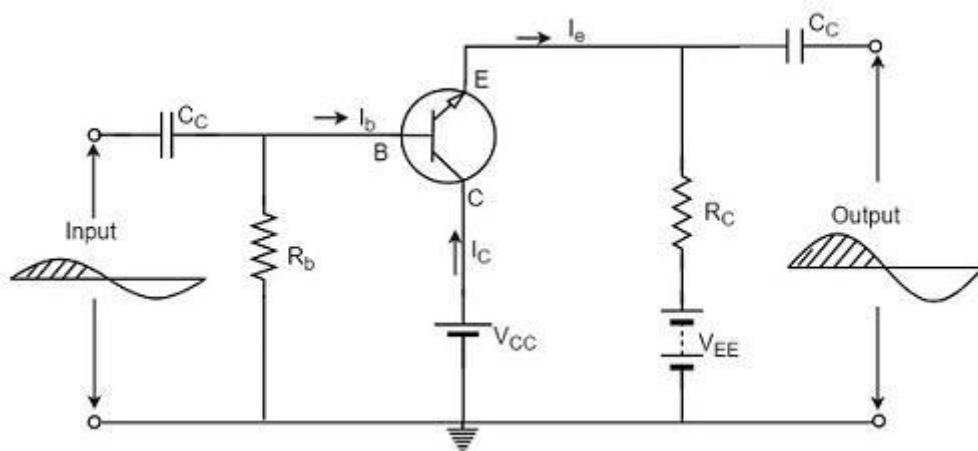
CE amplifier has a high input impedance and lower output impedance than CB amplifier. The voltage gain and power gain are also high in CE amplifier and hence this is mostly used in Audio amplifiers.

CC AMPLIFIER

The amplifier circuit that is formed using a CC configured transistor combination is called as CC amplifier.

Construction

The common collector amplifier circuit using NPN transistor is as shown below, the input signal being applied at base collector junction and the output signal being taken from emitter collector junction.

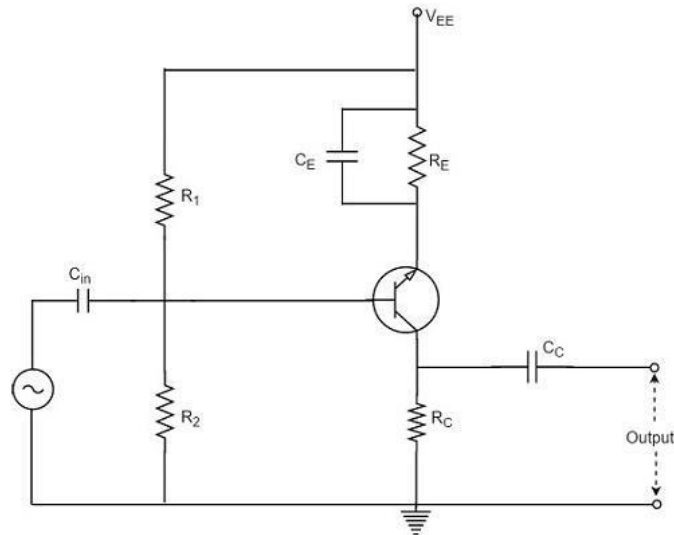


The emitter base junction is forward biased by V_{EE} and collector base junction is reverse biased by V_{CC} . The Q-values of I_b and I_e are adjusted by R_b and R_e .

Operation

When no input is applied, the quiescent conditions are formed and no output is present. When positive half of the signal is being applied, the forward bias is increased because V_{be} is positive with respect to collector or ground. With this, the base current I_B and the collector current I_C are increased.

The following circuit diagram shows a CC amplifier with self-bias circuit.



Consequently, the voltage drop across R_E i.e. the output voltage is increased. As a result, positive half cycle is obtained. As the input and output are in phase, there is no phase reversal.

If CC configuration is considered for amplification, though CC amplifier has better input impedance and lower output impedance than CE amplifier, the voltage gain of CC is very less which limits its applications to impedance matching only.

Comparison between CB CE CC Amplifiers

Let us compare the characteristic details of CB, CE, and CC amplifiers.

Characteristic	CE	CB	CC
Input resistance	Low (1K to 2K)	Very low (30-150 Ω)	High (20-500 K Ω)
Output resistance	Large (≈ 50 K)	High (≈ 500 K)	Low (50-1000 K Ω)
Current gain	B high	$\alpha < 1$	High ($1 + \beta$)

Voltage gain	High ($\approx$ 1500)	High ($\approx$ 1500)	Less than one
Power gain	High ($\approx$ 10,000)	High ($\approx$ 7500)	Low (250-500)
Phase between input and output	reversed	same	same

Due to the compatibility and characteristic features, the common-emitter configuration is mostly used in amplifier circuits

UNIT V

FET AMPLIFIERS

FET AC Equivalent Circuit

Now that the important parameters of an ac equivalent circuit have been introduced and discussed, a model for the FET transistor in the ac domain can be constructed. The control of I_d by V_{gs} is included as a current source $g_m V_{gs}$ connected from drain to source as shown in Fig. 9.8. The current source has its arrow pointing from drain to source to establish a 180° phase shift between output and input voltages as will occur in actual operation.

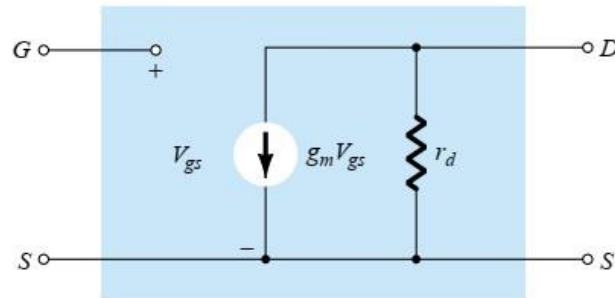


Figure 9.8 FET ac equivalent circuit.

The input impedance is represented by the open circuit at the input terminals and the output impedance by the resistor r_d from drain to source. Note that the gate to source voltage is now represented by V_{gs} (lower-case subscripts) to distinguish it from dc levels. In addition, take note of the fact that the source is common to both input and output circuits while the gate and drain terminals are only in “touch” through the controlled current source $g_m V_{gs}$.

In situations where r_d is ignored (assumed sufficiently large to other elements of the network to be approximated by an open circuit), the equivalent circuit is simply a current source whose magnitude is controlled by the signal V_{gs} and parameter g_m —clearly a voltage-controlled device.

9.3 JFET FIXED-BIAS CONFIGURATION

Now that the FET equivalent circuit has been defined, a number of fundamental FET small-signal configurations will be investigated. The approach will parallel the ac analysis of BJT amplifiers with a determination of the important parameters of Z_i , Z_o , and A_v for each configuration.

The *fixed-bias* configuration of Fig. 9.10 includes the coupling capacitors C_1 and C_2 that isolate the dc biasing arrangement from the applied signal and load; they act as short-circuit equivalents for the ac analysis.

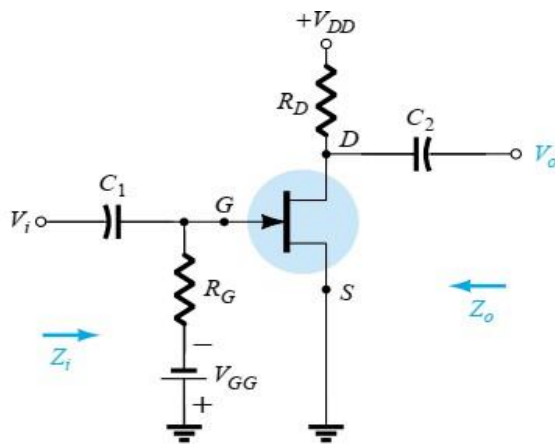


Figure 9.10 JFET fixed-bias configuration.

Once the level of g_m and r_d are determined from the dc biasing arrangement, specification sheet, or characteristics, the ac equivalent model can be substituted between the appropriate terminals as shown in Fig. 9.11. Note that both capacitors have the short-circuit equivalent because the reactance $X_C = 1/(2\pi fC)$ is sufficiently small compared to other impedance levels of the network, and the dc batteries V_{GG} and V_{DD} are set to zero volts by a short-circuit equivalent.

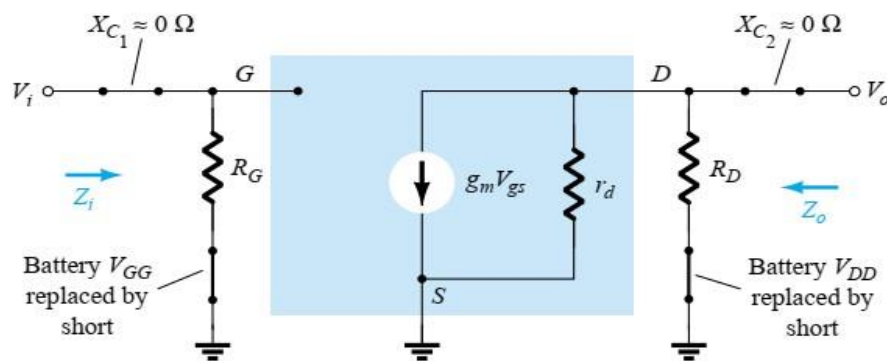


Figure 9.11 Substituting the JFET ac equivalent circuit unit into the network of Fig. 9.10.

The network of Fig. 9.11 is then carefully redrawn as shown in Fig. 9.12. Note the defined polarity of V_{gs} , which defines the direction of $g_m V_{gs}$. If V_{gs} is negative, the direction of the current source reverses. The applied signal is represented by V_i and the output signal across R_D by V_o .

Z_i : Figure 9.12 clearly reveals that

$$Z_i = R_G$$

(9.13)

because of the open-circuit equivalence at the input terminals of the JFET.

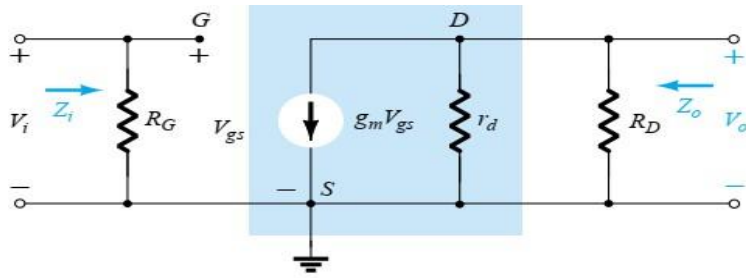


Figure 9.12 Redrawn network of Fig. 9.11.

Z_o : Setting $V_i = 0$ V as required by the definition of Z_o will establish V_{gs} as 0 V also. The result is $g_m V_{gs} = 0$ mA, and the current source can be replaced by an open-circuit equivalent as shown in Fig. 9.13. The output impedance is

$$Z_o = R_D \parallel r_d \quad (9.14)$$

If the resistance r_d is sufficiently large (at least 10:1) compared to R_D , the approximation $r_d \parallel R_D \cong R_D$ can often be applied and

$$Z_o \cong R_D \quad r_d \geq 10R_D \quad (9.15)$$

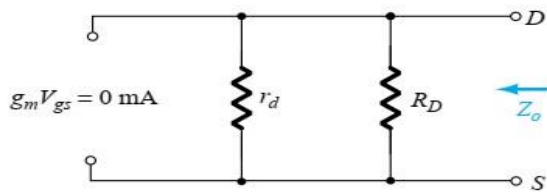


Figure 9.13 Determining Z_o .

A_v : Solving for V_o in Fig. 9.12, we find

$$V_o = -g_m V_{gs} (r_d \parallel R_D)$$

but

$$V_{gs} = V_i$$

and

$$V_o = -g_m V_i (r_d \parallel R_D)$$

so that

$$A_v = \frac{V_o}{V_i} = -g_m (r_d \parallel R_D) \quad (9.16)$$

If $r_d \geq 10R_D$:

$$A_v = \frac{V_o}{V_i} = -g_m R_D \quad r_d \geq 10R_D \quad (9.17)$$

Phase Relationship: The negative sign in the resulting equation for A_v clearly reveals a phase shift of 180° between input and output voltages.

9.4 JFET SELF-BIAS CONFIGURATION

Bypassed R_S

The fixed-bias configuration has the distinct disadvantage of requiring two dc voltage sources. The *self-bias* configuration of Fig. 9.15 requires only one dc supply to establish the desired operating point.

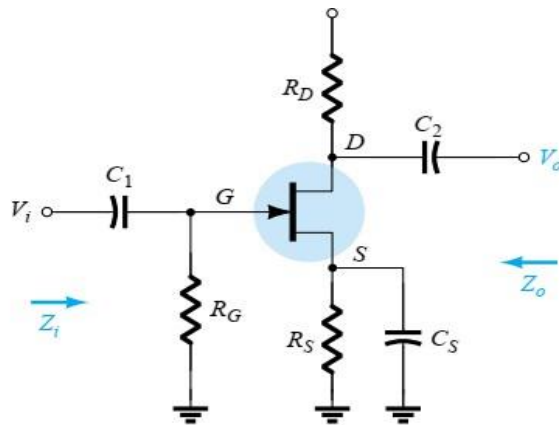


Figure 9.15 Self-bias JFET configuration

The capacitor C_S across the source resistance assumes its short-circuit equivalent for dc, allowing R_S to define the operating point. Under ac conditions, the capacitor assumes the short-circuit state and “short circuits” the effects of R_S . If left in the circuit, the gain will be reduced as will be shown in the paragraphs to follow.

The JFET equivalent circuit is established in Fig. 9.16 and carefully redrawn in Fig. 9.17.

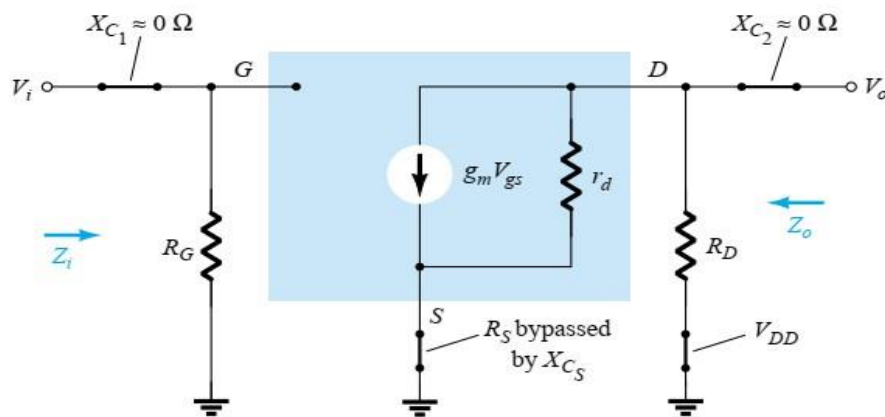


Figure 9.16 Network of Fig. 9.15 following the substitution of the JFET ac equivalent circuit.

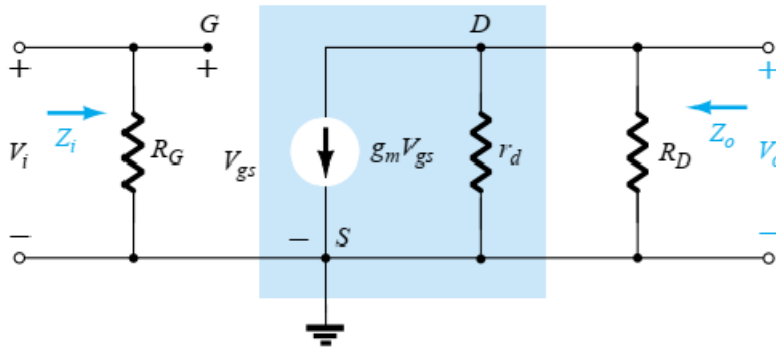


Figure 9.17 Redrawn network of Fig. 9.16.

Since the resulting configuration is the same as appearing in Fig. 9.12, the resulting equations Z_i , Z_o , and A_v will be the same.

Z_i :

$$Z_i = R_G \quad (9.18)$$

Z_o :

$$Z_o = r_d \parallel R_D \quad (9.19)$$

If $r_d \geq 10R_D$,

$$Z_o \cong R_D \quad r_d \geq 10R_D \quad (9.20)$$

A_v :

$$A_v = -g_m(r_d \parallel R_D) \quad (9.21)$$

If $r_d \geq 10R_D$,

$$A_v = -g_m R_D \quad r_d \geq 10R_D \quad (9.22)$$

Phase relationship: The negative sign in the solutions for A_v again indicates a phase shift of 180° between V_i and V_o .

Unbypassed R_S

If C_S is removed from Fig 9.15, the resistor R_S will be part of the ac equivalent circuit as shown in Fig. 9.18. In this case, there is no obvious way to reduce the network to lower its level of complexity. In determining the levels of Z_i , Z_o , and A_v , one must simply be very careful with notation and defined polarities and direction. Initially, the resistance r_d will be left out of the analysis to form a basis for comparison.

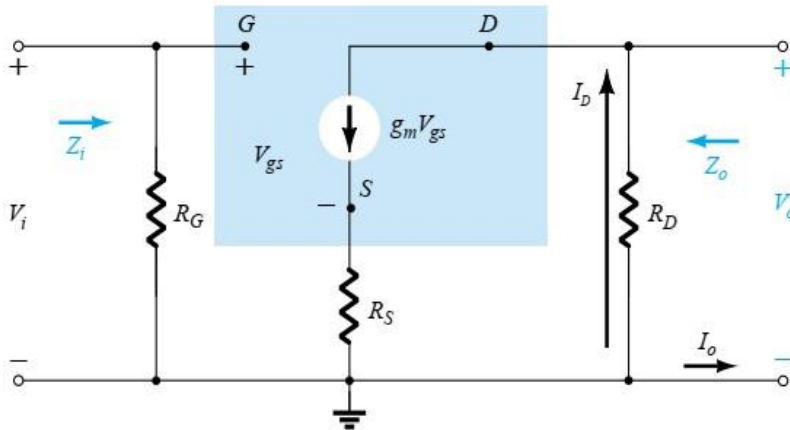


Figure 9.18 Self-bias JFET configuration including the effects of R_S with $r_d = \infty\Omega$.

Z_i : Due to the open-circuit condition between the gate and output network, the input remains the following:

$$Z_i = R_G \quad (9.23)$$

Z_o : The output impedance is defined by

$$Z_o = \left. \frac{V_o}{I_o} \right|_{V_i = 0}$$

Setting $V_i = 0$ V in Fig. 9.18 will result in the gate terminal being at ground potential (0 V). The voltage across R_G is then 0 V, and R_G has been effectively “shorted out” of the picture.

Applying Kirchhoff's current law will result in:

$$I_o + I_D = g_m V_{gs}$$

with

$$V_{gs} = -(I_o + I_D)R_S$$

so that

$$I_o + I_D = -g_m (I_o + I_D)R_S = -g_m I_o R_S - g_m I_D R_S$$

or

$$I_o[1 + g_m R_S] = -I_D[1 + g_m R_S]$$

and

$$I_o = -I_D \quad (\text{the controlled current source } g_m V_{gs} = 0 \text{ A for the applied conditions})$$

Since

$$V_o = -I_D R_D$$

then

$$V_o = -(-I_o)R_D = I_o R_D$$

and

$$Z_o = \frac{V_o}{I_o} = R_D$$

(9.24)

$r_d = \infty \Omega$

If r_d is included in the network, the equivalent will appear as shown in Fig. 9.19.

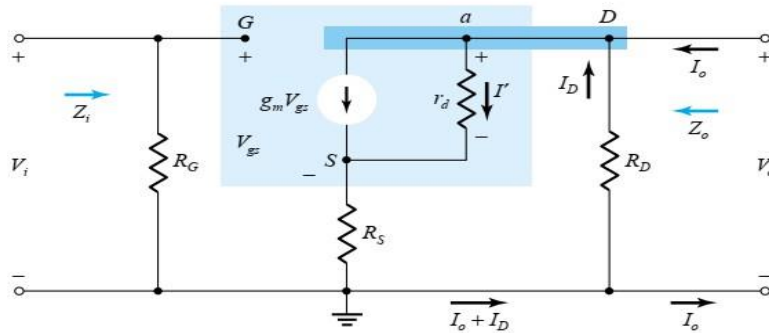


Figure 9.19 Including the effects of r_d in the self-bias JFET configuration.

Since

$$Z_o = \frac{V_o}{I_o} \Big|_{V_i = 0 \text{ V}} = -\frac{I_D R_D}{I_o}$$

we should try to find an expression for I_o in terms of I_D .

Applying Kirchhoff's current law:

$$I_o = g_m V_{gs} + I_{r_d} - I_D$$

but

$$V_{r_d} = V_o + V_{gs}$$

and

$$I_o = g_m V_{gs} + \frac{V_o + V_{gs}}{r_d} - I_D$$

or

$$I_o = \left(g_m + \frac{1}{r_d}\right) V_{gs} - \frac{I_D R_D}{r_d} - I_D \text{ using } V_o = -I_D R_D$$

Now,

$$V_{gs} = -(I_D + I_o) R_S$$

so that

$$I_o = -\left(g_m + \frac{1}{r_d}\right)(I_D + I_o) R_S - \frac{I_D R_D}{r_d} - I_D$$

with the result that

$$I_o \left[1 + g_m R_S + \frac{R_S}{r_d}\right] = -I_D \left[1 + g_m R_S + \frac{R_S}{r_d} + \frac{R_D}{r_d}\right]$$

$$-I_D \left[1 + g_m R_S + \frac{R_S}{r_d} + \frac{R_D}{r_d}\right]$$

or

$$I_o = \frac{-I_D \left[1 + g_m R_S + \frac{R_S}{r_d} + \frac{R_D}{r_d}\right]}{1 + g_m R_S + \frac{R_S}{r_d}}$$

and

$$Z_o = \frac{V_o}{I_o} = \frac{-I_D R_D}{-I_D \left(1 + g_m R_S + \frac{R_S}{r_d} + \frac{R_D}{r_d} \right)}$$

$$1 + g_m R_S + \frac{R_S}{r_d}$$

and finally,

$$Z_o = \frac{\left[1 + g_m R_S + \frac{R_S}{r_d} \right]}{\left[1 + g_m R_S + \frac{R_S}{r_d} + \frac{R_D}{r_d} \right]} R_D \quad (9.25a)$$

For $r_d \geq 10 R_D$, $\left(1 + g_m R_S + \frac{R_S}{r_d} \right) \gg \frac{R_D}{r_d}$ and $1 + g_m R_S + \frac{R_S}{r_d} + \frac{R_D}{r_d} \cong 1 + g_m R_S + \frac{R_S}{r_d}$ and

$$Z_o = R_D \quad r_d \geq 10 R_D \quad (9.25b)$$

A_v : For the network of Fig. 9.19, an application of Kirchhoff's voltage law on the input circuit will result in

$$V_i - V_{gs} - V_{R_S} = 0$$

$$V_{gs} = V_i - I_D R_S$$

The voltage across r_d using Kirchhoff's voltage law is

$$V_o - V_{R_S}$$

and

$$I' = \frac{V_o - V_{R_S}}{r_d}$$

so that an application of Kirchhoff's current law will result in

$$I_D = g_m V_{gs} + \frac{V_o - V_{R_S}}{r_d}$$

Substituting for V_{gs} from above and substituting for V_o and V_{R_S} we have

$$I_D = g_m [V_i - I_D R_S] + \frac{(-I_D R_D) - (I_D R_S)}{r_d}$$

so that

$$I_D \left[1 + g_m R_S + \frac{R_D + R_S}{r_d} \right] = g_m V_i$$

or

$$I_D = \frac{g_m V_i}{1 + g_m R_S + \frac{R_D + R_S}{r_d}}$$

The output voltage is then

$$V_o = -I_D R_D = - \frac{g_m R_D V_i}{1 + g_m R_S + \frac{R_D + R_S}{r_d}}$$

and

$$A_v = \frac{V_o}{V_i} = - \frac{g_m R_D}{1 + g_m R_S + \frac{R_D + R_S}{r_d}} \quad (9.26)$$

Again, if $r_d \geq 10(R_D + R_S)$,

$$A_v = \frac{V_o}{V_i} = - \frac{g_m R_D}{1 + g_m R_S} \quad r_d \geq 10(R_D + R_S) \quad (9.27)$$

Phase Relationship: The negative sign in Eq. (9.26) again reveals that a 180° phase shift will exist between V_i and V_o .

9.5 JFET VOLTAGE-DIVIDER CONFIGURATION

The popular voltage-divider configuration for BJTs can also be applied to JFETs as demonstrated in Fig. 9.21.

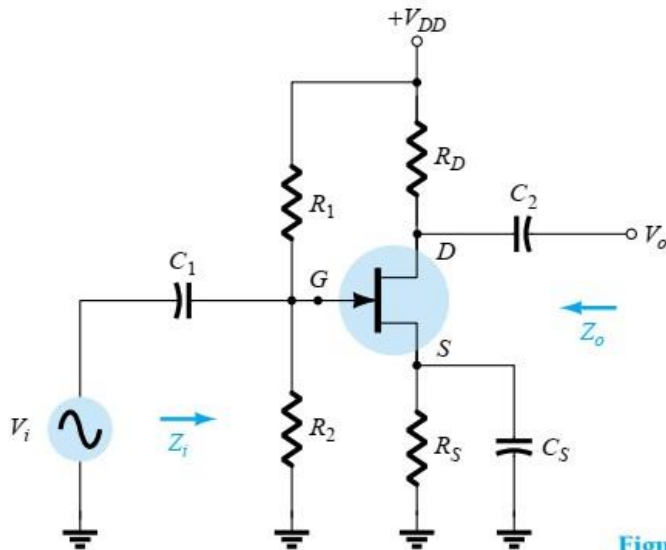


Figure 9.21 JFET voltage-divider configuration.

Substituting the ac equivalent model for the JFET will result in the configuration of Fig. 9.22. Replacing the dc supply V_{DD} by a short-circuit equivalent has grounded one end of R_1 and R_D . Since each network has a common ground, R_1 can be brought down in parallel with R_2 as shown in Fig. 9.23. R_D can also be brought down to ground but in the output circuit across r_d . The resulting ac equivalent network now has the basic format of some of the networks already analyzed.

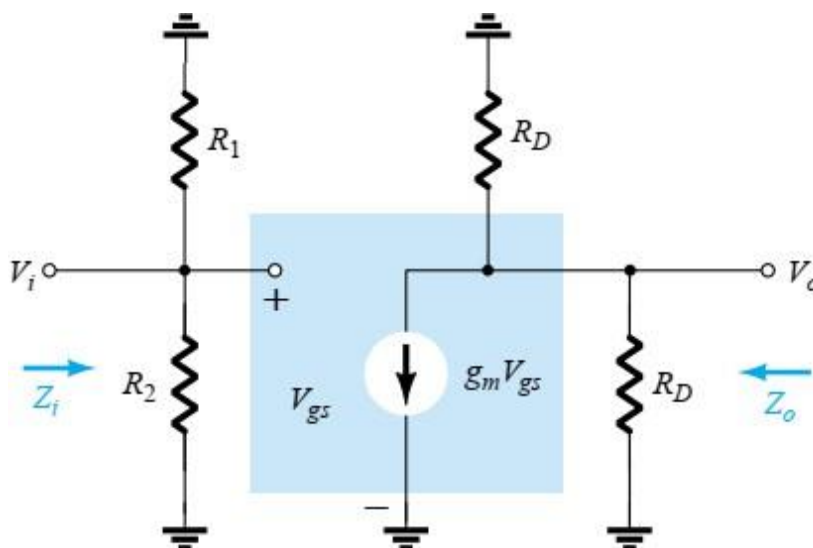


Figure 9.22 Network of Fig. 9.21 under ac conditions.

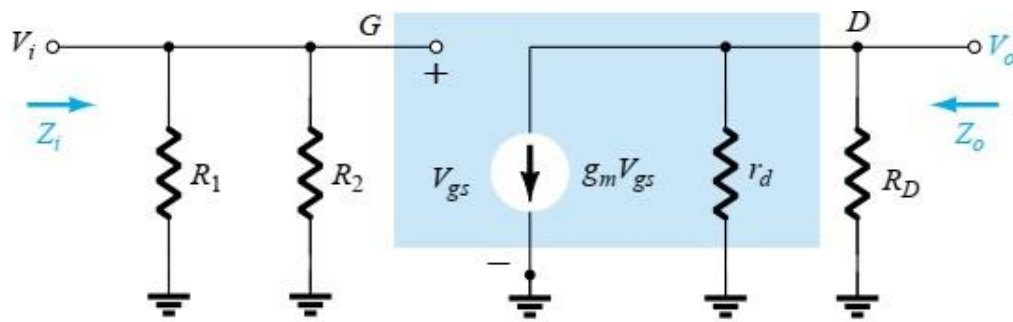


Figure 9.23 Redrawn network of Fig. 9.22.

Z_i : R_1 and R_2 are in parallel with the open-circuit equivalence of the JFET resulting in

$$Z_i = R_1 \parallel R_2 \quad (9.28)$$

Z_o : Setting $V_i = 0$ V will set V_{gs} and $g_m V_{gs}$ to zero and

$$Z_o = r_d \parallel R_D \quad (9.29)$$

For $r_d \geq 10R_D$,

$$Z_o \cong R_D \quad r_d \geq 10R_D \quad (9.30)$$

A_v :

$$V_{gs} = V_i$$

and

$$V_o = -g_m V_{gs} (r_d \parallel R_D)$$

so that

$$A_v = \frac{V_o}{V_i} = \frac{-g_m V_{gs} (r_d \parallel R_D)}{V_{gs}}$$

and

$$A_v = \frac{V_o}{V_i} = -g_m (r_d \parallel R_D) \quad (9.31)$$

If $r_d \geq 10R_D$,

$$A_v = \frac{V_o}{V_i} \cong -g_m R_D \quad r_d \geq 10R_D \quad (9.32)$$

Note that the equations for Z_o and A_v are the same as obtained for the fixed-bias and self-bias (with bypassed R_S) configurations. The only difference is the equation for Z_i , which is now sensitive to the parallel combination of R_1 and R_2 .

9.6 JFET SOURCE-FOLLOWER (COMMON-DRAIN) CONFIGURATION

The JFET equivalent of the BJT emitter-follower configuration is the source-follower configuration of Fig. 9.24. Note that the output is taken off the source terminal and, when the dc supply is replaced by its short-circuit equivalent, the drain is grounded (hence, the terminology common-drain).

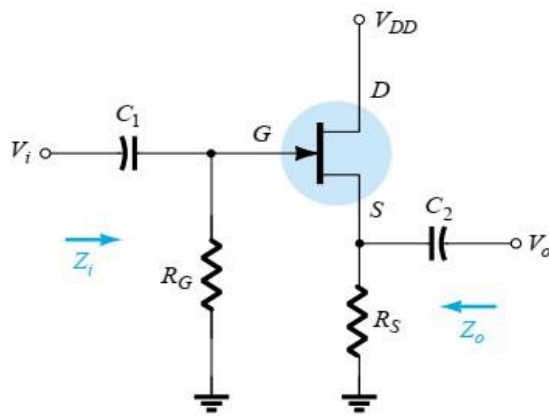


Figure 9.24 JFET source-follower configuration.

Substituting the JFET equivalent circuit will result in the configuration of Fig. 9.25. The controlled source and internal output impedance of the JFET are tied to ground at one end and R_S on the other, with V_o across R_S . Since $g_m V_{gs}$, r_d , and R_S are connected to the same terminal and ground, they can all be placed in parallel as shown in Fig. 9.26. The current source reversed direction but V_{gs} is still defined between the gate and source terminals.

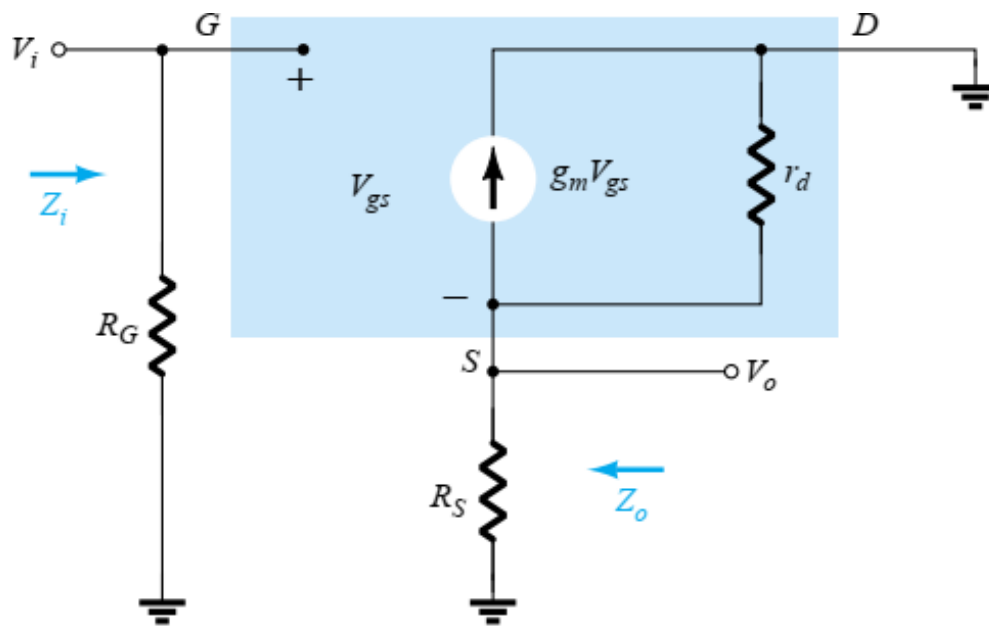


Figure 9.25 Network of Fig. 9.24 following the substitution of the JFET ac equivalent model.

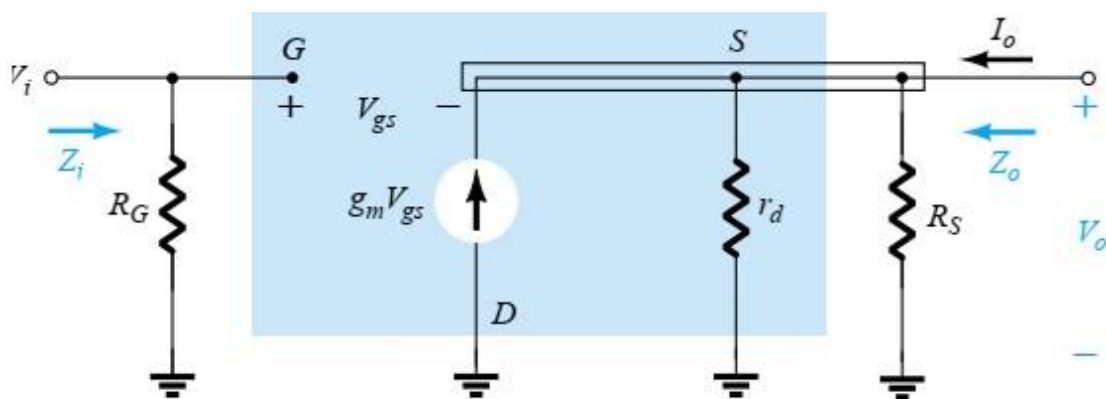


Figure 9.26 Network of Fig. 9.25 redrawn.

Z_i : Figure 9.26 clearly reveals that Z_i is defined by

$$Z_i = R_G \quad (9.33)$$

Z_o : Setting $V_i = 0$ V will result in the gate terminal being connected directly to ground as shown in Fig. 9.27. The fact that V_{gs} and V_o are across the same parallel network results in $V_o = -V_{gs}$.

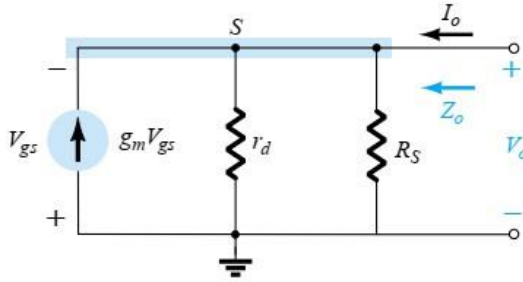


Figure 9.27 Determining Z_o for the network of Fig. 9.24.

Applying Kirchhoff's current law at node s ,

$$\begin{aligned} I_o + g_m V_{gs} &= I_{r_d} + I_{R_S} \\ &= \frac{V_o}{r_d} + \frac{V_o}{R_S} \end{aligned}$$

The result is

$$\begin{aligned} I_o &= V_o \left[\frac{1}{r_d} + \frac{1}{R_S} \right] - g_m V_{gs} \\ &= V_o \left[\frac{1}{r_d} + \frac{1}{R_S} \right] - g_m [-V_o] \\ &= V_o \left[\frac{1}{r_d} + \frac{1}{R_S} + g_m \right] \end{aligned}$$

$$\text{and } Z_o = \frac{V_o}{I_o} = \frac{V_o}{V_o \left[\frac{1}{r_d} + \frac{1}{R_S} + g_m \right]} = \frac{1}{\frac{1}{r_d} + \frac{1}{R_S} + g_m} = \frac{1}{\frac{1}{r_d} + \frac{1}{R_S} + \frac{1}{1/g_m}}$$

which has the same format as the total resistance of three parallel resistors. Therefore,

$$Z_o = r_d \parallel R_S \parallel 1/g_m \quad (9.34)$$

For $r_d \geq 10R_S$,

$$Z_o \cong R_S \parallel 1/g_m \quad r_d \geq 10R_S \quad (9.35)$$

A_v : The output voltage V_o is determined by

$$V_o = g_m V_{gs} (r_d \parallel R_S)$$

and applying Kirchhoff's voltage law around the perimeter of the network of Fig. 9.26 will result in

$$V_i = V_{gs} + V_o$$

and

$$V_{gs} = V_i - V_o$$

so that

$$V_o = g_m (V_i - V_o) (r_d \parallel R_S)$$

or

$$V_o = g_m V_i (r_d \parallel R_S) - g_m V_o (r_d \parallel R_S)$$

and

$$V_o [1 + g_m (r_d \parallel R_S)] = g_m V_i (r_d \parallel R_S)$$

so that

$$A_v = \frac{V_o}{V_i} = \frac{g_m (r_d \parallel R_S)}{1 + g_m (r_d \parallel R_S)} \quad (9.36)$$

In the absence of r_d or if $r_d \geq 10R_S$

$$A_v = \frac{V_o}{V_i} \cong \frac{g_m R_S}{1 + g_m R_S} \quad r_d \geq 10R_S \quad (9.37)$$

Since the bottom of Eq. (9.36) is larger than the numerator by a factor of one, the gain can never be equal to or greater than one (as encountered for the emitter-follower BJT network).

Phase Relationship: Since A_v of Eq. (9.36) is a positive quantity, V_o and V_i are in phase for the JFET source-follower configuration.

9.7 JFET COMMON-GATE CONFIGURATION

The last JFET configuration to be analyzed in detail is the common-gate configuration of Fig. 9.29, which parallels the common-base configuration employed with BJT transistors.

Substituting the JFET equivalent circuit will result in Fig. 9.30. Note the continuing requirement that the controlled source $g_m V_{gs}$ be connected from drain to source with r_d in parallel. The isolation between input and output circuits has obviously been lost since the gate terminal is now connected to the common ground of the network. In addition, the resistor connected between input terminals is no longer R_G but the resistor R_S connected from source to ground. Note also the location of the controlling voltage V_{gs} and the fact that it appears directly across the resistor R_S .

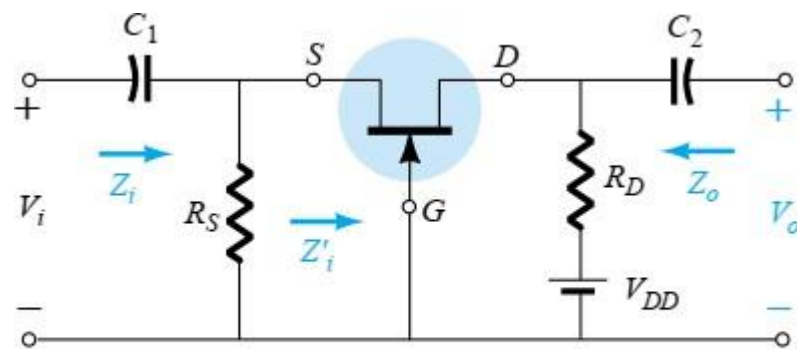


Figure 9.29 JFET common-gate configuration.

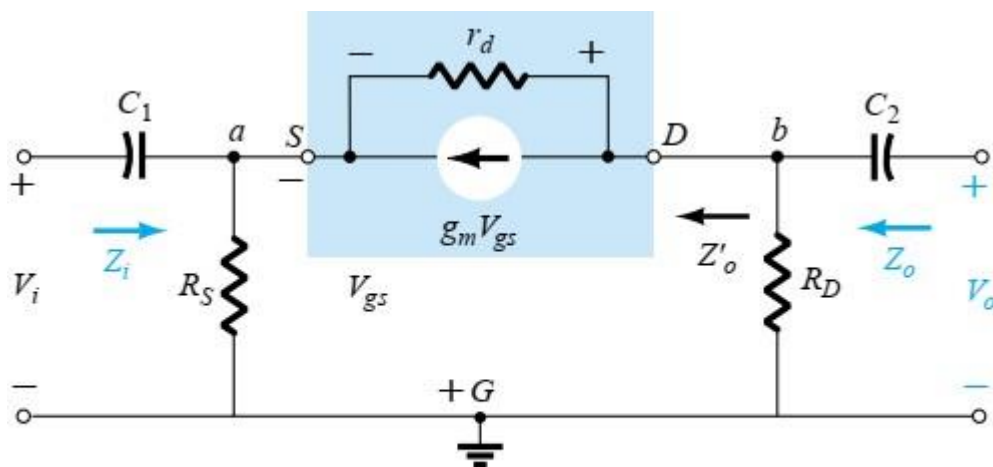


Figure 9.30 Network of Fig. 9.29 following substitution of JFET ac equivalent model.

Z_i : The resistor R_S is directly across the terminals defining Z_i . Let us therefore find the impedance Z'_i of Fig. 9.29, which will simply be in parallel with R_S when Z_i is defined.

The network of interest is redrawn as Fig. 9.31. The voltage $V' = -V_{gs}$. Applying Kirchhoff's voltage law around the output perimeter of the network will result in

$$V' - V_{r_d} - V_{R_D} = 0$$

and

$$V_{r_d} = V' - V_{R_D} = V' - I'R_D$$

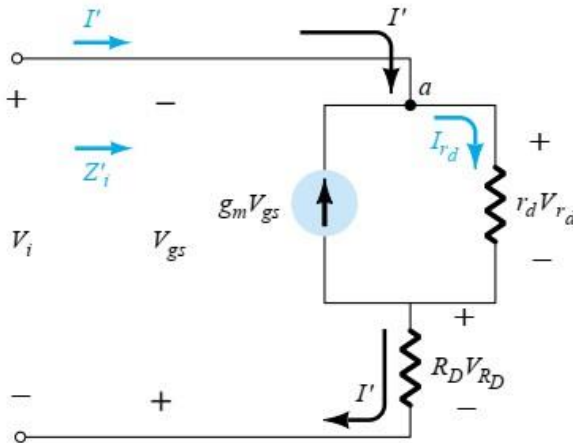


Figure 9.31 Determining Z'_i for the network of Fig. 9.29.

Applying Kirchhoff's current law at node a results in

$$I' + g_m V_{gs} = I_{r_d}$$

and

$$I' = I_{r_d} - g_m V_{gs} = \frac{(V' - I'R_D)}{r_d} - g_m V_{gs}$$

or

$$I' = \frac{V'}{r_d} - \frac{I'R_D}{r_d} - g_m[-V']$$

so that

$$I' \left[1 + \frac{R_D}{r_d} \right] = V' \left[\frac{1}{r_d} + g_m \right]$$

and

$$Z'_i = \frac{V'}{I'} = \frac{\left[1 + \frac{R_D}{r_d} \right]}{\left[g_m + \frac{1}{r_d} \right]} \quad (9.38)$$

or

$$Z'_i = \frac{V'}{I'} = \frac{r_d + R_D}{1 + g_m r_d}$$

and

$$Z_i = R_S \parallel Z'_i$$

results in

$$Z_i = R_S \parallel \left[\frac{r_d + R_D}{1 + g_m r_d} \right] \quad (9.39)$$

If $r_d \geq 10R_D$, Eq. (9.38) permits the following approximation since $R_D/r_d \ll 1$ and $1/r_d \ll g_m$:

$$Z'_i = \frac{\left[1 + \frac{R_D}{r_d} \right]}{\left[g_m + \frac{1}{r_d} \right]} \cong \frac{1}{g_m}$$

and

$$Z_i \cong R_S \parallel 1/g_m \quad r_d \geq 10R_D \quad (9.40)$$

Z_o : Substituting $V_i = 0$ V in Fig. 9.30 will “short-out” the effects of R_S and set V_{gs} to 0 V. The result is $g_m V_{gs} = 0$, and r_d will be in parallel with R_D . Therefore,

$$Z_o = R_D \parallel r_d \quad (9.41)$$

For $r_d \geq 10R_D$,

$$Z_o \cong R_D \quad r_d \geq 10R_D \quad (9.42)$$

A_v : Figure 9.30 reveals that

$$V_i = -V_{gs}$$

and

$$V_o = I_D R_D$$

The voltage across r_d is

$$V_{r_d} = V_o - V_i$$

and

$$I_{r_d} = \frac{V_o - V_i}{r_d}$$

Applying Kirchhoff's current law at node b in Fig. 9.30 results in

and

$$\begin{aligned}
 I_{r_d} + I_D + g_m V_{gs} &= 0 \\
 I_D &= -I_{r_d} - g_m V_{gs} \\
 &= -\left[\frac{V_o - V_i}{r_d} \right] - g_m [-V_i] \\
 I_D &= \frac{V_i - V_o}{r_d} + g_m V_i
 \end{aligned}$$

so that

$$\begin{aligned}
 V_o = I_D R_D &= \left[\frac{V_i - V_o}{r_d} + g_m V_i \right] R_D \\
 &= \frac{V_i R_D}{r_d} - \frac{V_o R_D}{r_d} + g_m R_D V_i
 \end{aligned}$$

and

$$V_o \left[1 + \frac{R_D}{r_d} \right] = V_i \left[\frac{R_D}{r_d} + g_m R_D \right]$$

with

$$A_v = \frac{V_o}{V_i} = \frac{\left[g_m R_D + \frac{R_D}{r_d} \right]}{\left[1 + \frac{R_D}{r_d} \right]} \quad (9.43)$$

For $r_d \geq 10R_D$, the factor R_D/r_d of Eq. (9.43) can be dropped as a good approximation and

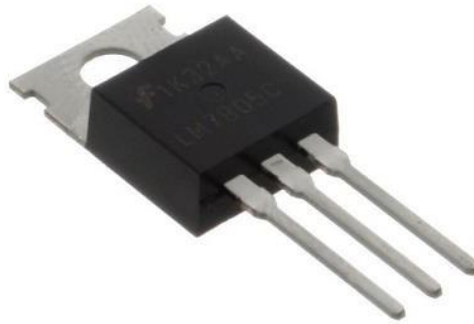
$$A_v = g_m R_D \quad r_d \geq 10R_D \quad (9.44)$$

Phase Relationship: The fact that A_v is a positive number will result in an *in-phase* relationship between V_o and V_i for the common-gate configuration.

MOSFET

FETs have a few disadvantages like high drain resistance, moderate input impedance and slower operation. To overcome these disadvantages, the MOSFET which is an advanced FET is invented.

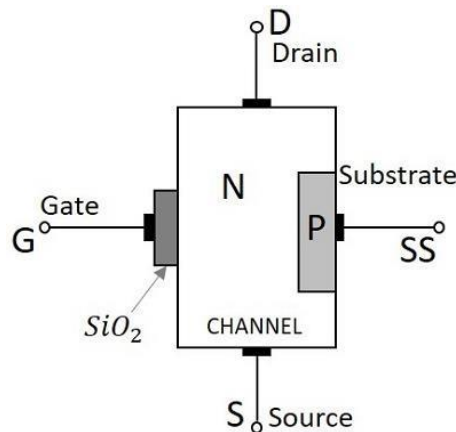
MOSFET stands for Metal Oxide Silicon Field Effect Transistor or Metal Oxide Semiconductor Field Effect Transistor. This is also called as IGFET meaning Insulated Gate Field Effect Transistor. The FET is operated in both depletion and enhancement modes of operation. The following figure shows how a practical MOSFET looks like.



Construction of a MOSFET

The construction of a MOSFET is a bit similar to the FET. An oxide layer is deposited on the substrate to which the gate terminal is connected. This oxide layer acts as an insulator (SiO_2 insulates from the substrate), and hence the MOSFET has another name as IGFET. In the construction of MOSFET, a lightly doped substrate, is diffused with a heavily doped region. Depending upon the substrate used, they are called as **P-type** and **N-type** MOSFETs.

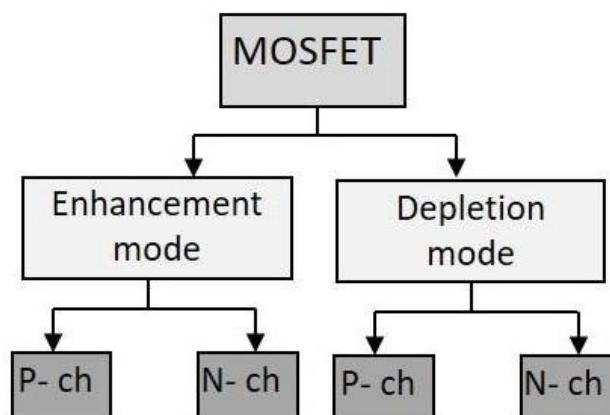
The following figure shows the construction of a MOSFET.



The voltage at gate controls the operation of the MOSFET. In this case, both positive and negative voltages can be applied on the gate as it is insulated from the channel. With negative gate bias voltage, it acts as **depletion MOSFET** while with positive gate bias voltage it acts as an **Enhancement MOSFET**.

Classification of MOSFETs

Depending upon the type of materials used in the construction, and the type of operation, the MOSFETs are classified as in the following figure.

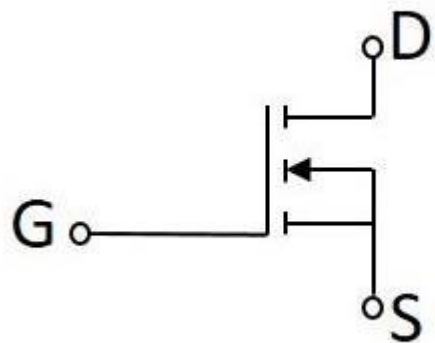


- P- ch = P- channel
- N- ch = N- channel

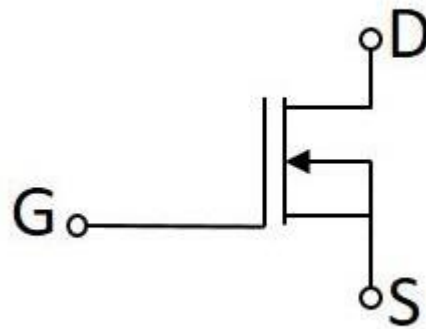
After the classification, let us go through the symbols of MOSFET.

The **N-channel MOSFETs** are simply called as **NMOS**. The symbols for N-channel MOSFET are as given below.

Symbols of N-Channel MOSFET



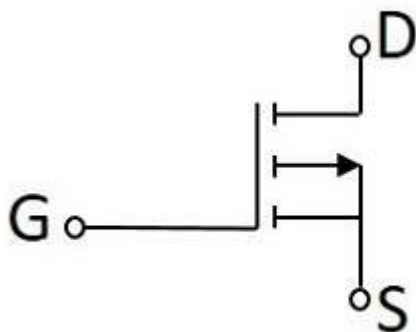
Enhancement Mode



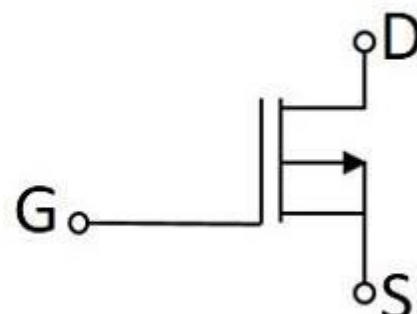
Depletion Mode

The **P-channel MOSFETs** are simply called as **PMOS**. The symbols for P-channel MOSFET are as given below.

Symbols of P-Channel MOSFET



Enhancement Mode

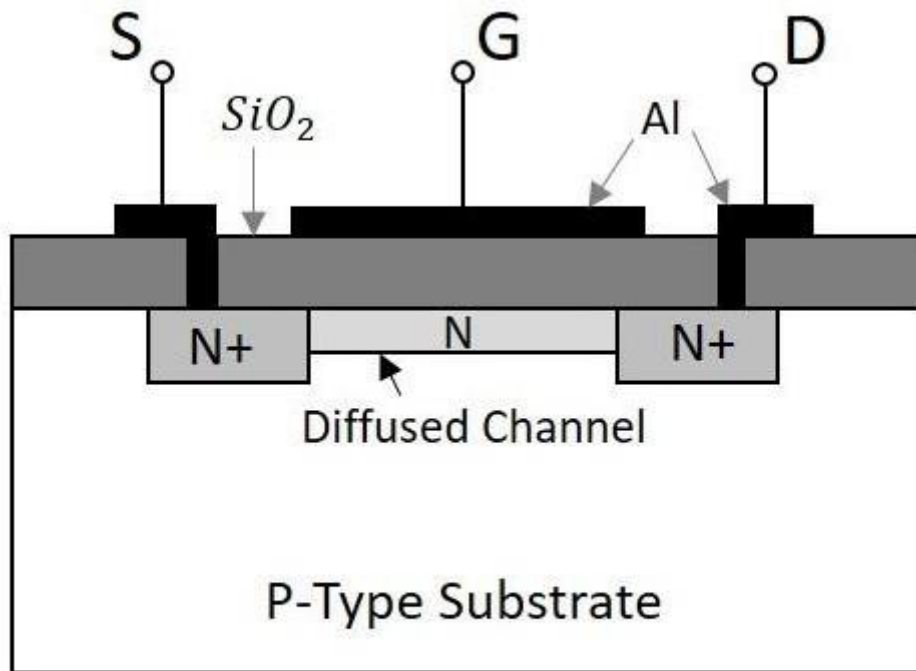


Depletion Mode

Now, let us go through the constructional details of an N-channel MOSFET. Usually an N-channel MOSFET is considered for explanation as this one is mostly used. Also, there is no need to mention that the study of one type explains the other too.

Construction of N- Channel MOSFET

Let us consider an N-channel MOSFET to understand its working. A lightly doped P-type substrate is taken into which two heavily doped N-type regions are diffused, which act as source and drain. Between these two N⁺ regions, there occurs diffusion to form an N-channel, connecting drain and source.



Structure of N-channel MOSFET

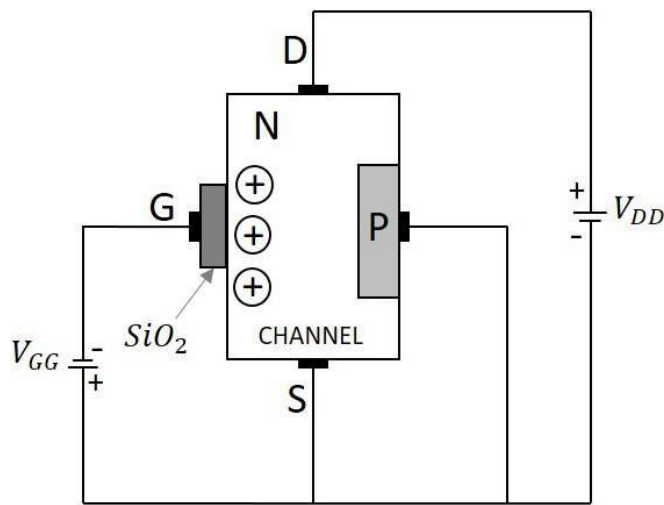
A thin layer of **Silicon dioxide (SiO_2)** is grown over the entire surface and holes are made to draw ohmic contacts for drain and source terminals. A conducting layer of **aluminium** is laid over the entire channel, upon this SiO_2 layer from source to drain which constitutes the gate. The SiO_2 **substrate** is connected to the common or ground terminals.

Because of its construction, the MOSFET has a very less chip area than BJT, which is 5% of the occupancy when compared to bipolar junction transistor. This device can be operated in modes. They are depletion and enhancement modes. Let us try to get into the details.

Working of N - Channel depletion mode MOSFET

For now, we have an idea that there is no PN junction present between gate and channel in this, unlike a FET. We can also observe that, the diffused channel N between two N⁺ regions, the **insulating dielectric SiO_2** and the aluminium metal layer of the gate together form a **parallel plate capacitor**.

If the NMOS has to be worked in depletion mode, the gate terminal should be at negative potential while drain is at positive potential, as shown in the following figure.



Working of MOSFET in depletion mode

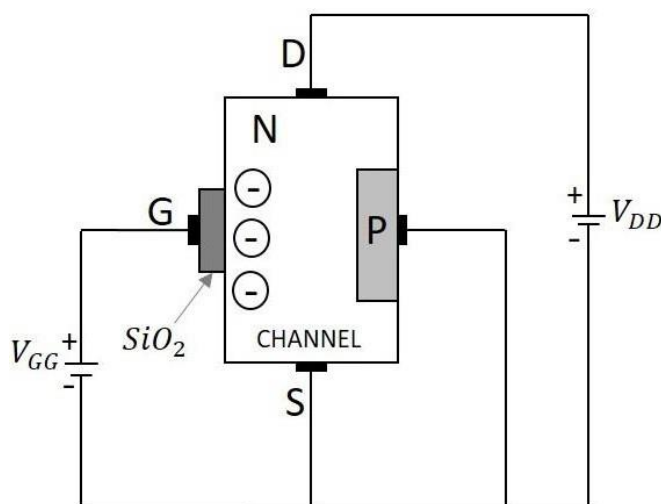
When no voltage is applied between gate and source, some current flows due to the voltage between drain and source. Let some negative voltage is applied at V_{GG} . Then the minority carriers i.e. holes, get attracted and settle near SiO_2 layer. But the majority carriers, i.e., electrons get repelled.

With some amount of negative potential at V_{GG} a certain amount of drain current I_D flows through source to drain. When this negative potential is further increased, the electrons get depleted and the current I_D decreases. Hence the more negative the applied V_{GG} , the lesser the value of drain current I_D will be.

The channel nearer to drain gets more depleted than at source like in FET and the current flow decreases due to this effect. Hence it is called as depletion mode MOSFET.

Working of N-Channel MOSFET Enhancement Mode

The same MOSFET can be worked in enhancement mode, if we can change the polarities of the voltage V_{GG} . So, let us consider the MOSFET with gate source voltage V_{GG} being positive as shown in the following figure.



Working of MOSFET in Enhancement mode

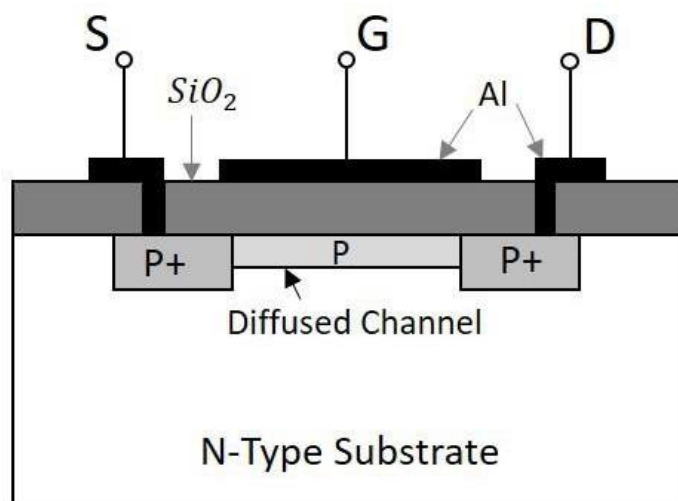
When no voltage is applied between gate and source, some current flows due to the voltage between drain and source. Let some positive voltage is applied at V_{GG} . Then the minority

carriers i.e. holes, get repelled and the majority carriers i.e. electrons gets attracted towards the SiO_2 layer.

With some amount of positive potential at V_{GG} a certain amount of drain current I_D flows through source to drain. When this positive potential is further increased, the current I_D increases due to the flow of electrons from source and these are pushed further due to the voltage applied at V_{GG} . Hence the more positive the applied V_{GG} , the more the value of drain current I_D will be. The current flow gets enhanced due to the increase in electron flow better than in depletion mode. Hence this mode is termed as **Enhanced Mode MOSFET**.

P - Channel MOSFET

The construction and working of a PMOS is same as NMOS. A lightly doped **n-substrate** is taken into which two heavily doped **P+ regions** are diffused. These two P+ regions act as source and drain. A thin layer of SiO_2 is grown over the surface. Holes are cut through this layer to make contacts with P+ regions, as shown in the following figure.



Structure of P-channel MOSFET

Working of PMOS

When the gate terminal is given a negative potential at V_{GG} than the drain source voltage V_{DD} , then due to the P+ regions present, the hole current is increased through the diffused P channel and the PMOS works in **Enhancement Mode**.

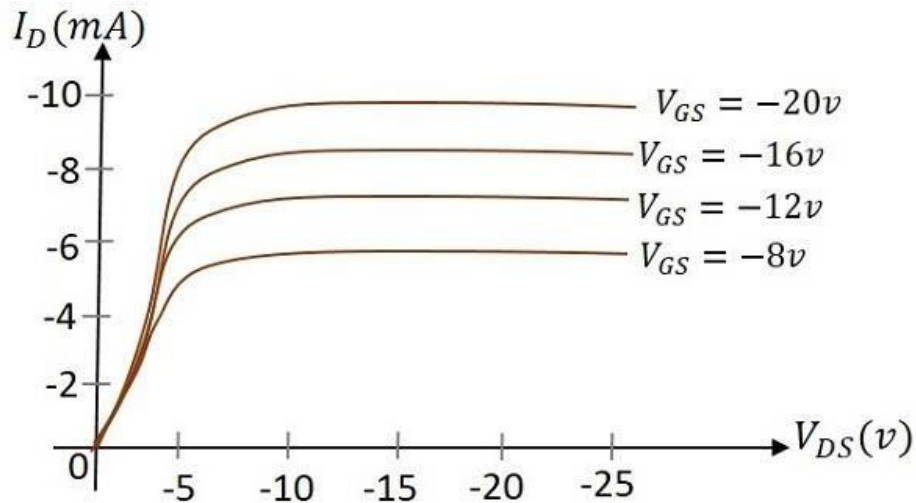
When the gate terminal is given a positive potential at V_{GG} than the drain source voltage V_{DD} , then due to the repulsion, the depletion occurs due to which the flow of current reduces. Thus PMOS works in **Depletion Mode**. Though the construction differs, the working is similar in both the type of MOSFETs. Hence with the change in voltage polarity both of the types can be used in both the modes.

This can be better understood by having an idea on the drain characteristics curve.

Drain Characteristics

The drain characteristics of a MOSFET are drawn between the drain current I_D and the drain source voltage V_{DS} . The characteristic curve is as shown below for different values of inputs.

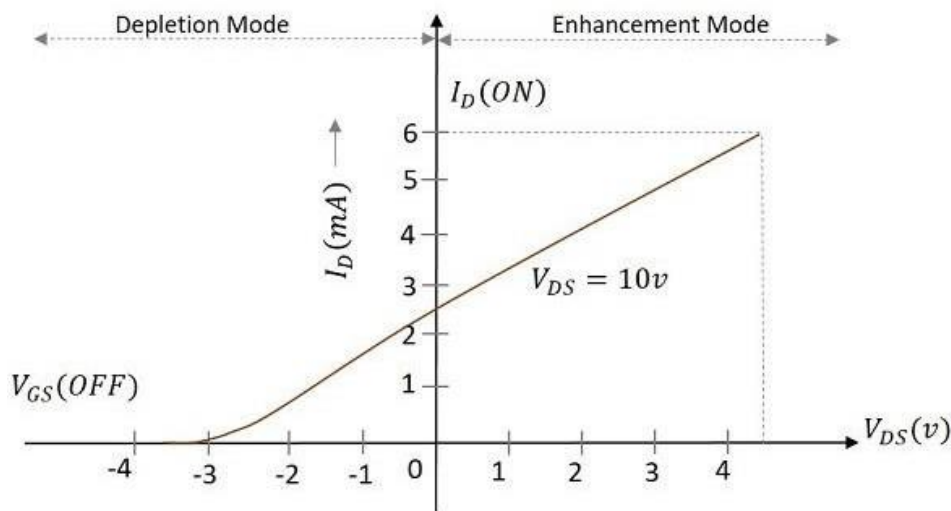
Drain Characteristics of MOSFET



Actually when V_{DS} is increased, the drain current I_D should increase, but due to the applied V_{GS} , the drain current is controlled at certain level. Hence the gate current controls the output drain current.

Transfer Characteristics

Transfer characteristics define the change in the value of V_{DS} with the change in I_D and V_{GS} in both depletion and enhancement modes. The below transfer characteristic curve is drawn for drain current versus gate to source voltage.



Transfer Characteristics of a MOSFET

Comparison between BJT, FET and MOSFET

Now that we have discussed all the above three, let us try to compare some of their properties.

TERMS	BJT	FET	MOSFET
Device type	Current controlled	Voltage controlled	Voltage Controlled
Current flow	Bipolar	Unipolar	Unipolar
Terminal s	Not interchangeable	Interchangeable	Interchangeable
Operational modes	No modes	Depletion mode only	Both Enhancement and Depletion modes
Input impedance	Low	High	Very high
Output resistance	Moderate	Moderate	Low
Operational speed	Low	Moderate	High
Noise	High	Low	Low
Thermal stability	Low	Better	High

So far, we have discussed various electronic components and their types along with their construction and working. All of these components have various uses in the electronics field