

Government polytechnic kendrapara



DEPARTMENT OF ELECTRONICS AND TELECOMMUNICATION ENGINEERING LECTURE NOTES

Semester : 4th

Subject: ANALOG ELECTRONICS & LINEAR IC (TH-4)

**Prepared by : Rabindra Kumar Satapathy,
(PTGF. Electronics & Telecommunication Engg.)**

UNIT-1: PN JUNCTION DIODE

SEMICONDUCTOR:

Semiconductors (*e.g. germanium, silicon etc.*) are those substances whose electrical conductivity lies in between conductors and insulators. In terms of energy band, the valence band is almost filled and conduction band is almost empty. Further, the energy gap between valence and conduction bands is very small. The semiconductor has:

- Filled valence band
- Empty conduction band
- Small energy gap or forbidden gap (1 eV) between valence and conduction bands.
- Semiconductor virtually behaves as an insulator at low temperatures. However, even at room temperature, some electrons cross over to the conduction band, imparting little conductivity (i.e. conductor).

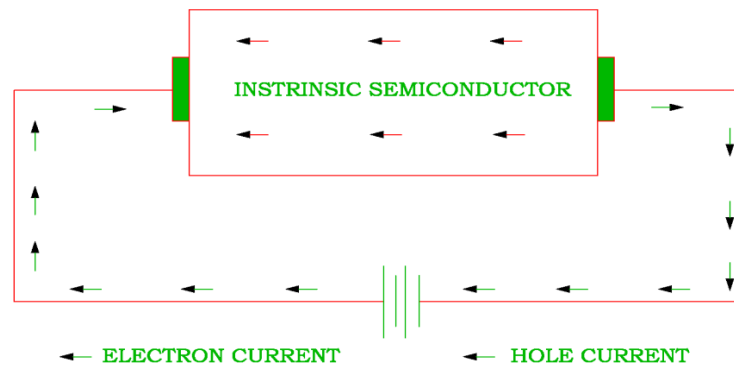
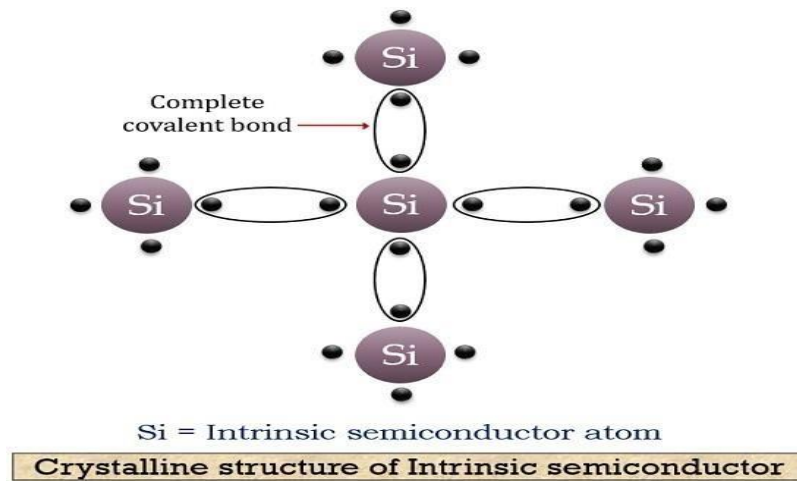
Types of semiconductors:-

Semiconductors are classified into two types:-

- ❖ Intrinsic semiconductors
- ❖ Extrinsic semiconductors
- Extrinsic semiconductors are also of two types:-
 - P-type semiconductors
 - N-type semiconductors

Intrinsic semiconductors

- A semiconductor in an extremely pure form is known as an intrinsic semiconductor. When electric field is applied across an intrinsic semiconductor, the current conduction takes place by two processes i.e. by free electrons and holes.
- The free electrons are produced due to the breaking up of some covalent bonds by thermal energy. At the same time, holes are created in the covalent bonds.
- Under the influence of electric field, conduction through the semiconductor is by both free electrons and holes. Therefore, the total current inside the semiconductor is the sum of currents due to free electrons and holes.



Extrinsic semiconductors

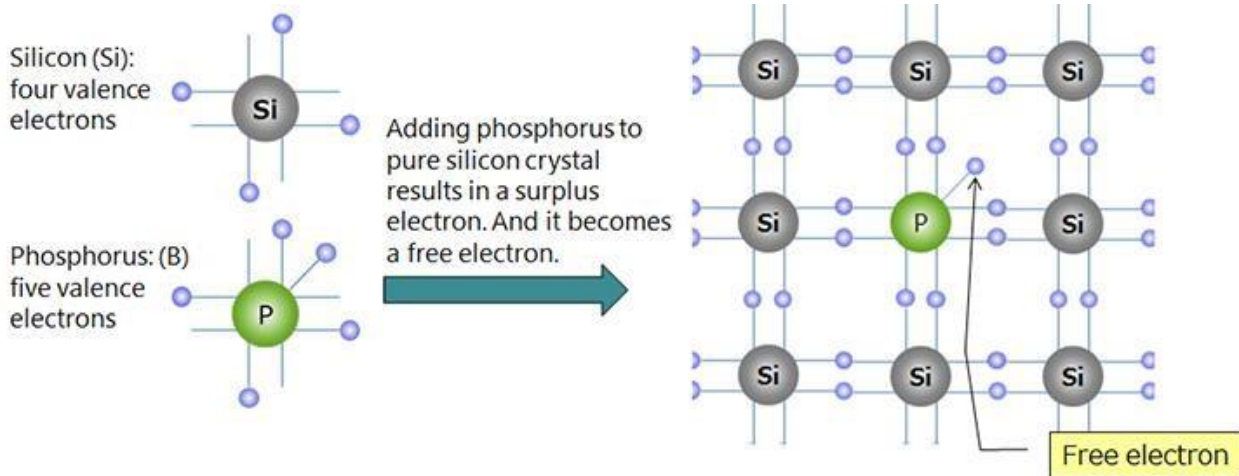
An extrinsic semiconductor is a semiconductor doped by addition of small amount impurity. The process of adding impurities to a semiconductor is known as doping. The purpose of adding impurity is to increase either the number of free electrons or holes in the semiconductor crystal. Depending upon the type of impurity added, extrinsic semiconductors are classified into:

- n-type semiconductor
- p-type semiconductor

N-type semiconductors

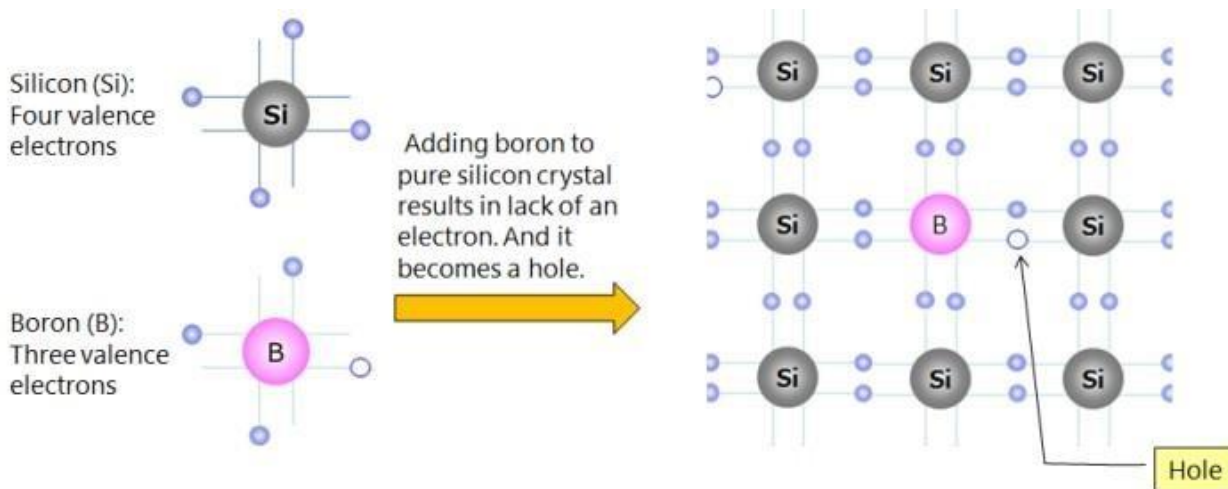
When a small amount of pentavalent impurity is added to a pure semiconductor, it is known as n-type semiconductor. Examples of pentavalent impurities are arsenic, antimony, Bismuth and Phosphorous etc. Such impurities are known as donor impurities

because they donate or provide free electrons to the semiconductor crystal. Electrons are said to be the majority carriers whereas holes are the minority carriers.



P-type semiconductors

When a small amount of trivalent impurity is added to a pure semiconductor, it is called p-type Semiconductor. Examples of trivalent impurities are gallium, indium, boron etc. The addition of trivalent impurity provides a large number of holes in the semiconductor. Such impurities are known as acceptor impurities because the holes created can accept the electrons. In a p type semiconductor holes are the majority carriers and electrons are the minority carriers.



PN JUNCTION:-

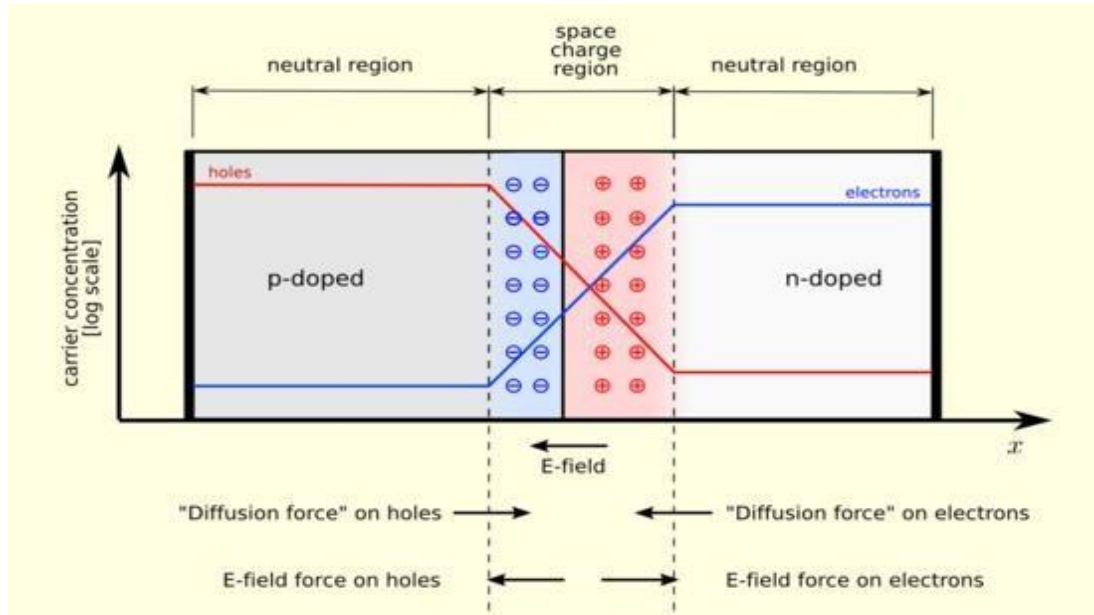
Zero Biased Junction Diode

When a diode is connected in a Zero Bias condition, no external potential energy is applied to the PN junction. However if the diodes terminals are shorted together, a few holes (majority carriers) in the P-type material with enough energy to overcome the potential barrier will move across the junction against this barrier potential. This is known as the “Forward Current” and is referenced as I_F . Likewise, holes generated in the N-type material (minority carriers), find this situation favorable and move across the junction in the opposite direction. This is known as the “Reverse Current” and is referenced as I_R . This transfer of electrons and holes back and forth across the PN junction is known as diffusion.

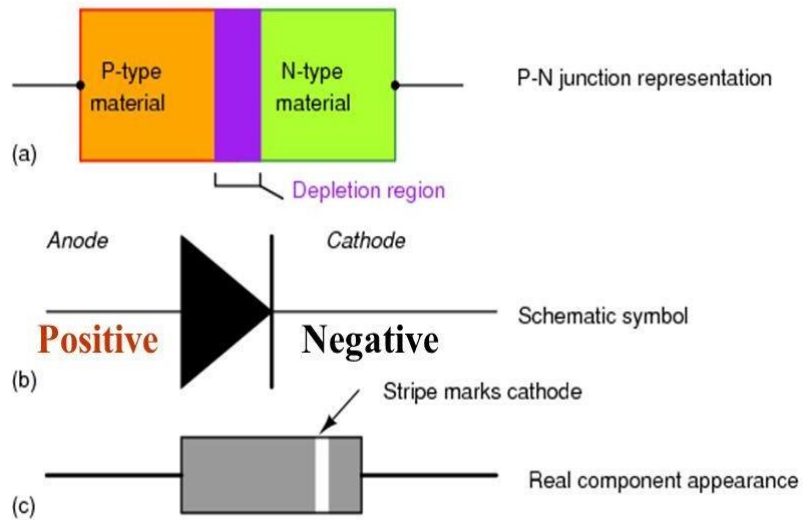
The potential barrier that now exists discourages the diffusion of any more majority carriers across the junction. However, the potential barrier helps minority carriers (few free electrons in the P-region and few holes in the N-region) to drift across the junction.

Then an “Equilibrium” or balance will be established when the majority carriers are equal and both moving in opposite directions, so that the net result is zero current flowing in the circuit. When this occurs the junction is said to be in a state of “Dynamic Equilibrium”.

The minority carriers are constantly generated due to thermal energy so this state of equilibrium can be broken by raising the temperature of the PN junction causing an increase in the generation of minority carriers, thereby resulting in an increase in leakage current but an electric current cannot flow since no circuit has been connected to the PN junction.

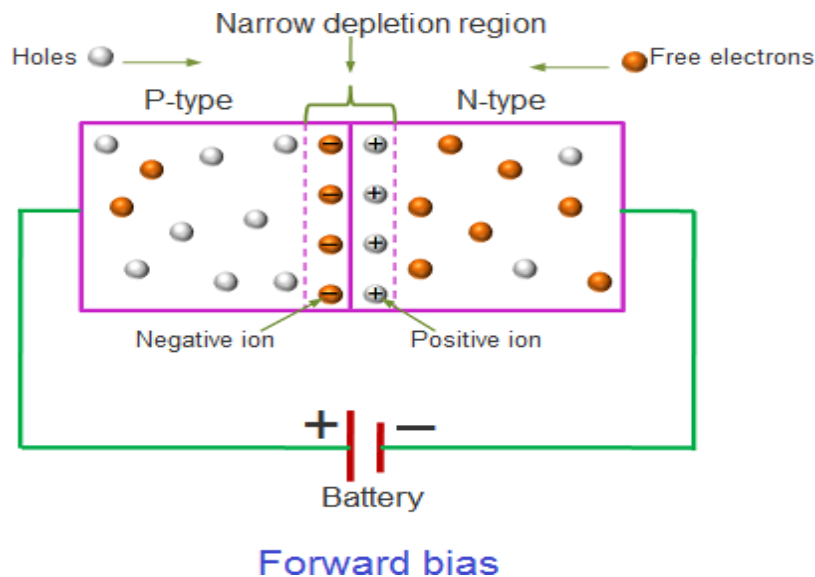


Diode Symbol



FORWARD BIASING:-

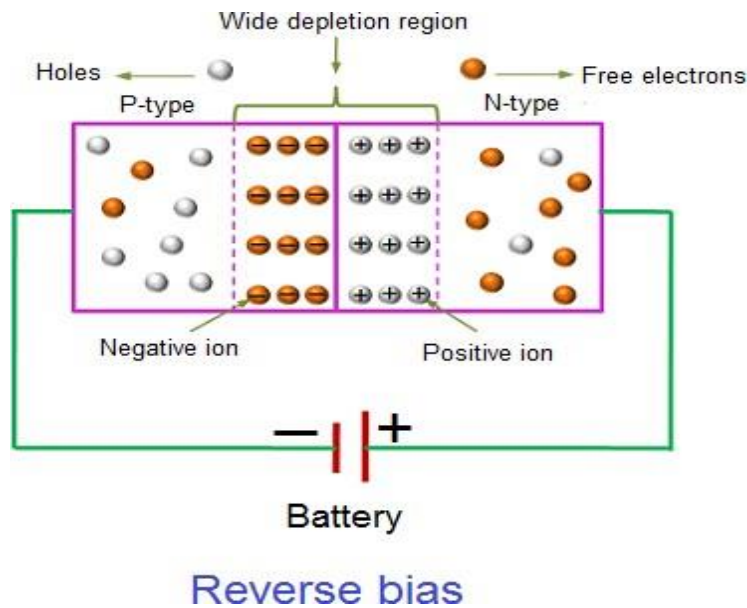
- A P-N junction diode is said to be forward biased when the positive terminal of a cell or battery is connected to the p-side of the junction and the negative terminal to the n side.
 - When diode is forward-biased the depletion region narrows and consequently, the potential barrier is lowered.
 - This causes the majority charge carriers of each region to cross into the other region.
 - The electrons travel from the n-side to the p-side and go to the positive terminal of the battery.
 - The holes that travel from the p-side to the n-side combine with the electrons injected into the n-region from the negative terminal of the battery.
- With the increase in forward bias voltage, the depletion region eventually becomes thin, thus reducing electrical resistance.
- This causes electrons to pass through the junction resulting in an exponential rise in current.
 - This way the diode conducts when forward-biased.



REVERSE BIASING:-

- A pn-junction diode is said to be reverse biased when the positive terminal of a cell or battery is connected to the n-side of the junction and the negative terminal to the p-side.
- When reverse biased, the depletion region widens and the potential barrier is increased.
- The polarity of the battery extracts the majority charge carriers of each region.
- The holes in the p-region from the electrons injected into the p-region from the negative terminal of the battery.
- The electrons in the n-region go to the positive terminal of the battery.

- This way, the majority charge carrier concentration in each region decreases against the equilibrium values and the reverse biased junction diode has a high resistance.
- Thus, the diffusion current across the junction becomes zero.
- Thus, the diode does not conduct when reverse biased and is said to be in a quiescent or non-conducting state i.e., it acts as an open switch.

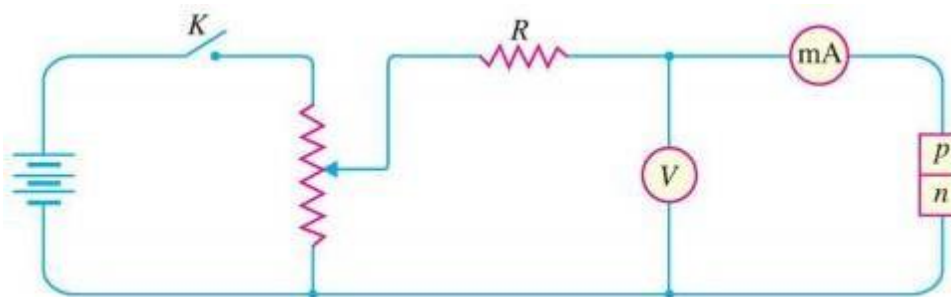


VI CHARACTERISTICS OF PN JUNCTION:-

Volt-ampere (V-I) characteristics of a pn junction or semiconductor diode is the curve between voltage across the junction and the current through the circuit.

Normally the voltage is taken along the x-axis and current along y-axis.

The circuit connection for determining the V-I characteristics of a pn junction is shown in the figure below.



The characteristics can be explained under three cases, such as :

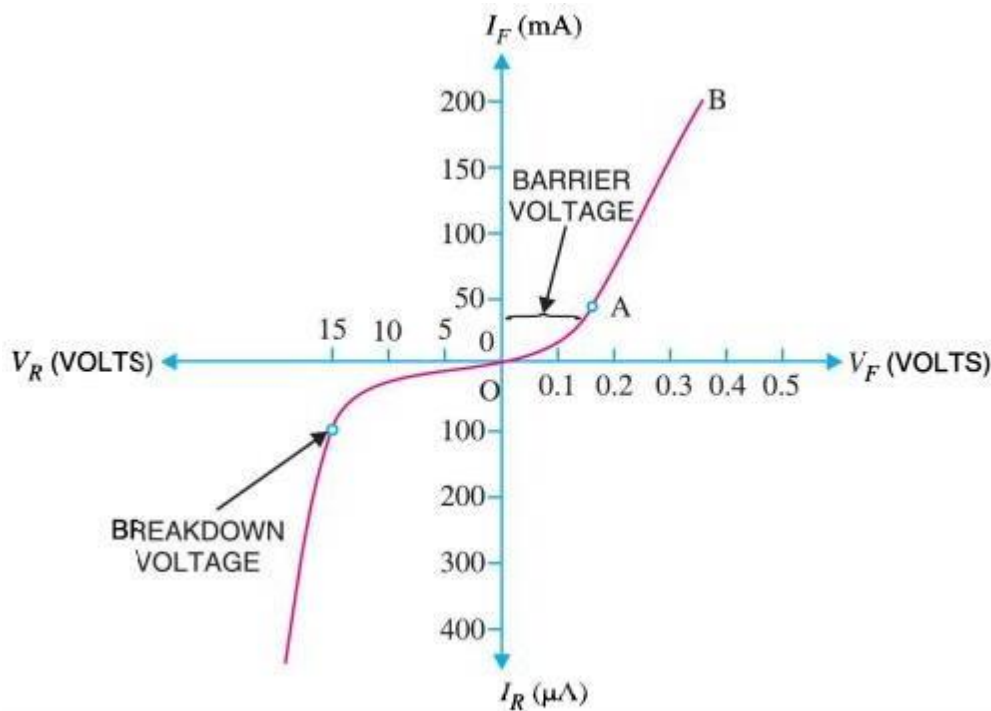
1. Zero bias
2. Forward bias
3. Reverse bias

Case-1: Zero Bias

In zero bias condition, no external voltage is applied to the pn junction i.e. the circuit is open at K.

Hence, the potential barrier at the junction does not permit current flow.

Therefore, the circuit current is zero at $V=0$ V, as indicated by point O in figure below.



Case-2: Forward Bias

In forward biased condition, p-type of the pn junction is connected to the positive terminal and n-type is connected to the negative terminal of the external voltage.

This results in reduced potential barrier.

At some forward voltage i.e. 0.7 V for Si and 0.3 V for Ge, the potential barrier is almost eliminated and the current starts flowing in the circuit.

From this point, the current increases with the increase in forward voltage. Hence a curve OB is obtained with forward bias as shown in figure above.

From the forward characteristics, it can be noted that at first i.e. region OA, the current increases very slowly and the curve is non-linear. It is because in this region the external voltage applied to the pn junction is used in overcoming the potential barrier.

However, once the external voltage exceeds the potential barrier voltage, the potential barrier is eliminated and the pn junction behaves as an ordinary conductor. Hence, the curve AB rises very sharply with the increase in external voltage and the curve is almost linear.

Case-3: Reverse Bias

In reverse bias condition, the p-type of the pn junction is connected to the negative terminal and n-type is connected to the positive terminal of the external voltage.

This results in increased potential barrier at the junction.

Hence, the junction resistance becomes very high and as a result practically no current flows through the circuit.

However, a very small current of the order of μA , flows through the circuit in practice. This is known as reverse saturation current (I_s) and it is due to the minority carriers in the junction.

As we already know, there are few free electrons in p-type material and few holes in n-type material. These free electrons in p-type and holes in n-type are called minority carriers.

The reverse bias applied to the pn junction acts as forward bias to their minority carriers and hence, small current flows in the reverse direction.

If the applied reverse voltage is increased continuously, the kinetic energy of the minority carriers may become high enough to knock out electrons from the semiconductor atom.

At this stage breakdown of the junction may occur. This is characterized by a sudden increase of reverse current and a sudden fall of the resistance of barrier region. This may destroy the junction permanently.

DC LOAD LINE-

A circuit supplied with dc power as the external source of the circuit. There exist both alternating and direct currents in the circuit. The reactive components of the circuits are

made zero and the straight line is drawn above the voltage-current characteristics curves. Hence these results in the formation of intersecting point referred to an operating point. The straight that is drawn for this purpose is defined as the DC load line.

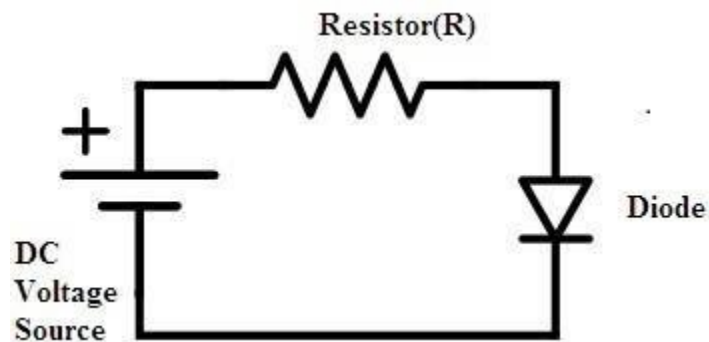
DC Load Line of a Diode and Its Equation

DC load line for a non-linear device is drawn by making the reactive components as zero. Hence a diode is considered as a non-linear device and its voltage and current characteristics are exponentially related to each other. The formation of the intersection point for the characteristic curve and the straight line or dc-load line can be analyzed better by considering the example for the diode as in forward bias condition.

Let us consider a diode connected to the resistor(R), source of voltage (V_{DD}) in series. The diode is connected in forward bias so that the forward current and the forward voltages flowing through the circuit. As per the Kirchhoff's current law, the current flowing through the diode (I_D) and the resistor (I_R) is equal.

$$I_D = I_R$$

Analysis of the circuit is done by applying Kirchhoff's voltage law (KVL). This law results in the formation of the final equation for the dc load line. Here the dc voltage is the biasing voltage of the circuit by keeping any further reactive components as zero.



Diode-operating-in-forward-bias-for-the-analysis-of-dc-load

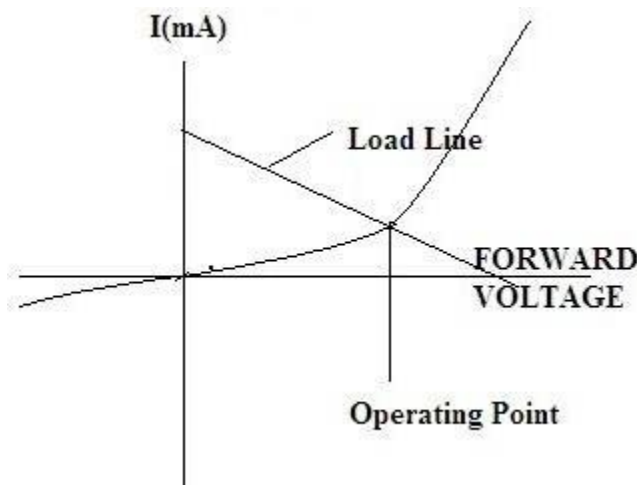
Once the Kirchhoff's voltage law is applied to the circuit an equation is obtained for voltages and currents in the circuit.

$$V_{DD} = V_D + I_D R$$

$$V_D = I_D R - V_{DD}$$

Where V_{DD} , is the applied dc source voltage and V_D is the voltage across the diode. Hence the above can be considered as the equation for the diode. The voltage and current characteristics of the diode in forward bias condition can be drawn. By our previous analysis on the condition of the diode in forward bias applied a voltage and the generated current in the circuit are exponentially related to each other.

After a certain cut-off voltage, the diode starts operating in forward bias condition. To this slope, the technique is considered and a straight line on the v-i characteristics is drawn. The slope here for the above general circuit for the diode is V_{DD}/R .



Dc-load-line-and-the-formation-of-operating-point

In this way, the analysis for the dc-load line is done for the non-linear device like a diode. Depending on the type of non-linear device some part of the analysis differs but the technique remains the same. This type of method comes under the graphical analysis because here the characteristics curve is considered for the formation of the dc-load line.

IMPORTANT TERMS SUCH AS IDEAL DIODE, KNEE VOLTAGE

Ideal Diode:

An ideal diode is one kind of an electrical component that performs like an ideal conductor when voltage is applied in forward bias and like an ideal insulator when the voltage is applied in reverse bias. So when +ve voltage is applied across the anode toward the cathode, the diode performs forward current immediately. When a voltage is applied in reverse bias, then it performs no current at all. This diode operates like a switch. When the diode is in forward bias, it works like a closed switch. Whereas, if an ideal diode is in reverse bias, then it works like an open switch.

Knee Voltage:

The forward voltage at which the current through the junction starts increasing rapidly, is called knee voltage or cut-in voltage. It is generally 0.6V for a diode.

JUNCTIONS BREAKDOWN:

- ✓ The Avalanche Breakdown and Zener Breakdown are two different mechanisms by which a PN junction breaks.
- ✓ The Zener and Avalanche breakdown both occur in diode under reverse bias.
- ✓ The avalanche breakdown occurs because of the ionization of electrons and hole pairs whereas the Zener diode occurs because of heavy doping.

AVALANCHE BREAKDOWN:

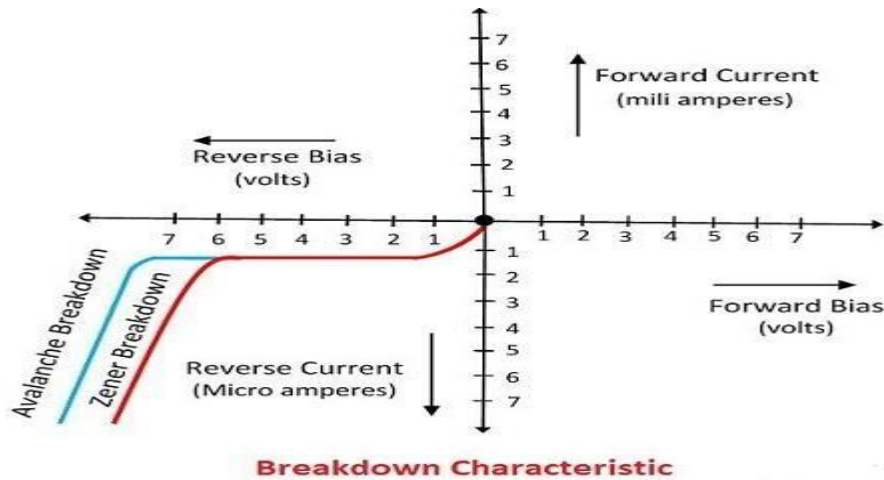
The mechanism of avalanche breakdown occurs because of the reverse saturation current. The P-type and N-type material together forms the PN-junction. The depletion region develops at the junction where the P and N-type material contact. The P and N-type materials of the PN junction are not perfect, and they have some impurities in it, i.e., the p-type material has some electrons, and the N-type material has some hole in it. The width of the depletion region varies. Their width depends on the bias applied to the terminal of P and N region. The reverse bias increases the electrical field across the depletion region. When the high electric field exists across the depletion, the velocity of minority charge carrier crossing the depletion region increases. These carriers collide with the atoms of the crystal. Because of the violent collision, the charge carrier takes out the electrons from the atom. The collision increases the electron-hole pair. As the electron-hole induces in the high electric field, they are quickly separated and collide with the other atoms of the crystals. The process is continuous, and the electric field becomes so much higher than the reverse current starts flowing in the PN junction. The process is known as the Avalanche breakdown. After the breakdown, the junction cannot regain its original position because the diode is completely burnt off.

ZENER BREAKDOWN:

The PN junction is formed by the combination of the p-type and the n-type semiconductor material. The combination of the P-type and N-type regions creates the depletion region. The width of the depletion region depends on the doping of the P and N-type semiconductor material. If the material is heavily doped, the width of the depletion region becomes very thin. The phenomenon of the Zener breakdown occurs in the very thin depletion region. The thin depletion region has more numbers of free electrons. The reverse bias applied across the PN junction develops the electric field intensity across the depletion region. The strength of the electric field intensity becomes very high. The electric field intensity increases the kinetic energy of the free charge carriers. Thus the carrier starts jumping from one region to another. These energetic charge carriers collide with the atoms of the p-type and n-type material and produce the electron-hole pairs. The reverse current starts flowing in the junction because of which depletion region becomes entirely vanishes. This process is known as the Zener breakdown. In Zener breakdown, the

junction is not completely damaged. The depletion region regains its original position after the removal of the reverse voltage. The voltage of Zener breakdown is less than the Avalanche breakdown.

Breakdown Characteristic Graph



The graphical representation of the Avalanche and Zener breakdown is shown in the figure below.

RECTIFIER CIRCUITS

RECTIFIERS AND CLASSIFICATION OF RECTIFIER:-

Rectifier is an electronic device which converts the alternating current to unidirectional current, in other words rectifier converts the AC voltage to DC voltage. We use rectifier in almost all the electronic devices mostly in the power

supply section to convert the main voltage into DC voltage. Every electronic device will work on the DC voltage supply only.

Rectifiers are classified according to the period of conduction.

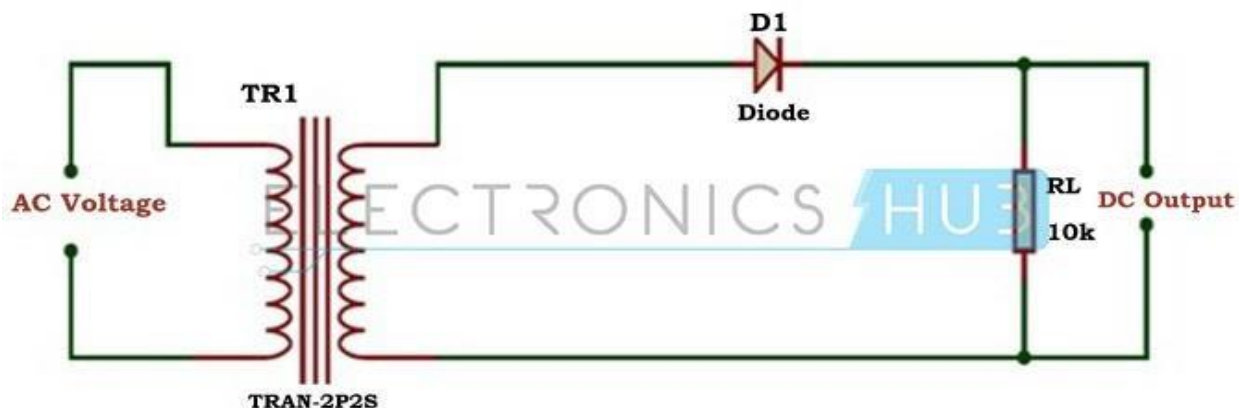
They are

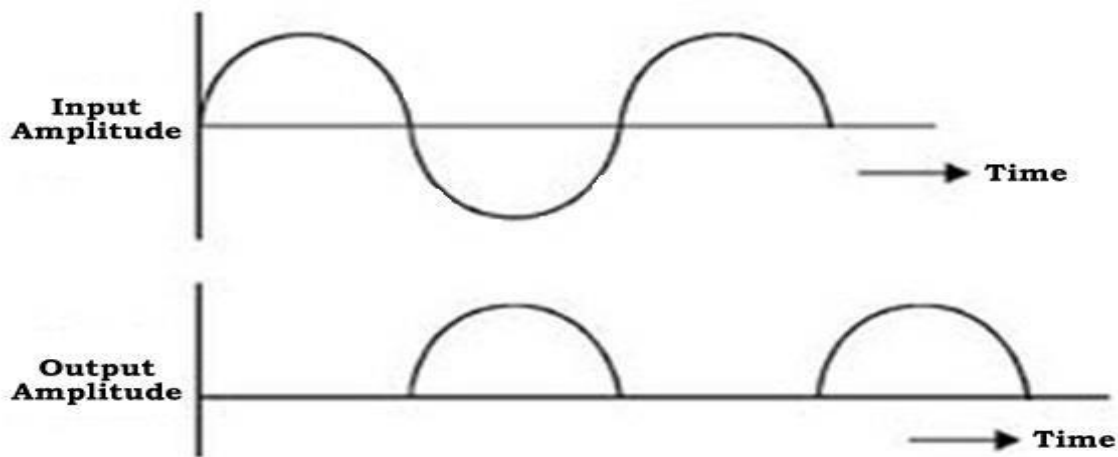
1. Half Wave Rectifier
2. Full Wave Rectifier

ANALYSIS OF HALF WAVE, FULL WAVE CENTRE TAPPED AND BRIDGE RECTIFIERS:

Half Wave Rectifier:

The half wave rectifier is a type of rectifier that rectifies only half cycle of the waveform. This article describes the half wave rectifier circuit working. The half rectifier consist a step down transformer, a diode connected to the transformer and a load resistance connected to the cathode end of the diode. The circuit diagram of half wave transformer is shown below:





The main supply voltage is given to the transformer which will increase or decrease the voltage and give to the diode. In most of the cases we will decrease the supply voltage by using the step down transformer here also the output of the step down transformer will be in AC. This decreased AC voltage is given to the diode which is connected serial to the secondary winding of the transformer, diode is electronic component which will allow only the forward bias current and will not allow the reverse bias current. From the diode we will get the pulsating DC and give to the load resistance R_L .

Working of Half Wave Rectifier:

The input given to the rectifier will have both positive and negative cycles. The half rectifier will allow only the positive half cycles and omit the negative half cycles. So first we will see how half wave rectifier works in the positive half cycles.

Positive Half Cycle:

- In the positive half cycles when the input AC power is given to the primary winding of the step down transformer, we will get the decreased voltage at the secondary winding which is given to the diode.
- The diode will allow current flowing in clock wise direction from anode to cathode in the forward bias (diode conduction will take place in forward bias) which will generate only the positive half cycle of the AC.
- The diode will eliminate the variations in the supply and give the pulsating DC voltage to the load resistance R_L . We can get the pulsating DC at the Load resistance.

Negative Half Cycle:

- In the negative half cycle the current will flow in the anti-clockwise direction and the diode will go in to the reverse bias. In the reverse bias the diode will not conduct so, no current is flown from anode to cathode, and we cannot get any power at the load resistance.
- Only small amount of reverse current is flown from the diode but this current is almost negligible. And voltage across the load resistance is also zero.

Characteristics of Half Wave Rectifier:

There are some characteristics to the half wave rectifier they are

Efficiency: The efficiency is defined as the ratio of input AC to the output DC.

Efficiency, $\eta = P_{dc} / P_{ac}$

DC power delivered to the load, $P_{dc} = I_{dc}^2 R_L = (I_{max}/\pi)^2 R_L$

AC power input to the transformer, $P_{ac} = \text{Power dissipated in junction of diode} + \text{Power dissipated in load resistance } R_L$

$$= I_{rms}^2 R_F + I_{rms}^2 R_L = \{I_{MAX}^2/4\}[R_F + R_L]$$

Rectification Efficiency, $\eta = P_{dc} / P_{ac} = \{4/\pi^2\}[R_L / (R_F + R_L)] = 0.406 / \{1 + R_F/R_L\}$

If R_F is neglected, the efficiency of half wave rectifier is 40.6%.

2. Ripple factor: It is defined as the amount of AC content in the output DC. It is nothing but amount of AC noise in the output DC. Less the ripple factor, performance of the rectifier is more. The ripple factor of half wave rectifier is about 1.21. It can be calculated as follows:

The effective value of the load current I is given as sum of the rms values of harmonic currents I_1, I_2, I_3, I_4 and DC current I_{dc} .

$$I^2 = I_{dc}^2 + I_1^2 + I_2^2 + I_4^2 = I_{dc}^2 + I_{ac}^2$$

Ripple factor, is given as $\gamma = I_{ac} / I_{dc} = (I^2 - I_{dc}^2) / I_{dc} = \{(I_{rms} / I_{dc})^2 - 1\} = K_f^2 - 1)$

Where K_f is the form factor of the input voltage. Form factor is given as

$$K_f = I_{rms} / I_{avg} = (I_{max}/2) / (I_{max}/\pi) = \pi/2 = 1.57$$

$$\text{So, ripple factor, } \gamma = (1.57^2 - 1) = 1.21$$

3. Peak Inverse Voltage: It is defined as the maximum voltage that a diode can withstand in reverse bias. During the reverse bias as the diode do not conduct total voltage drops across the diode. Thus peak inverse voltage is equal to the input voltage V_s .

4. Transformer Utilization Factor (TUF): The TUF is defined as the ratio of DC power is delivered to the load and the AC rating of the transformer secondary. Half wave rectifier has around 0.287 and full wave rectifier has around 0.693.

$$\begin{aligned} \checkmark \text{ TUF} &= \frac{P_{dc}}{P_{ac}} \\ &= \frac{(\frac{I_m}{2})^2 R_L}{\frac{\pi}{2\sqrt{2}} V_m I_m} \\ &= 0.287 \end{aligned}$$

5. Voltage regulation

The variation of d.c output voltage as function of d.c load current is called regulation.

$$\text{V.R in \% age} = \frac{V_{NL} - V_{FL}}{V_{FL}} * 100$$

Where, V_{NL} =DC voltage across load resistance when minimum current flows through it.

V_{FL} =DC voltage across load resistance when maximum current flows through it.

6. Form factor

It is the ratio of the rms value to the average value.

$$\text{Form factor} = \frac{\text{Rms value}}{\text{Average value}} = \frac{I_m/2}{I_m/\pi} = 1.57$$

7. Output DC Voltage

The output voltage (V_{DC}) across the load resistor is denoted by:

$$V_{DC} = \frac{V_{Smax}}{\pi}, \text{ where } V_{Smax} = \text{maximum amplitude of secondary voltage}$$

8. RMS value of Half Wave Rectifier

To derive the RMS value of half wave rectifier, we need to calculate the current across the load. If the instantaneous load current is equal to $i_L = I_m \sin \omega t$, then the average of load current (I_{DC}) is equal to:

$$I_{dc} = \frac{1}{2\pi} \int_0^{\pi} I_m \sin \omega t = \frac{I_m}{\pi}$$

Where I_m is equal to the peak instantaneous current across the load (I_{max}). Hence the output DC current (I_{DC}) obtained across the load is:

$$I_{DC} = \frac{I_{max}}{\pi}, \text{ where } I_{max} = \text{maximum amplitude of dc current}$$

For a half-wave rectifier, the RMS load current (I_{rms}) is equal to the average current (I_{DC}) multiple by $\pi/2$. Hence the RMS value of the load current (I_{rms}) for a half wave rectifier is:

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

Where $I_m = I_{max}$ which is equal to the peak instantaneous current across the load.

Half wave rectifier is mainly used in the low power circuits. It has very low performance when it is compared with the other rectifiers.

FULL WAVE RECTIFIER:-

Full wave rectifier rectifies the full cycle in the waveform i.e. it rectifies both the positive and negative cycles in the waveform. This Full wave rectifier has an advantage over the half wave i.e. it has average output higher than that of half wave rectifier. The number of AC components in the output is less than that of the input.

The full wave rectifier can be further divided mainly into following types.

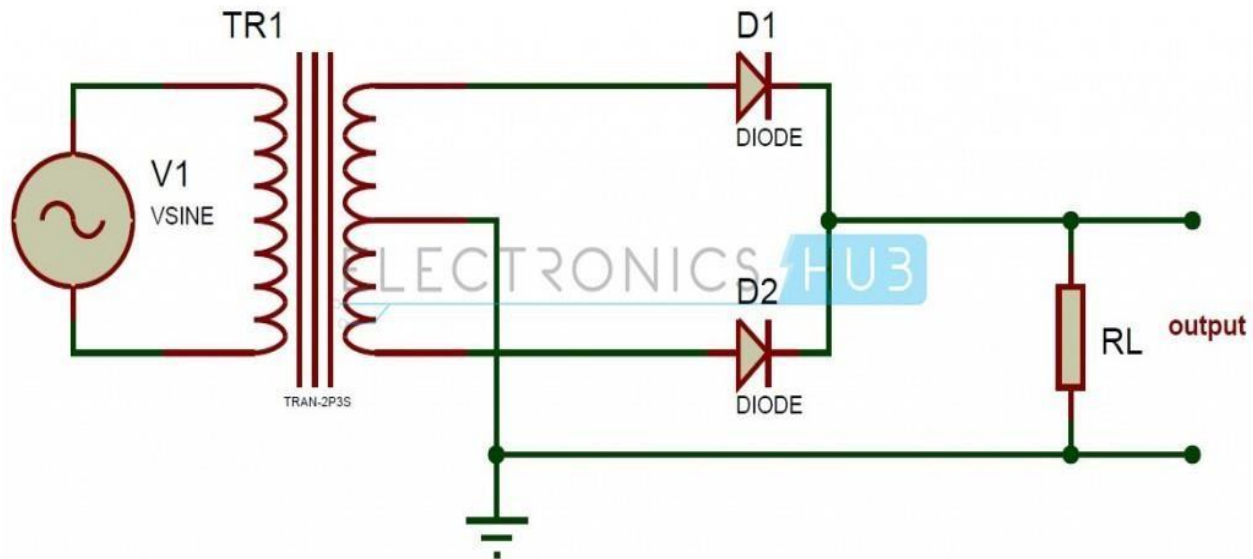
1. Center Tapped Full Wave Rectifier
2. Full Wave Bridge Rectifier

Center Tapped Full Wave Rectifier

Center tap is the contact made at the middle of the winding of the transformer.

In the center tapped full wave rectifier two diodes were used. These are connected to the center tapped secondary winding of the transformer. Above circuit diagram shows the center tapped full wave rectifier. It has two diodes. The positive terminal of two diodes is connected to the two ends of the transformer. Center tap divides the total secondary voltage into equal parts.

Center Tapped Full Wave Rectifier Working:



The primary winding of the center tap transformer is applied with the AC voltage. Thus the two diodes connected to the secondary of the transformer conduct alternatively. For the positive half cycle of the input diode D1 is connected to the positive terminal and D2 is connected to the negative terminal. Thus diode D1 is in forward bias and the diode D2 is reverse biased. Only diode D1 starts conducting and thus current flows from diode and it appears across the load R_L . So positive cycle of the input appears at the load.

During the negative half cycle the diode D2 is applied with the positive cycle. D2 starts conducting as it is in forward bias. The diode D1 is in reverse bias and this does not conduct. Thus current flows from diode D2 and hence negative cycle is also rectified, it appears at the load resistor R_L .

By comparing the current flow through load resistance in the positive and negative half cycles, it can be concluded that the direction of the current flow is same. Thus the frequency of rectified output voltage is two times the input frequency. The output that is rectified is not pure, it consists of a DC component and a lot of AC components of very low amplitudes.

Peak Inverse Voltage (PIV) of Centre Tap Full Wave Rectifier:

PIV is defined as the maximum possible voltage across a diode during its reverse bias. During the first half that is positive half of the input, the diode D1 is forward bias and thus conducts providing no resistance at all. Thus, the total voltage V_s appears in the upper-half of the ac supply, provided to the load resistance R . Similarly, in the case of diode D2 for the lower half of the transformer total secondary voltage developed appears at the load. The amount of voltage that drops across the two diodes in reverse bias is given as

$$\text{PIV of D2} = V_m + V_m = 2V_m$$

$$\text{PIV of D1} = 2V_m$$

V_m is the voltage developed across upper and lower halves.

Peak Current

The peak current is the instantaneous value of the voltage applied to the rectifier. It can be written as

$$V_s = V_m \sin \omega t$$

Let us assume that the diode has a forward resistance of R_F ohms and a reverse resistance is equal to infinity, thus current flowing through the load resistance R_L is given as

$$I_m = V_m / (R_F + R_L)$$

Transformer Utilization Factor:

This can be calculated by considering primary and secondary windings separately. Its value is 0.693. This can be used to determine transformer secondary rating.

Output Current:

Since the current is same through the load resistance R_L in the two halves of the ac cycle, magnitude of dc current I_{dc} , which is equal to the average value of ac current,

can be obtained by integrating the current I_1 between 0 and π or current I_2 between π and 2π .

$$I_{dc} = \frac{1}{\pi} \int_0^{\pi} I_1(\omega t) d(\omega t) = \frac{1}{\pi} \int_0^{\pi} I_{max} \sin \omega t d(\omega t) = \frac{2I_m}{\pi}$$

DC output voltage:

Average value or dc value of voltage across the load is given by

$$I_{dc} = \frac{1}{\pi} \int_0^{\pi} I_1(\omega t) d(\omega t) = \frac{1}{\pi} \int_0^{\pi} I_{max} \sin \omega t d(\omega t) = \frac{2I_m}{\pi}$$

Root Mean Square (RMS) value of current:

RMS value of current flowing through the load resistance is given as

$$I_{RMS}^2 = \frac{1}{\pi} \int_0^{\pi} I_1^2(\omega t) d(\omega t) = \frac{I_m^2}{2}$$

Or

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

Root Mean Square (RMS) Value of output voltage:

RMS value of voltage across the load is given by:

$$V_{load_{rms}} = I_{rms} * R_{load}$$

Rectification efficiency:

$$P_{dc} = I_{dc}^2 R_L = \left(\frac{2I_m}{\pi}\right)^2 R_L$$

As power input to the transformer = power dissipated at the diode + power dissipated at the in load resistance R_L .

$$\text{Rectification efficiency } \eta = \frac{P_{dc}}{P_{ac}} = \frac{\left(\frac{2I_m}{\pi}\right)^2 R_L}{\left\{\frac{I_m^2}{2}\right\} [R_F + R_{load}]}$$

Ripple factor:

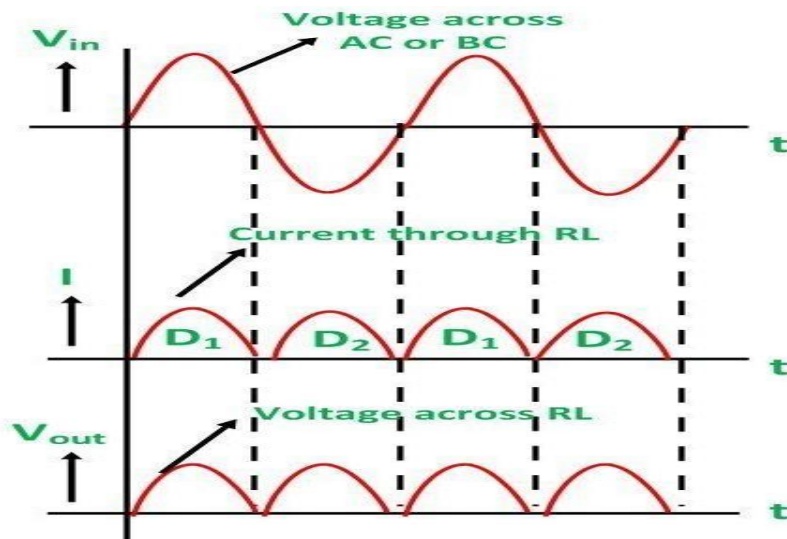
Form factor K_f of the rectified output voltage of a full wave rectifier is given as

$$K_f = I_{rms}/I_{avg} = \frac{I_m \sqrt{2}}{(2I_m/\pi)} = 1.11$$

Regulation:

The dc output voltage is given by

$$\begin{aligned} V_{dc} &= I_{dc} R_L = 2/(\pi I_m R_L) \\ &= 2V_{sm} R_L / (R_F + R_L) \\ &= (2V_{sm}/\pi) - (I_{dc} R_F) \end{aligned}$$



Advantages:-

- ✓ Output is obtained for both cycles of input ac voltages.
- ✓ Efficiency is higher than that of half wave rectifier.

Disadvantages:-

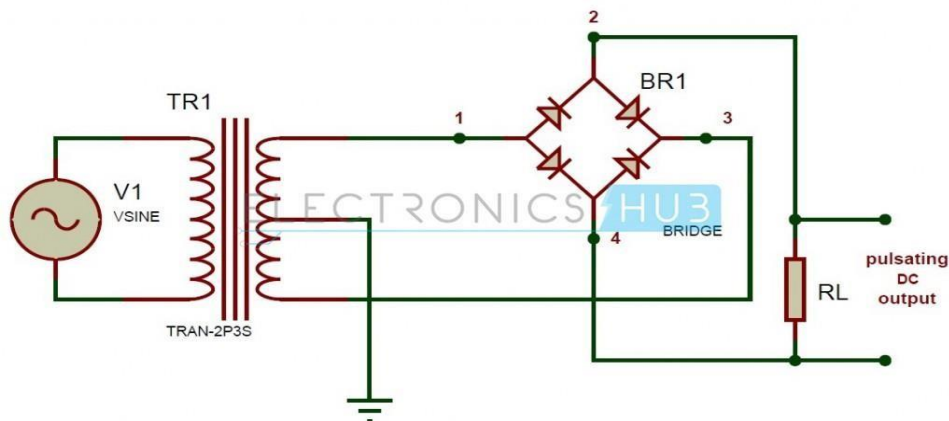
- ✓ Locating center tap on the secondary winding is difficult.
- ✓ The diodes used have high PIV.
- ✓ The d.c output is small as each diode utilizes only one half of the transformer secondary voltage.

FULL WAVE BRIDGE RECTIFIER

Bridge is a type of electrical circuit. Bridge rectifier is a type of rectifier in which diodes were arranged in the form of a bridge. This provides full wave rectification and is of low cost. So it is used in many applications.

Bridge Rectifier:

In bridge rectifier four diodes are used. These are connected as shown in the circuit diagram. The four diodes are connected in the form of a bridge to the transformer and the load as shown.



Working of Bridge Rectifier:

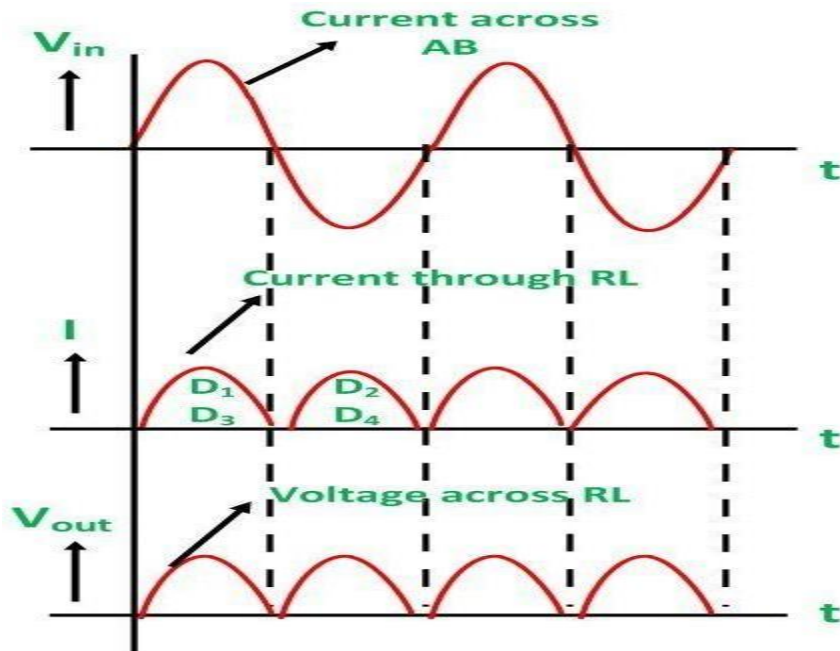
The working of a bridge rectifier is simple. The circuit diagram of bridge rectifier is given above. The secondary winding of the transformer is connected to the two diametrically opposite points of the bridge at points 1 and 3. Assume that a load is connected at the output. The load R_{Load} is connected to bridge through points 2 and 4.

During first half cycle of the AC input, the upper portion of the transformer secondary winding is positive with respect to the lower portion. Thus during the first half cycle diodes D_1 and D_4 are forward biased. Current flows through path 1-2, enter into the load R_L . It returns back flowing through path 4-3. During this half input cycle, the diodes D_2 and D_3 are reverse biased. Hence there is no current flow through the path 2-3 and 1-4.

During the next cycle lower portion of the transformer is positive with respect to the upper portion. Hence during this cycle diodes D_2 and D_3 are forward biased. Current flows through the path 3-2 and flows back through the path 4-1. The diodes

D1 and D4 are reverse biased. So there is no current flow through the path 1-2 and 3-4. Thus negative cycle is rectified and it appears across the load.

Peak Inverse Voltage (PIV) of a bridge rectifier = Maximum of Secondary Voltage



Advantages:-

- ✓ PIV is one half that of centre tap circuit.
- ✓ Output is twice that of centre tap circuit.
- ✓ Need for centre tapped transformer is eliminated.

Disadvantages:-

- ✓ Requires 4 diodes which increase the cost.

Bridge Rectifier Applications:

- Because of their low cost compared to center tapped they are widely used in power supply circuit.
- This can be used to detect the amplitude of modulated radio signal.
- Bridge rectifiers can be used to supply polarized voltage in welding.

TRANSISTORS

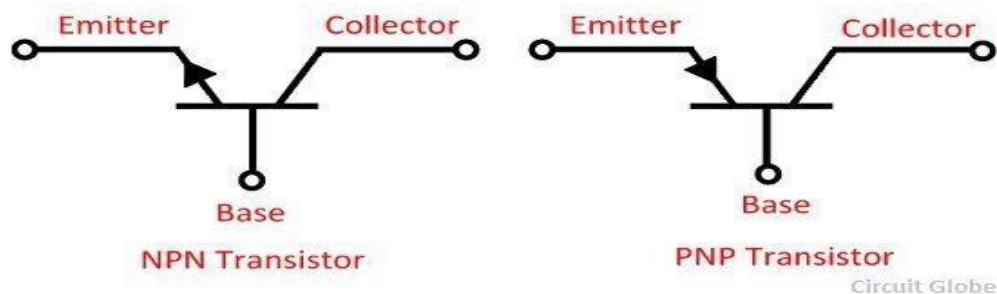
PRINCIPLE OF BIPOLAR JUNCTION TRANSISTOR:

The transistor is a semiconductor device which transfers a weak signal from low resistance circuit to high resistance circuit. The words trans mean transfer property and istor mean resistance property offered to the junctions. In other words, it is a switching device which regulates and amplifies the electrical signal like voltage or current.

The transistor consists of two PN diodes connected back to back. It has three terminals namely emitter, base and collector. The base is the middle section which is made up of thin layers. The right part of the diode is called emitter diode and the left part is called collector-base diode. These names are given as per the common terminal of the transistor. The emitter-base junction of the transistor is connected to forward bias and the collector-base junction is connected in reverse bias which offers a high resistance.

Transistor Symbols

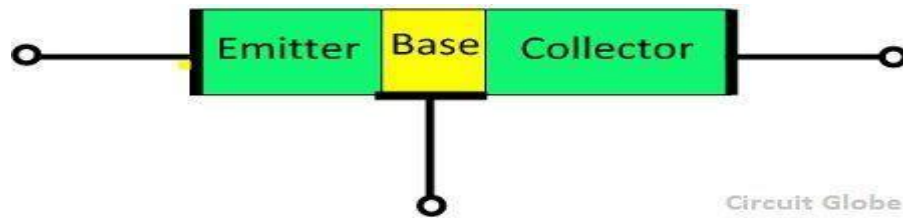
There are two types of transistor, namely NPN transistor and PNP transistor. The transistor which has two blocks of n-type semiconductor material and one block of P-type semiconductor material is known as NPN transistor. Similarly, if the material has one layer of N-type material and two layers of P-type material then it is called PNP transistor. The symbol of NPN and PNP is shown in the figure below.



The arrow in the symbol indicates the direction of flow of conventional current in the emitter with forward biasing applied to the emitter-base junction. The only difference between the NPN and PNP transistor is in the direction of the current.

Transistor Terminals

The transistor has three terminals namely, emitter, collector and base. The terminals of the diode are explained below in details.



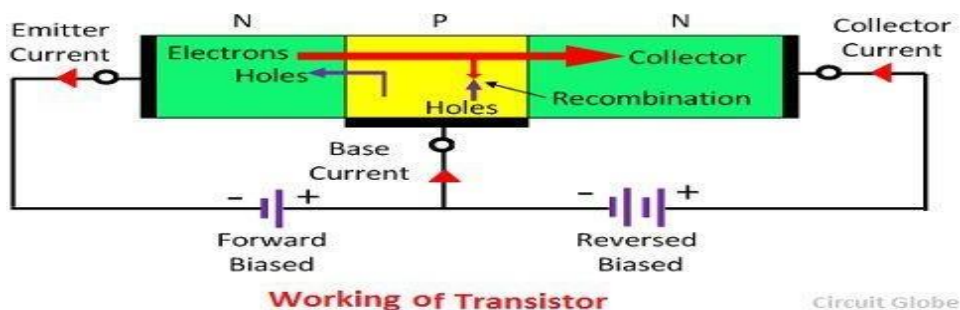
Emitter – The section that supplies the large section of majority charge carrier is called emitter. The emitter is always connected in forward biased with respect to the base so that it supplies the majority charge carrier to the base. The emitter-base junction injects a large amount of majority charge carrier into the base because it is heavily doped and moderate in size.

Collector – The section which collects the major portion of the majority charge carrier supplied by the emitter is called a collector. The collector-base junction is always in reverse bias. Its main function is to remove the majority charges from its junction with the base. The collector section of the transistor is moderately doped, but larger in size so that it can collect most of the charge carrier supplied by the emitter.

Base – The middle section of the transistor is known as the base. The base forms two circuits, the input circuit with the emitter and the output circuit with the collector. The emitter-base circuit is in forward biased and offered the low resistance to the circuit. The collector-base junction is in reverse bias and offers the higher resistance to the circuit. The base of the transistor is lightly doped and very thin due to which it offers the majority charge carrier to the base.

Working of Transistor

Usually, silicon is used for making the transistor because of their high voltage rating, greater current and less temperature sensitivity. The emitter-base section kept in forward biased constitutes the base current which flows through the base region. The magnitude of the base current is very small. The base current causes the electrons to move into the collector region or create a hole in the base region.



The base of the transistor is very thin and lightly doped because of which it has less number of electrons as compared to the emitter. The few electrons of the emitter are combined with the hole of the base region and the remaining electrons are moved towards the collector region and constitute the collector current. Thus we can say that the large collector current is obtained by varying the base region.

DIFFERENT MODES OF OPERATION OF TRANSISTOR:

When the emitter junction is in forward biased and the collector junction is in reverse bias, then it is said to be in the active region. The transistor has two junctions which can be biased in different ways. The different working conduction of the transistor is shown in the table below.

CONDITION	EMITTER JUNCTION (EB)	COLLECTOR JUNCTION (CB)	REGION OF OPERATION
FR	Forward-biased	Reversed-biased	Active
FF	Forward-biased	Forward-biased	Saturation
RR	Reversed-biased	Reversed-biased	Cut-off
RF	Reversed-biased	Forward-biased	Inverted

FR – In this case, the emitter-base junction is connected in forward biased and the collector-base junction is connected in reverse biased. The transistor is in the active region and the collector current depends on the emitter current. The transistor, which operates in this region, is used for amplification.

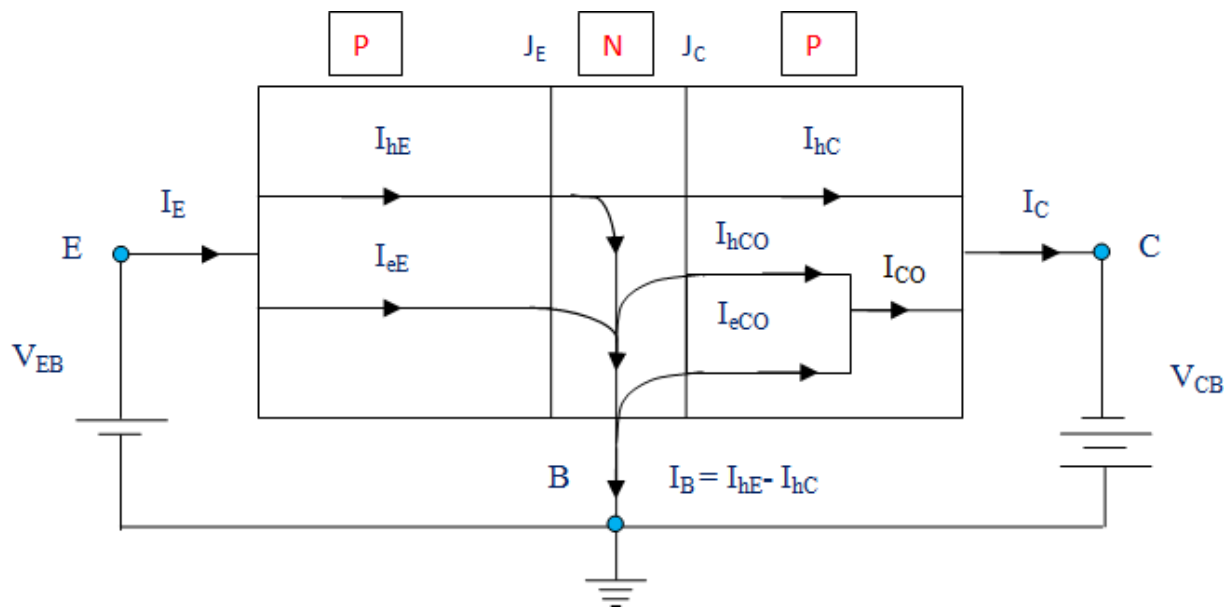
FF – In this condition, both the junction is in forward biased. The transistor is in saturation and the collector current becomes independent of the base current. The transistors act like a closed switch.

RR – Both the current are in reverse biased. The emitter does not supply the majority charge carrier to the base and carriers current are not collected by the collector. Thus the transistors act like a closed switch.

RF – The emitter-base junction is in reverse bias and the collector-base junction is kept in forward biased. As the collector is lightly doped as compared to the emitter junction it does not supply the majority charge carrier to the base. Hence poor transistor action is achieved.

CURRENT COMPONENTS OF IN A TRANSISTOR:

The conduction of current in NPN transistor is owing to electrons and in PNP transistor, it is owing to holes. The direction of current flow will be in opposite direction. Here, we can discuss the current components in a PNP transistor with common base configuration. The emitter-base junction (J_E) is forward biased and the collector-base junction (J_C) is reversed biased as shown in figure. All the current components related to this transistor are shown here.



The current arrives the transistor through the emitter and this current is called emitter current (I_E). This current consists of two constituents – **Hole current** (I_{hE}) and **Electron current** (I_{eE}). I_{eE} is due to passage of electrons from base to emitter and I_{hE} is due to passage of holes from emitter to base.

$$I_E = I_{hE} + I_{eE}$$

Normally, the emitter is heavily doped compared to base in industrial transistor. So, the Electron current is negligible compared to Hole current. Thus we can conclude that, the whole emitter current in this transistor is due to the passage of holes from the emitter to the base.

Some of the holes which are crossing the junction J_E (emitter junction) combines with the electrons present in the base (N-type). Thus, every holes crossing J_E will not arrive at J_C . The remaining holes will reach the collector junction which produces the hole current component, I_{hC} . There will be bulk recombination in the base and the current leaving the base will be

$$I_B = I_{hE} - I_{hC}$$

The electrons in the base which are lost by the recombination with holes (injected into the base across J_E) are refilled by the electrons that enter into the base region. The holes which are arriving at the collector junction (J_C) will cross the junction and it will go into the collector region.

When the emitter circuit is open circuited, then $I_E = 0$ and $I_{hC} = 0$. In this condition, the base and collector will perform as reverse biased diode. Here, the collector current, I_C will be same as reverse saturation current (I_{CO} or I_{CBO}).

I_{CO} is in fact a small reverse current which passes through the PN junction diode. This is due to thermally generated minority carriers which are pushed by barrier potential. This reverse current increase; if the junction is reverse biased and it will have the same direction as the collector current. This current attains a saturation value (I_0) at moderate reverse biased voltage.

When the emitter junction is at forward biased (in active operation region), then the collector current will become

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

The α is the large signal current gain which is a fraction of the emitter current which comprises of I_{hC} .

When the emitter is at closed condition, then $I_E \neq 0$ and collector current will be

$$I_C = I_{CO} + I_{hC}$$

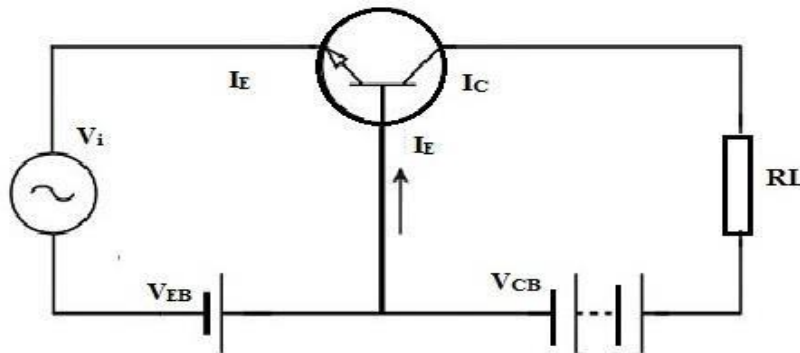
In a PNP transistor, the reverse saturation current (I_{CBO}) will comprises of the current due to the holes passing through the collector junction from the base to collector region (I_{hCO}) and the current due to the electrons which are passing through the collector junction in the opposite direction (I_{eCO}).

$$\text{Therefore, } I_{CO} = I_{hCO} + I_{eCO}$$

The total current entering into the transistor will be equal to the total current leaving the transistor (according to Kirchhoff's current law).

$$\text{So, } I_E = I_C + I_B \text{ or } I_E = -(I_C + I_B)$$

TRANSISTOR AS AN AMPLIFIER:-



A transistor can be used as **an amplifier** by enhancing the weak signal's strength. With the help of the following transistor amplifier circuit, one can get an idea about how the transistor circuit works as an amplifier circuit.

In the below circuit, the input signal can be applied across the emitter-base junction and the output across the R_c load connected in the collector circuit.

For accurate amplification, always remember that the input is connected in forward-biased whereas the output is connected in reverse-biased. For this reason, in addition to the signal, we apply DC voltage (V_{EE}) in the input circuit as shown in the above circuit.

Generally, the input circuit includes low resistance as a result; a little change will occur in signal voltage at the input which leads to a significant change within the emitter current. Because of the transistor action, emitter current change will cause the same change within the collector circuit.

At present, the flow of collector current through an R_c generates a huge voltage across it. Therefore, the applied weak signal at the input circuit will come out in the amplified form at the collector circuit in the output. In this method, the transistor performs as an amplifier.

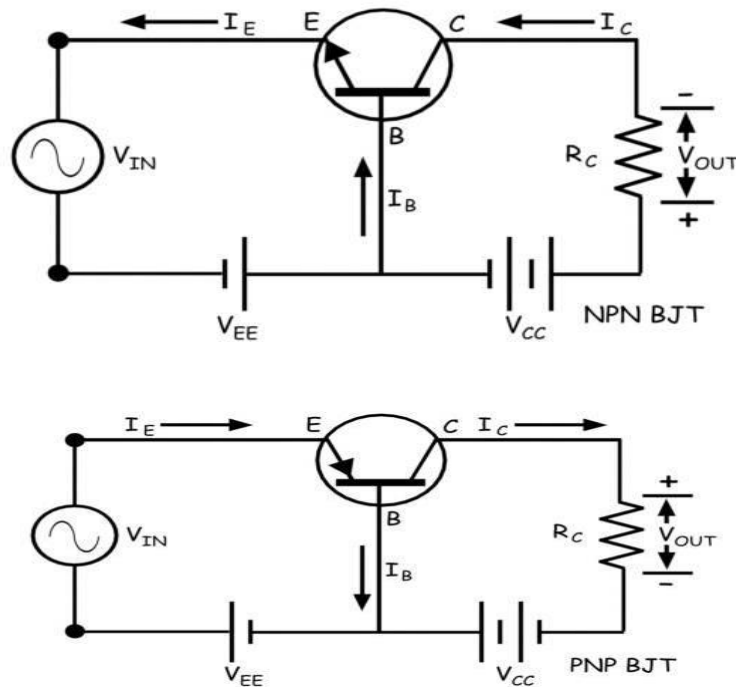
TRANSISTOR CIRCUIT CONFIGURATION & ITS CHARACTERISTICS:

A transistor can be connected in a circuit in the following three ways:

- a) common base connection
- b) common emitter connection
- c) common collector connection

COMMON BASE CONNECTION

Here the base terminal is common to both input and output circuit. The common base configurations or modes are as shown in the figure below. Here, the common base mode of NPN transistor and PNP transistor are shown separately.



Here emitter-base circuit is taken as input circuit and collector base circuit as output circuit.

Current Gain

Here the input current is emitter current I_E and output current is collector current I_C . The current gain is considered as when we only consider the dc biasing voltages of the circuit and no alternating signal is applied in the input.

$$\alpha_{(dc)} = \frac{I_C}{I_E}$$

Now if we consider the alternating signal applied to the input then the current amplification factor (α) at a constant collector-base voltage, would be

$$\alpha = \frac{\Delta I_C}{\Delta I_E}$$

Here it is seen that neither of current gain and current amplification factor has value more than unity since collector current in no way can be more than emitter current. But as we know that the emitter current and collector current are nearly equal in a bipolar junction transistor, these ratios would be very near to unity. The value generally ranges from 0.9 to even 0.99.

Expression of Collector Current

- If the emitter circuit is open, there will be no emitter current ($I_C = 0$). But in this condition, there will be a tiny current flowing through the collector region. This is because of flow of minority charge carriers and this is the reverse leakage current.
- As this current flow through collector and base keeping the emitter terminal open, the current is denoted as I_{CBO} . In small power rated transistor the reverse leakage current I_{CBO} is quite small and generally, we neglect it during calculations but in high power rated transistor this leakage current cannot be neglected.
- This current is highly dependent on the temperature so at high temperatures the reverse leakage current I_{CBO} cannot be neglected during calculations.

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

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

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

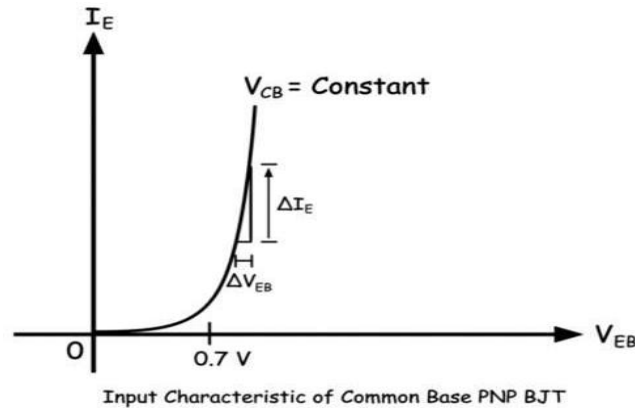
$$\Rightarrow I_C = \frac{\alpha}{1 - \alpha} I_B + \frac{1}{1 - \alpha} I_{CBO}$$

This expression proves that collector current also depends on base current.

Characteristic of Common Base Connection

Input Characteristic

This is drawn between input current and input voltage of the transistor itself. The input current is emitter current (I_E) and the input voltage is emitter-base voltage (V_{EB}). After crossing emitter-base junction forward barrier potential emitter current (I_E) starts increasing rapidly with increasing emitter-base voltage (V_{EB}).

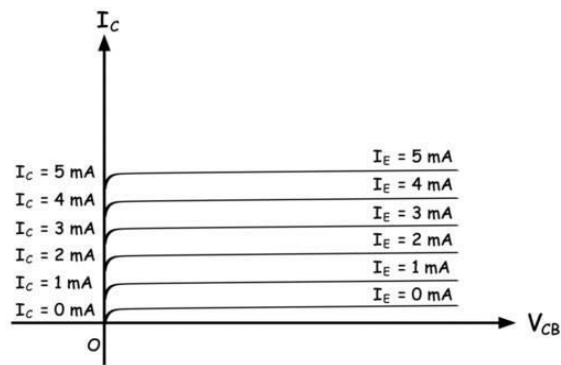


The input resistance of the circuit is the ratio of change in emitter-base voltage (ΔV_{EB}) to emitter current (ΔI_E) at a constant collector-base voltage ($V_{CB} = \text{Constant}$). As the change in emitter current is quite large compared to the change in emitter-base voltage ($\Delta I_E \gg \Delta V_{EB}$), the input resistance of the common base transistor is quite small.

$$r_o = \frac{\Delta V_{CB}}{\Delta I_C} \text{ When, } I_E = \text{Constant}$$

Output Characteristic

- Collector current gets only constant value when there is sufficient reverse biased established between base and collector region. This is why there is a rise of collector current with an increase of collector-base voltage when this voltage has very low value.
- But after a certain collector-base voltage the collector-base junction gets sufficient reverse biased and hence the collector current becomes constant for a certain emitter current and it entirely depends on the emitter current.

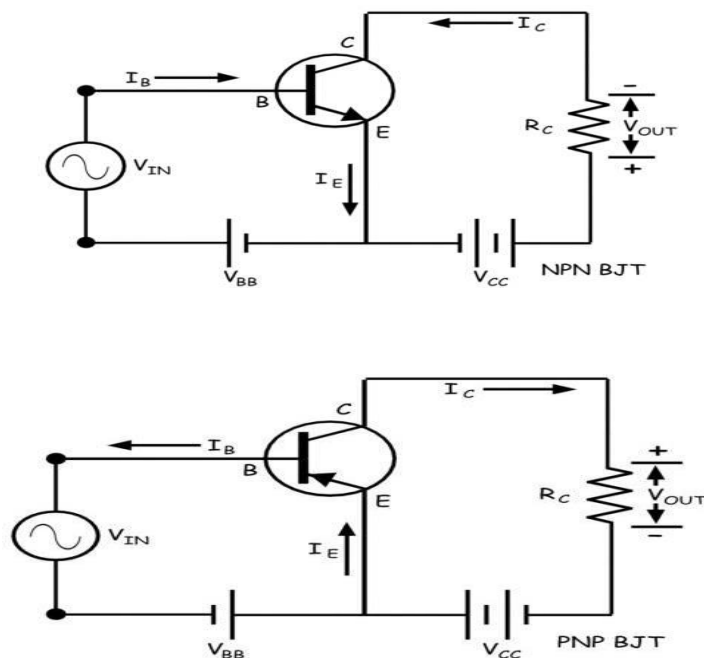


The ratio of change in collector-base voltage to the change in collector current is defined as the output resistance of common base mode of a transistor. Naturally, the value of output resistance is very high in the common base mode of a transistor.

$$r_o = \frac{\Delta V_{CB}}{\Delta I_C} \text{ When, } I_E = \text{Constant}$$

COMMON EMITTER CONNECTION

Common Emitter Transistor is the most commonly used transistor connection. Here the emitter terminal is common for both input and output circuit. The circuit connected between base and emitter is the input circuit and the circuit connected between collector and emitter is the output circuit. The common emitter mode of NPN transistor and PNP transistor are shown separately in the figure below.



Current Gain

In common emitter configuration, the input current is base current (I_B) and the output current is collector current (I_C). In bipolar junction transistor, the base current controls the collector current. The ratio of change in collector current (ΔI_C) to change in base current (ΔI_B) is defined as the current gain of common emitter transistor. In a bipolar junction transistor, the emitter current (I_E) is the sum of the base current (I_B) and collector current (I_C).

If base current changes, the collector current also changes and as a result the emitter current gets also changed accordingly.

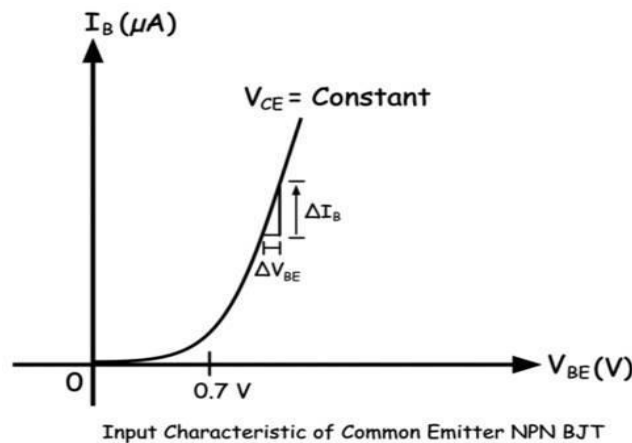
$$\beta = \frac{\Delta I_C}{\Delta I_B}$$

Again the ratio of change of collector current to the corresponding change in emitter current is denoted by α

As the value of base current is quite low compared to the collector current ($I_B \ll I_C$), the current gain in a common emitter transistor is quite high and it ranges from 20 to 500.

Characteristic of Common Emitter Transistor

- In common emitter mode of the transistor, there are two circuits – input circuit and the output circuit. In the input circuit, the parameters are base current and base-emitter voltage.
- The characteristic curve drawn against variations of base current and base-emitter voltage is input characteristic of a common emitter transistor.

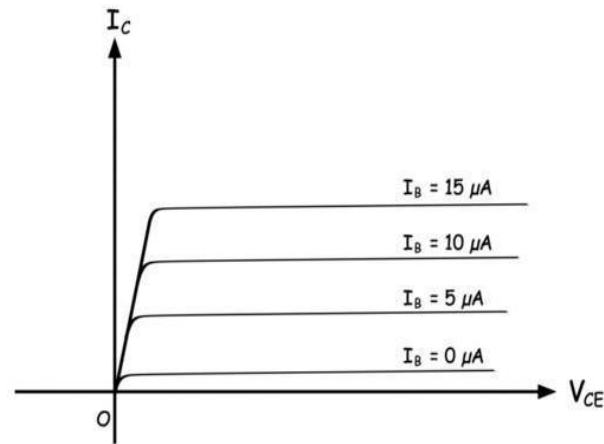


Input resistance of the circuit is:

$$r_i = \frac{\Delta V_{BE}}{\Delta I_B} \text{ When, } V_{CE} \text{ is constant}$$

Output Characteristic of Common Emitter Transistor

- The output characteristic is drawn against variations of output current and the output voltage of the transistor. The collector current is output current and collector-emitter voltage is the output voltage of the transistor.

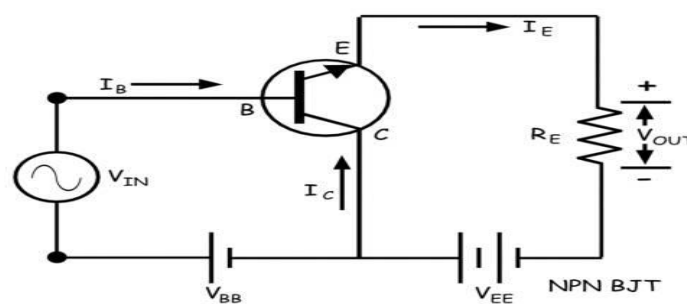


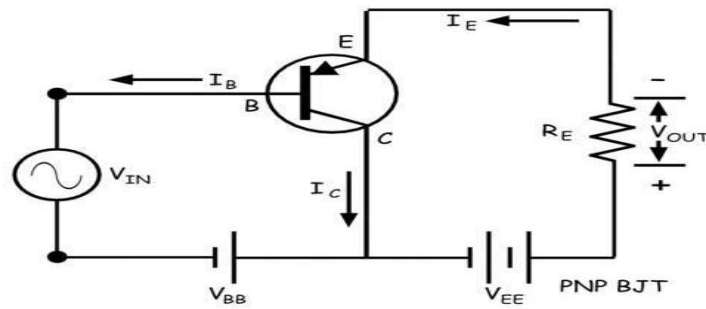
The output resistance would be

$$r_o = \frac{\Delta V_{CE}}{\Delta I_C} \text{ When, } I_B \text{ is constant}$$

COMMON COLLECTOR:

In common collector configuration the input circuit is between base and collector terminal and the output circuit is between emitter and collector terminal.





The ratio of change of emitter current to change of base current is defined as the current gain of common collector configuration. This is denoted as,

$$\frac{I_E}{I_B}$$

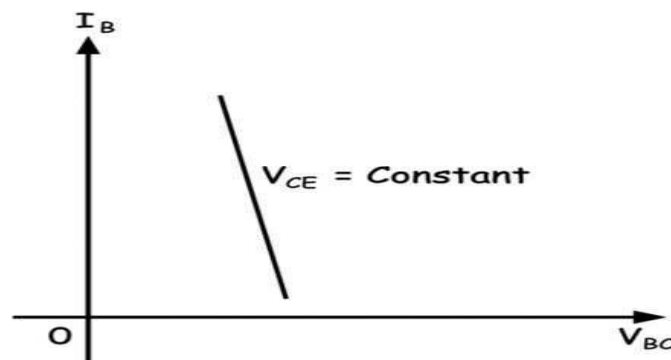
The current amplification factor of the circuit is the ratio of change of emitter current to change of base current when a time-varying signal is applied to the input.

$$\gamma = \frac{\Delta I_E}{\Delta I_B}$$

Input Characteristic of Common Collector Transistor

The input current is base current and input voltage of the transistor is base-collector voltage. The base-collector junction is reverse biased and hence with increasing base-collector voltage the reverse biasing of the junction increases. This causes base current to decrease slightly with the increase in base-collector voltage.

Input Characteristic of Common Collector Transistor

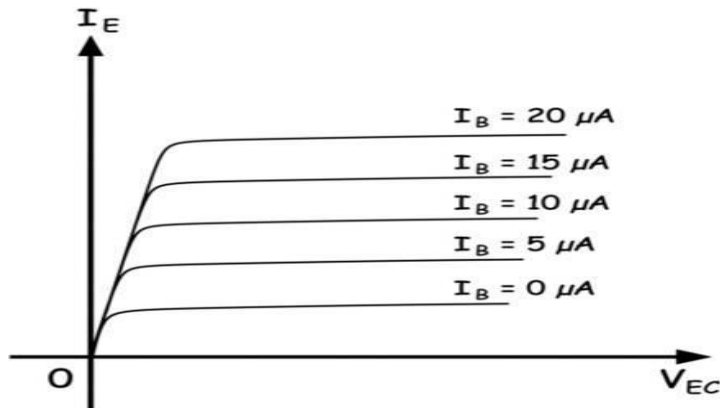


Output Characteristic of Common Collector Transistor

The output characteristic of a common collector transistor is nearly the same as the output characteristic of a common emitter transistor. The only difference that here

in the case of common collector configuration the output current is emitter current instead of collector current as in the case of common emitter configuration. Here also for a fixed base current, the emitter current increases linearly with increasing collector-emitter voltage up to a certain level of this voltage and then the emitter current gets almost constant irrespective of collector-emitter voltage. Although there would be a very slow increase of emitter current with the collector-emitter voltage as shown in the characteristic curve below.

Output Characteristic of Common Collector Transistor



TRANSISTOR BIASING:-

- ✓ Biasing is the process of providing DC voltage which helps in 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.
- ✓ The proper flow of zero signal collector current and the maintenance of proper collector emitter voltage during the passage of a signal is known as transistor biasing.
- ✓ The circuit which provides transistor biasing is known as biasing circuit.
- ✓ Transistor biasing is basically classified into 4 types:

(a) Fixed biasing

(b) Emitter stabilized biasing

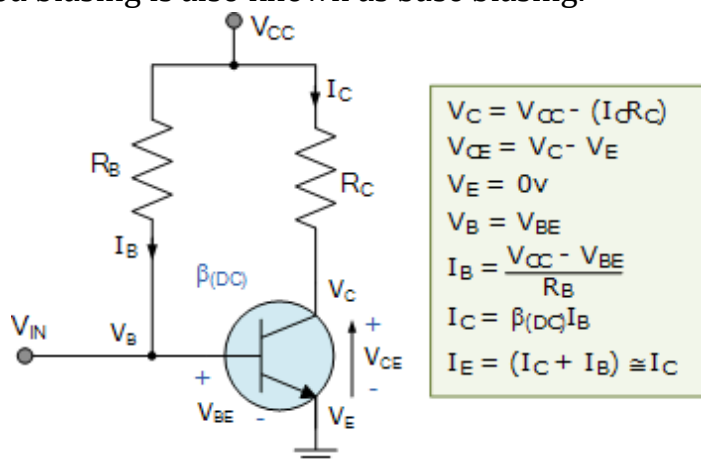
(c) Voltage divider biasing

(d) DC biasing with voltage feedback

DIFFERENT METHOD OF TRANSISTOR BIASING:

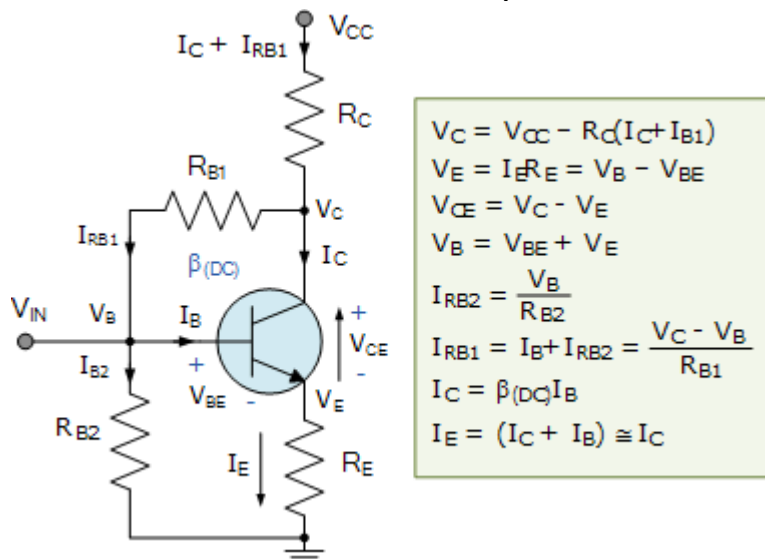
FIXED BIASING:

- ✓ Fixed biasing is also known as base biasing.



- ✓ The above circuit shown is called as a “fixed base bias circuit”, because the transistors base current, I_B remains constant for given values of V_{CC} , and therefore the transistors operating point must also remain fixed.
- ✓ This two resistor biasing network is used to establish the initial operating region of the transistor using a fixed current bias.
- ✓ This type of transistor biasing arrangement is also beta dependent biasing as the steady-state condition of operation is a function of the transistors beta β value, so the biasing point will vary over a wide range for transistors of the same type as the characteristics of the transistors will not be exactly the same.
- ✓ The emitter diode of the transistor is forward biased by applying the required positive base bias voltage via the current limiting resistor R_B .
- ✓ Assuming a standard bipolar transistor, the forward base-emitter voltage drop would be 0.7V. Then the value of R_B is simply: $(V_{CC} - V_{BE})/I_B$ where I_B is defined as I_C/β .
- ✓ With this single resistor type of biasing arrangement the biasing voltages and currents do not remain stable during transistor operation and can vary enormously.
- ✓ Also the operating temperature of the transistor can adversely affect the operating point.

EMITTER STABILIZED BIASING/COLLECTOR FEEDBACK BIASING

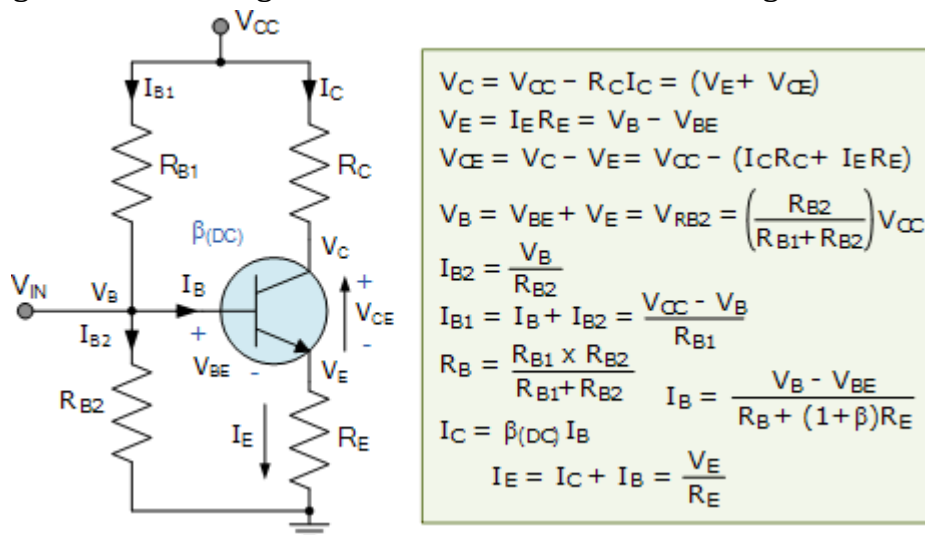


- ✓ This type of transistor biasing configuration, often called self-emitter biasing, uses both emitter and base-collector feedback to stabilize the collector current even further.
- ✓ This is because resistors R_{B1} and R_E as well as the base-emitter junction of the transistor are all effectively connected in series with the supply voltage, V_{CC} .

- ✓ The downside of this emitter feedback configuration is that it reduces the output gain due to the base resistor connection.
- ✓ The collector voltage determines the current flowing through the feedback resistor, R_{B1} producing what is called “degenerative feedback”.
- ✓ The current flowing from the emitter, I_E (which is a combination of $I_C + I_B$) causes a voltage drop to appear across R_E in such a direction, that it reverse biases the base-emitter junction.
- ✓ So if the emitter current increases, due to an increase in collector current, voltage drop $I \cdot R_E$ also increases. Since the polarity of this voltage reverse biases the base-emitter junction, I_B automatically decrease. Therefore the emitter current increase less than it would have done had there been no self-biasing resistor.
- ✓ Generally, resistor values are set so that the voltage dropped across the emitter resistor R_E is approximately 10% of V_{CC} and the current flowing through resistor R_{B1} is 10% of the collector current I_C .
- ✓ Thus this type of transistor biasing configuration works best at relatively low power supply voltages.

VOLTAGE DIVIDER BIASING:

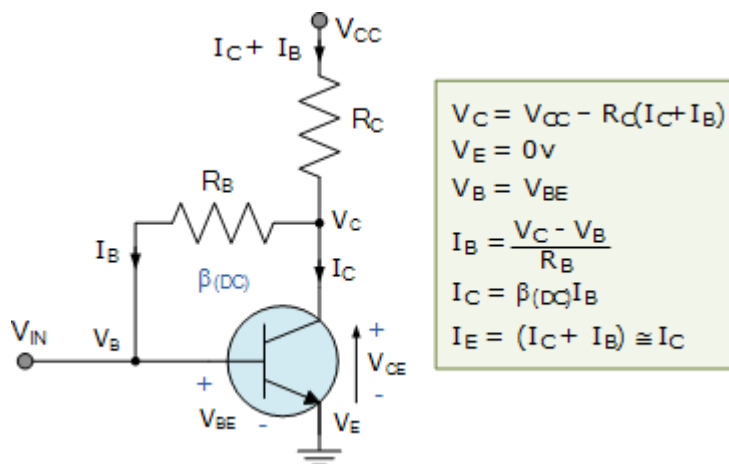
- ✓ Voltage divider biasing is also known as universal biasing.



- ✓ Here the common emitter transistor configuration is biased using a voltage divider network to increase stability.
- ✓ The name of this biasing configuration comes from the fact that the two resistors R_{B1} and R_{B2} form a voltage or potential divider network across the supply with their center point junction connected the transistors base terminal as shown.

- ✓ This voltage divider biasing configuration is the most widely used transistor biasing method.
- ✓ The emitter diode of the transistor is forward biased by the voltage value developed across resistor R_{B2} .
- ✓ The voltage divider network biasing makes the transistor circuit independent of changes in beta as the biasing voltages set at the transistors base, emitter, and collector terminals are not dependant on external circuit values.
- ✓ To calculate the voltage developed across resistor R_{B2} and the voltage applied to the base terminal we simply use the voltage divider formula for resistors in series.
- ✓ Generally the voltage drop across resistor R_{B2} is much less than for resistor R_{B1} . Clearly the transistors base voltage V_B with respect to ground will be equal to the voltage across R_{B2} .
- ✓ The amount of biasing current flowing through resistor R_{B2} is generally set to 10 times the value of the required base current I_B so that it is sufficiently high enough to have no effect on the voltage divider current or changes in Beta.

DC BIAS WITH COLLECTOR FEEDBACK:



This self-biasing collector feedback configuration is another beta dependent biasing method which requires two resistors to provide the necessary DC bias for the transistor. The collector to base feedback configuration ensures that the transistor is always biased in the active region regardless of the value of Beta (β). The DC base bias voltage is derived from the collector voltage V_C , thus providing good stability.

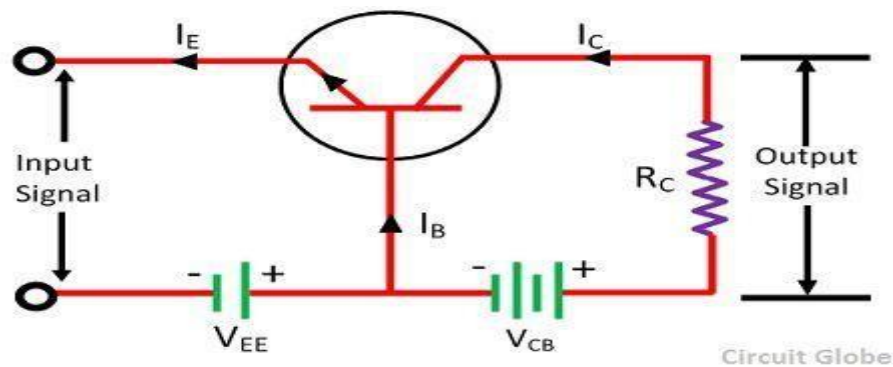
In this circuit, the base bias resistor, R_B is connected to the transistors collector C, instead of to the supply voltage rail, V_{CC} . Now if the collector current increases, the collector voltage drops, reducing the base drive and thereby automatically reducing the collector current to keep the transistors Q-point fixed. Therefore this method of collector feedback biasing produces negative feedback round the transistor as there is a direct feedback from the output terminal to the input terminal via resistor, R_B .

Since the biasing voltage is derived from the voltage drop across the load resistor, R_L , if the load current increases there will be a larger voltage drop across R_L , and a corresponding reduced collector voltage, V_C . This effect will cause a corresponding drop in the base current, I_B which in turn, brings I_C back to normal.

The opposite reaction will also occur when the transistors collector current reduces. Then this method of biasing is called self-biasing with the transistors stability using this type of feedback bias network being generally good for most amplifier designs.

PRACTICAL CIRCUIT OF TRANSISTOR AMPLIFIER

The transistor raises the strength of a weak signal and hence acts as an amplifier. The transistor amplifier circuit is shown in the figure below. The transistor has three terminals namely emitter, base and collector. The emitter and base of the transistor are connected in forward bias and the collector-base region is in reverse bias. The forward bias means the P-region of the transistor is connected to the positive terminal of the supply and the negative region is connected to the N-terminal and in reverse bias just opposite of it has occurred.



The input signal or weak signal is applied across the emitter-base and the output is obtained from the load resistor R_C which is connected in the collector circuit. The DC voltage V_{EE} is applied to the input circuit along with the input signal to achieve the amplification. The DC voltage V_{EE} keeps the emitter-base junction under the forward-biased condition regardless of the polarity of the input signal and is known as a bias voltage.

When a weak signal is applied to the input, a small change in signal voltage causes a change in emitter current (or we can say a change of 0.1V in signal voltage causes a change of 1mA in the emitter current) because the input circuit has very low resistance. This change is almost the same in collector current because of the transistor action.

In the collector circuit, a load resistor R_C of high value is connected. When collector current flows through such a high resistance, it produces a large voltage drop across it. Thus, a weak signal (0.1V) applied to the input circuit appears in the amplified form (10V) in the collector circuit.

TRANSISTOR LOAD LINE ANALYSIS:

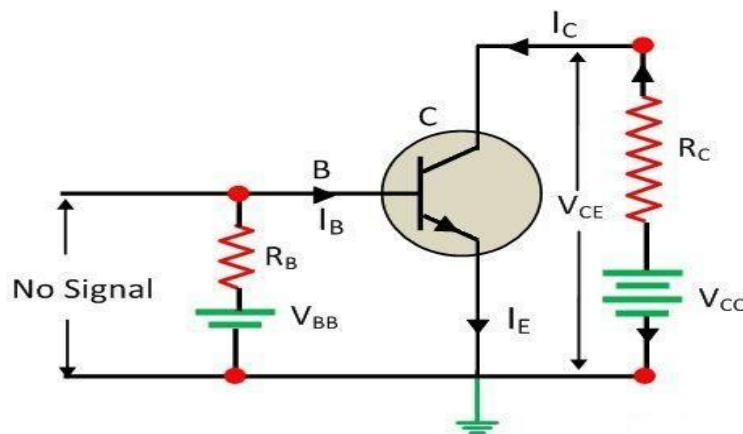
The load line analysis of a transistor means for the given value of collector-emitter voltage we find the value of collector current. This can be done by plotting the output characteristic and then determining the collector current I_C with respect to

collector-emitter voltage V_{CE} . The load line analysis can easily be obtained by determining the output characteristics of the load line analysis methods.

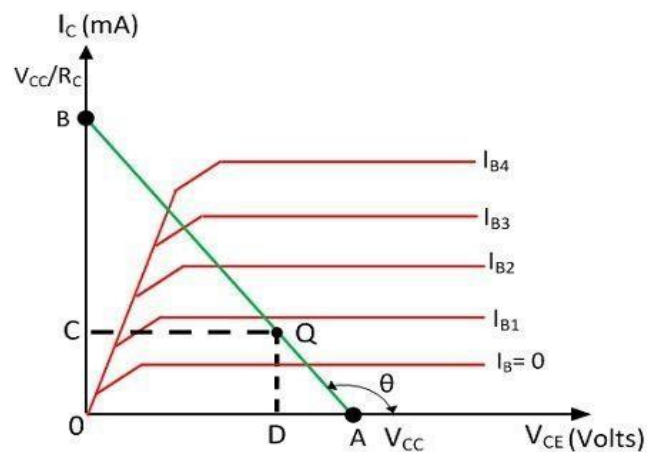
DC LOAD LINE & DC EQUIVALENT CIRCUIT: -

The DC load represents the desirable combinations of the collector current and the collector-emitter voltage. It is drawn when no signal is given to the input, and the transistor becomes bias.

Consider a CE NPN transistor circuit shown in the figure below where no signal is applied to the input side. For this circuit, DC condition will obtain, and the output characteristic of such a circuit is shown in the figure below.



The DC load line curve of the above circuit is shown in the figure below.



By applying Kirchhoff's voltage law to the collector circuit, we get,

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

$$V_{CE} = V_{CC} - I_C R_C \text{ ----- equation 1}$$

The above equation shows that the V_{CC} and R_C are the constant value, and it is the first-degree equation which is represented by the straight line on the output characteristic. This load line is known as a DC load line. The input characteristic is used to determine the locus of V_{CE} and I_C point for the given value of R_C . The end point of the line are located as

1. The collector-emitter voltage V_{CE} is maximum when the collector current $I_C = 0$ then from the equation (1) we get,

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

The first point A ($OA = V_{CC}$) on the collector-emitter voltage axis shown in the figure above.

2. The collector current I_C becomes maximum when the collector-emitter voltage $V_{CE} = 0$ then from the equation (1) we get,

$$0 = V_{CC} - I_C R_C$$

$$I_C = \frac{V_{CC}}{R_C}$$

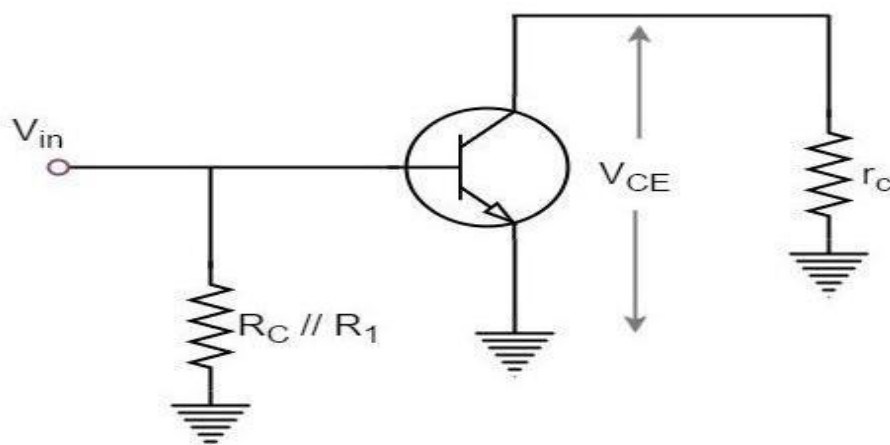
This gives the second point on the collector current axis as shown in the figure above.

By adding the points A and B, the DC load line is drawn. With the help of load line, any value of collector current can be determined.

AC LOAD LINE & AC EQUIVALENT CIRCUIT: -

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,

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

$$r_c = (r_c // R_1)$$

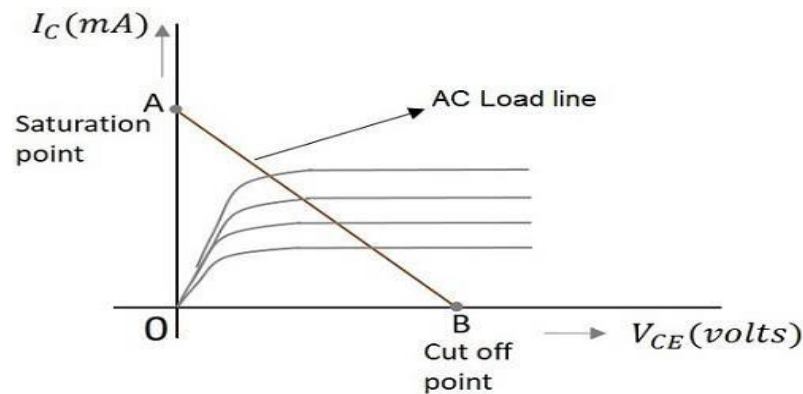
For a transistor to operate as an amplifier, it should stay in active region. The quiescent point is so chosen in such a way that the maximum input signal excursion is symmetrical on both negative and positive half cycles.

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

$$I_{C(sat)} = I_{CQ} + \left(\frac{V_{CEQ}}{r_c} \right)$$

The voltage V_{CE} at the cutoff point is

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

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

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

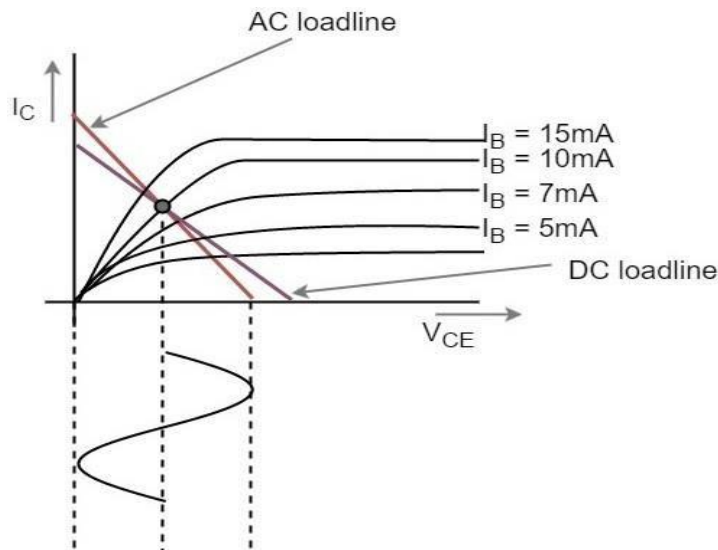
Hence by adding quiescent currents the end points of AC load line are

$$I_{(sat)} = I_{CQ} + V_{CEQ} / (R_c // R_1)$$

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

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.

HYBRID PARAMETERS OR H PARAMETERS

Hybrid parameters (also known as h parameters) are known as 'hybrid' parameters as they use Z parameters, Y parameters, voltage ratio, and current ratios to represent the relationship between voltage and current in a two port network.

H parameters are useful in describing the input-output characteristics of circuits where it is hard to measure Z or Y parameters (such as a transistor). H parameters encapsulate all the important linear characteristics of the circuit, so they are very useful for simulation purposes. The relationship between voltages and current in h parameters can be represented as:

$$V_1 = h_{11}I_1 + h_{12}V_2$$

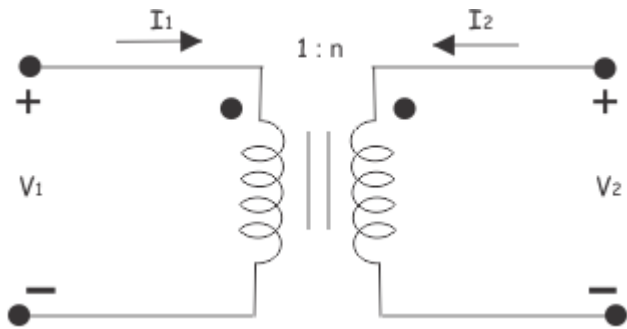
$$I_2 = h_{21}I_1 + h_{22}V_2$$

This can be represented in matrix form as:

$$\begin{bmatrix} V_1 \\ I_2 \end{bmatrix} = \begin{bmatrix} h_{11} & h_{12} \\ h_{21} & h_{22} \end{bmatrix} \begin{bmatrix} I_1 \\ V_2 \end{bmatrix}$$

To illustrate where h parameters are useful, take the case of an ideal transformer, where Z parameters cannot be used. Since here, the relations between voltages and current in that ideal transformer would be,

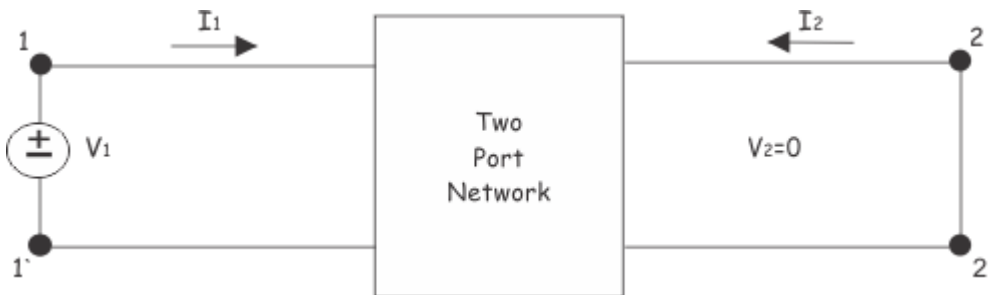
$$V_1 = \frac{1}{n}V_2 \text{ and } I_1 = -nI_2$$



Since, in an ideal transformer voltages cannot be expressed in terms of current, it is impossible to analyze a transformer with Z parameters because a transformer does not have Z parameters. The problem can be solved by using hybrid parameters (i.e. h parameters).

Determining h Parameters

Let us short circuit the output port of a two port network as shown below,



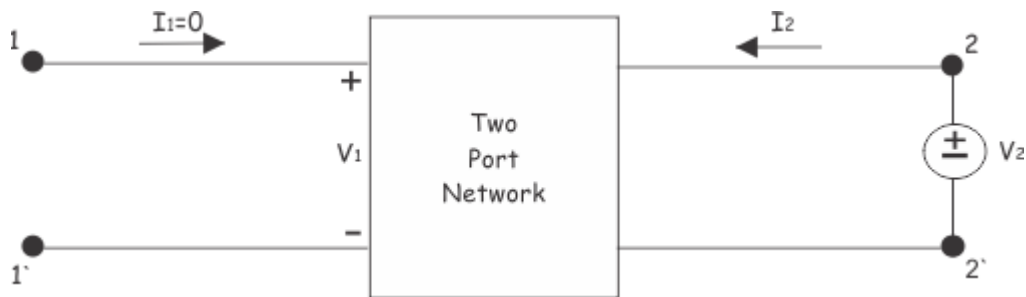
Now, ratio of input voltage to input current, at short circuited output port is:

$$\left. \frac{V_1}{I_1} \right|_{V_2=0} = h_{11}$$

This is referred to as the short circuit input impedance. Now, the ratio of the output current to input current at the short-circuited output port is:

$$\left. \frac{I_2}{I_1} \right|_{V_2=0} = h_{21}$$

This is called short-circuit current gain of the network. Now, let us open circuit the port 1. At that condition, there will be no input current ($I_1=0$) but open circuit voltage V_1 appears across the port 1, as shown below:



Now:

$$\left. \frac{V_1}{V_2} \right|_{I_1 = 0} = h_{12} = \text{open circuit reverse voltage gain}$$

This is referred as reverse voltage gain because, this is the ratio of input voltage to the output voltage of the network, but voltage gain is defined as the ratio of output voltage to the input voltage of a network.

Now:

$$\left. \frac{I_2}{V_2} \right|_{I_1 = 0} = h_{21}$$

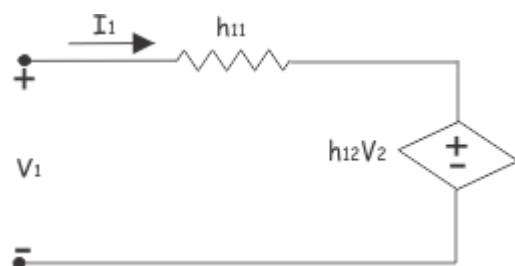
It is referred as open circuit output admittance.

To draw h parameter equivalent network of a two port network, first we have to write the equation of voltages and currents using h parameters. These are:

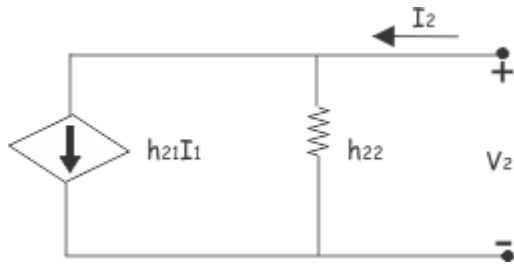
$$V_1 = h_{11}I_1 + h_{12}V_2 \dots\dots\dots(i)$$

$$I_2 = h_{21}I_1 + h_{22}V_2 \dots\dots\dots(ii)$$

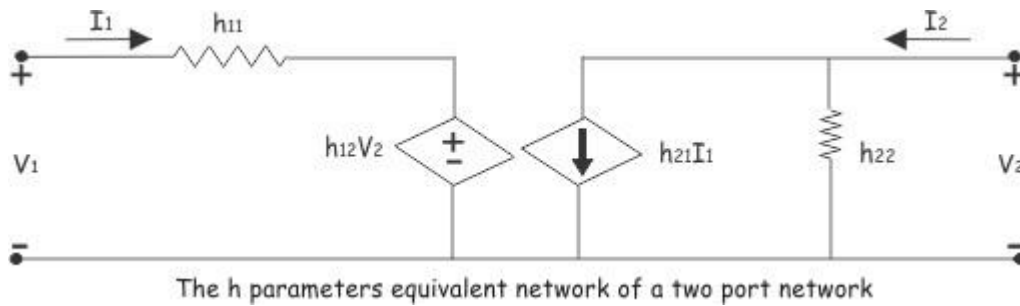
Equation (i) can be represented as a circuit based on Kirchhoff Voltage Law:



Equation (ii) can be represented as a circuit based on Kirchhoff Current Law:



Combining these two parts of the network we get:



GENERALISED APPROXIMATE MODEL:

In the analysis of transistor amplifier, we have as far used the exact h-model for the transistor. In practice, we may conveniently use an approximately h-model for the transistor which introduces error $< 10\%$ in most cases.

This much error may be conveniently tolerated since the h-parameters themselves are not steady but vary considerably for the same type of transistor. We first derive this approximate CE h-model.

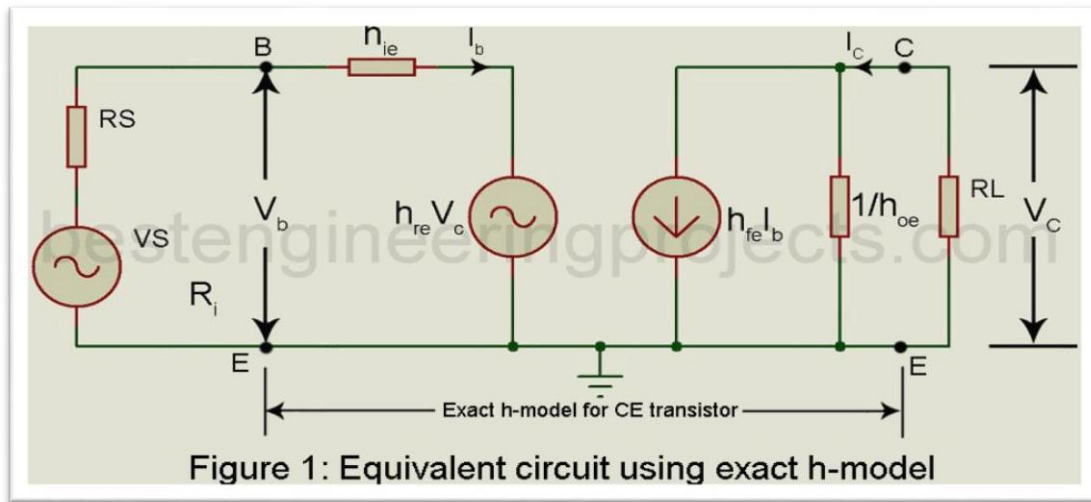
Figure 1 gives the equivalent circuit of CE amplifier using exact h-model for CE transistor.

The following steps are used to driving the approximate h-model:

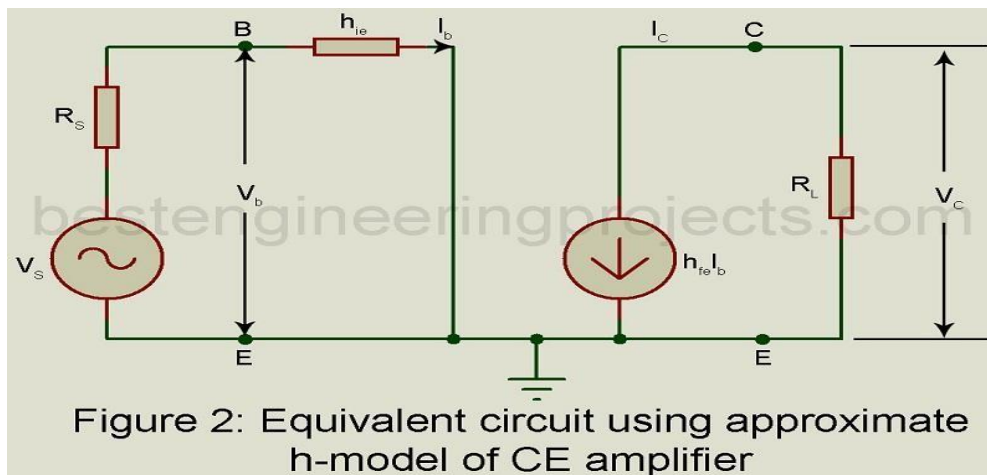
1. If $R_L < 0.1 \frac{1}{h_{oe}}$ and $h_{oe} \cdot R_L < 0.1$, then we may neglected, $\frac{1}{h_{oe}}$ being in parallel with R_L .

2. Having neglected h_{oe} , the collected current I_C equals $h_{fe} \cdot I_b$ and the magnitude of the dependent voltage generator in the emitter circuit is then given by,

$$h_{re} * |V_C| = h_{re} * I_C * R_L \approx h_{re} * h_{fe} * I_b * R_L \dots\dots(1)$$



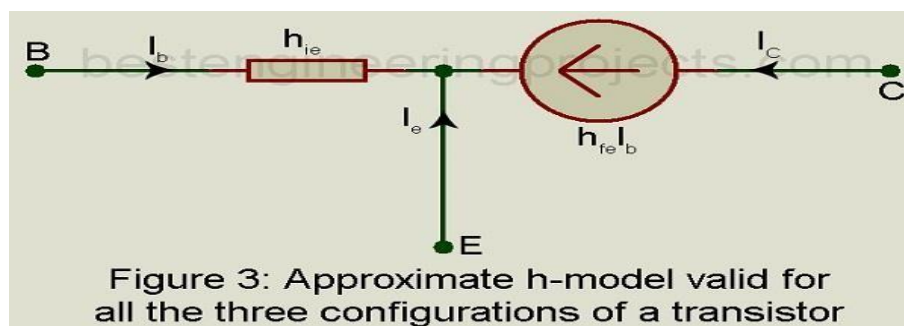
But $h_{re} * h_{fe} \approx 0$. Hence the voltage $h_{re} |V_c|$ in the emitter circuit may be neglected in comparison with the voltage drop $h_{ie} I_b$ provided that R_L is not very large. Then the approximate CE h-model reduces to the form shown in Figure 2.



ANALYSIS OF CB, CE, CC AMPLIFIER USING GENERALISED APPROXIMATE MODEL:

Approximate h-model Valid for all the three Configuration

The approximate CE h-model of Figure 2 is redrawn in figure 3. This model may be used for any of the three configurations by grounding the appropriate node and analysis done accordingly. It may be proved that the error in values of A_i , R_i , A_v or output terminal resistance $R_{ot} (= R_o || R_L)$ caused by use of approximate model does not exceed 10% if $h_{oe} * R_L < 0.1$.



Analysis of CE Amplifier using approximate h-model

Figure 2 gives the equivalent circuit of CE amplifier using approximate h-model for the transistor. For this equivalent circuit we get,

$$\text{Current gain } A_I = \frac{-h_{fe} X I_b}{I_b} = -h_{fe} \dots\dots(2)$$

$$\text{Input resistance } R_i = h_{ie}$$

$$\text{Voltage gain } A_V = A_I \times \frac{R_L}{R_i} = \frac{-h_{fe} X R_L}{h_{ie}} \dots\dots(3)$$

Output resistance R_0 : From this approximate equivalent circuit of figure 1(b) with $V_s = 0$ and with external voltage source connected across the output, we get $I_b = 0$ and therefore $I_c = 0$. Hence output resistance $R_0 = \infty$. However, in actual practice, R_0 lies between $40k\Omega$ and $80k\Omega$ depending on the value of R_S .

With load resistance $R_L = 4k\Omega$ (the maximum practical value), the output terminal resistance $R_t = R_L \parallel \infty = R_L = 4k\Omega$.

Condition $h_{oe} * R_L < 0.1$. For a typical transistor $h_{oe} = 25 * 10^{-6}S$. Hence to meet the condition that, $h_{oe} * R_L < 0.1$, we must use R_L less than $4k\Omega$.

Analysis of CB Amplifier using the Approximate Model

From figure 4 gives the equivalent circuit of a CB amplifier using the approximate model for the transistor as given in figure 2 with base grounded, the input applied between emitter and base and output obtained across load resistor R_L between the collector and the base.

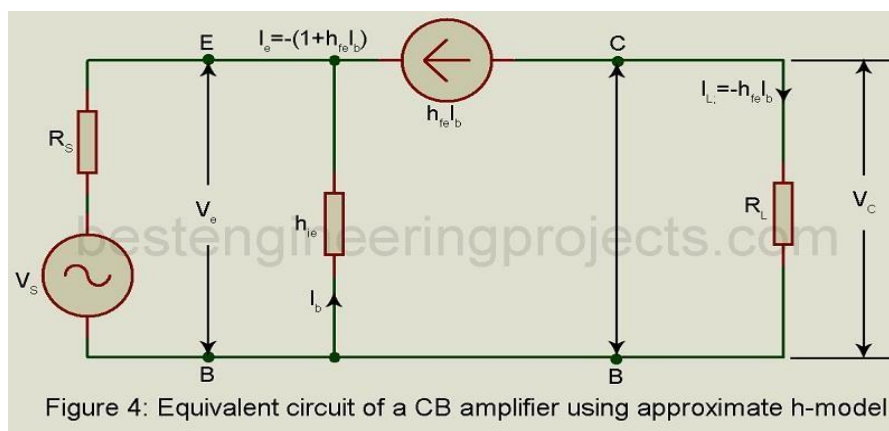


Figure 4: Equivalent circuit of a CB amplifier using approximate h-model

$$\text{Current gain } A_I = \frac{I_L}{I_e} = \frac{-h_{fe} X I_b}{-(1 + h_{fe}) X I_b} = \frac{h_{fe}}{1 + h_{fe}} \dots\dots(4)$$

Input resistance R_i : from figure 4,

$$V_e = -I_b * h_{ie} \dots\dots(5)$$

$$I_e = -(1 + h_{fe}) * I_b \dots\dots(6)$$

$$\text{Hence, } \frac{V_e}{I_e} = \frac{-I_b \times h_{ie}}{-I_b (1 + h_{fe})} = \frac{h_{ie}}{1 + h_{fe}} \dots (7)$$

Voltage Gain A_v : From figure 4,

$$\text{Hence, } A_v = \frac{V_e}{V_b} = \frac{-h_{fe} \times I_b \times R_L}{-I_b \times h_{ie}} = \frac{h_{fe} \times R_L}{h_{ie}} \dots (8)$$

Output resistance in the equivalent circuit of figure 3, with $V_s = 0$, we get $I_e = 0$. Hence, $I_b = 0$. Hence the output resistance $R_o = \infty$.

$$\text{Output Terminal Resistance } R_{ot} = R_o \parallel R_L = \infty \parallel R_L = R_L \dots (9)$$

Analysis of CC Amplifier (Emitter Follower) using approximate h-model

Figure 5 gives the equivalent circuit of an emitter follower using the approximate model as given in figure 3, with collector grounded, input signal applied between the base and the ground and the load impedance R_L connected between emitter and ground.

Current gain A_i : from the circuit of figure 5,

$$\text{Load current } I_L = (1 + h_{fe})I_b \dots (10)$$

$$\text{Hence Current gain } A_i = \frac{I_L}{I_b} = (1 + h_{fe}) \dots (11)$$

Input resistance R_i : from figure 5,

$$V_b = I_b \times h_{ie} + (1 + h_{fe})I_b \times R_L \dots (12)$$

$$\text{Hence, } R_i = \frac{V_b}{I_b} = h_{ie} + (1 + h_{fe})R_L \dots (13)$$

Voltage Gain A_v : From figure 5,

$$V_e = (1 + h_{fe})I_b \times R_L \dots (14) \text{Hence,}$$

$$A_v = \frac{V_e}{V_b} = \frac{(1 + h_{fe})I_b \times R_L}{I_b \times h_{ie} + (1 + h_{fe})I_b \times R_L} = \frac{(1 + h_{fe})R_L}{h_{ie} + (1 + h_{fe})R_L} = 1 - \frac{h_{ie}}{h_{ie} + (1 + h_{fe})R_L} = 1 - \frac{h_{ie}}{R_i}$$

Output Resistance from figure 5, Open circuit output voltage = V_s

Short circuit output current

$$I = (1 + h_{fe})I_b = \frac{(1 + h_{fe})V_s}{h_{ie} + R_s}$$

Hence output impedance

$$R_0 = \frac{\text{Open circuit output voltage}}{\text{Short circuit output current}} = \frac{h_{ie} + R_s}{1 + h_{fe}}$$

Output terminal Impedance $R_{ot} = R_0 \parallel R_L$

Expressions for current gain etc. for the three configurations using approximate h-model.

Expressions for A_i , R_i , A_v , R_0 and R_{ot} using Approximate h-model

Quantity	CE	CB	CC
A_i	$-h_{fe}$	$\frac{h_{fe}}{1 + h_{fe}}$	$1 + h_{fe}$
R_i	h_{ie}	$\frac{h_{ie}}{1 + h_{fe}}$	$h_{ie} + (1 + h_{fe})R_L$
A_v	$\frac{h_{fe} \times R_L}{h_{ie}}$	$\frac{h_{fe} \times R_L}{R_e}$	$1 - \frac{h_{ie}}{R_i}$
R_0	∞	∞	$\frac{h_{ie} + R_s}{1 + h_{fe}}$
R_{ot}	R_L	R_L	$R_0 \parallel R_L$

MULTISTAGE TRANSISTOR AMPLIFIERS:

In Multi-stage amplifiers, the output of first stage is coupled to the input of next stage using a coupling device. The following figure shows a two-stage amplifier connected in cascade.



If there are n numbers of stages, the product of voltage gains of those n stages will be the overall gain of that multistage amplifier circuit.

COUPLING:-

Coupling is a process in which the output of one stage is fed as input to the next stage. The main purpose of coupling is to:-

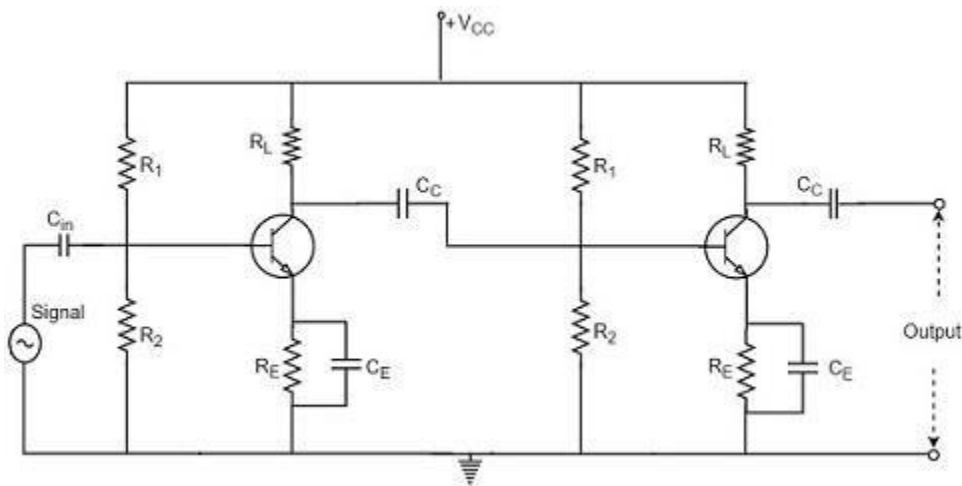
- 1) Transfer output of one stage to the input of next stage.
- 2) Isolate the dc condition of one stage from the next stage.

TYPES OF COUPLING:

Coupling is classified into three types. They are:-

- ✚ RC coupling
- ✚ Transformer coupling
- ✚ Direct coupling

RC COUPLED TRANSISTOR AMPLIFIER:-

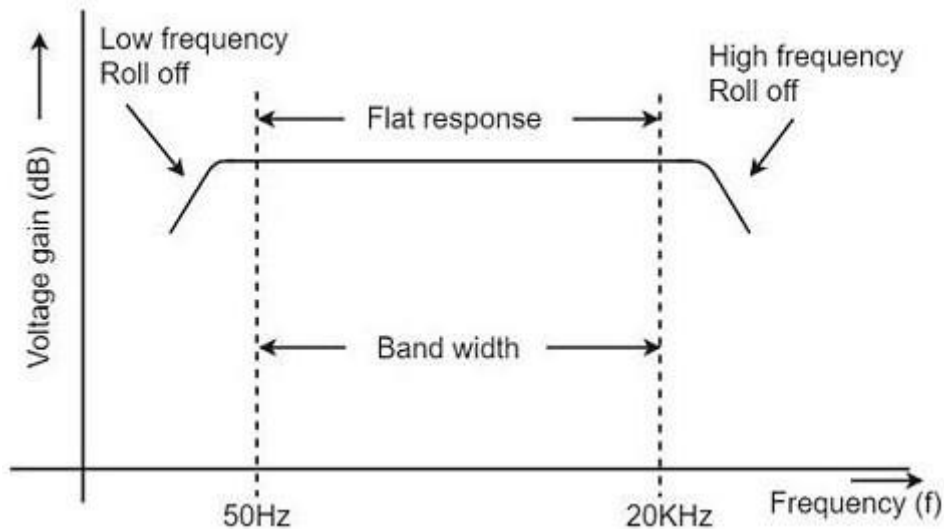


Operation of RC Coupled Amplifier

When an AC input signal is applied to the base of first transistor, it gets amplified and appears at the collector load R_L which is then passed through the coupling capacitor C_C to the next stage. This becomes the input of the next stage, whose amplified output again appears across its collector load. Thus the signal is amplified in stage by stage action.

Frequency Response of RC Coupled Amplifier

Frequency response curve is a graph that indicates the relationship between voltage gain and function of frequency. The frequency response of a RC coupled amplifier is as shown in the following graph.



From the above graph, it is understood that the frequency rolls off or decreases for the frequencies below 50Hz and for the frequencies above 20 KHz. whereas the voltage gain for the range of frequencies between 50Hz and 20 KHz is constant.

We know that,

$$X_C = 1/2\pi fC$$

It means that the capacitive reactance is inversely proportional to the frequency.

At Low frequencies (i.e. below 50 Hz)

The capacitive reactance is inversely proportional to the frequency. At low frequencies, the reactance is quite high. The reactance of input capacitor C_{in} and the coupling capacitor C_C are so high that only small part of the input signal is allowed. The reactance of the emitter by pass capacitor C_E is also very high during low frequencies. Hence it cannot shunt the emitter resistance effectively. With all these factors, the voltage gain rolls off at low frequencies.

At High frequencies (i.e. above 20 KHz)

Again considering the same point, we know that the capacitive reactance is low at high frequencies. So, a capacitor behaves as a short circuit, at high frequencies. As a result of this, the loading effect of the next stage increases, which reduces the voltage gain. Along with this, as the capacitance of emitter diode decreases, it increases the base current of the transistor due to which the current gain (β) reduces. Hence the voltage gain rolls off at high frequencies.

At Mid-frequencies (i.e. 50 Hz to 20 KHz)

The voltage gain of the capacitors is maintained constant in this range of frequencies, as shown in figure. If the frequency increases, the reactance of the

capacitor C_C decreases which tends to increase the gain. But this lower capacitance reactive increases the loading effect of the next stage by which there is a reduction in gain.

Due to these two factors, the gain is maintained constant.

Advantages of RC Coupled Amplifier

The following are the advantages of RC coupled amplifier.

- The frequency response of RC amplifier provides constant gain over a wide frequency range, hence most suitable for audio applications.
- The circuit is simple and has lower cost because it employs resistors and capacitors which are cheap.
- It becomes more compact with the upgrading technology.

Disadvantages of RC Coupled Amplifier

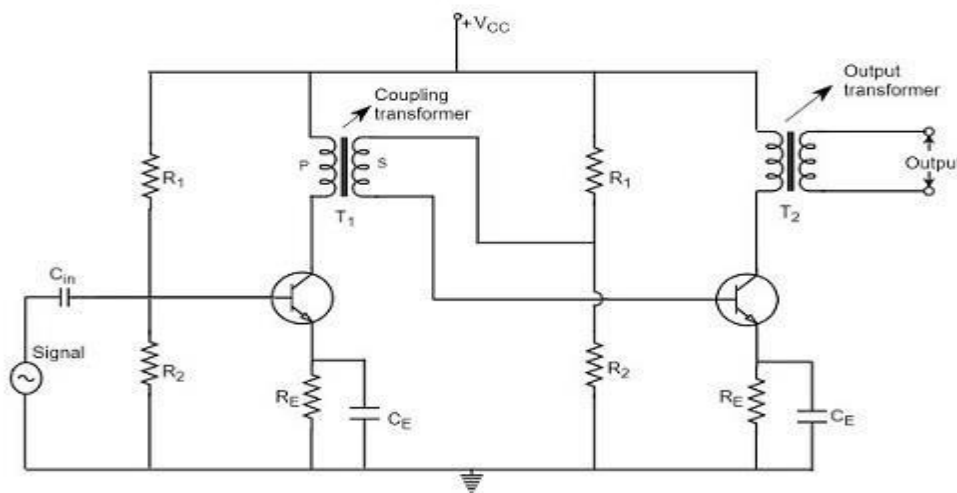
The following are the disadvantages of RC coupled amplifier.

- The voltage and power gain are low because of the effective load resistance.
- They become noisy with age.
- Due to poor impedance matching, power transfer will be low.

TRANSFORMER COUPLED AMPLIFIER:-

In a transformer-coupled amplifier, the stages of amplifier are coupled using a transformer. Let us go into the constructional and operational details of a transformer coupled amplifier.

The figure below shows the circuit diagram of transformer coupled amplifier.



Operation of Transformer Coupled Amplifier

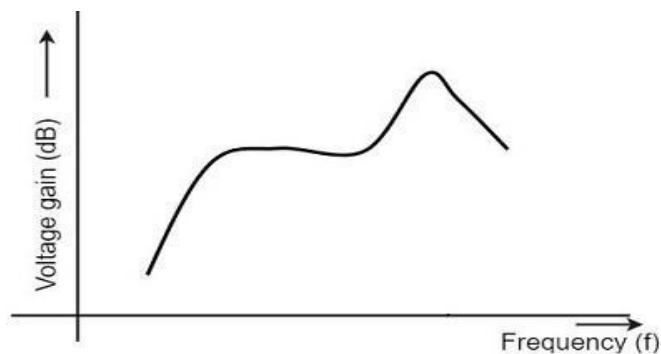
When an AC signal is applied to the input of the base of the first transistor then it gets amplified by the transistor and appears at the collector to which the primary of the transformer is connected.

The transformer which is used as a coupling device in this circuit has the property of impedance changing, which means the low resistance of a stage (or load) can be reflected as a high load resistance to the previous stage. Hence the voltage at the primary is transferred according to the turns ratio of the secondary winding of the transformer.

This transformer coupling provides good impedance matching between the stages of amplifier. The transformer coupled amplifier is generally used for power amplification.

Frequency Response of Transformer Coupled Amplifier

The figure below shows the frequency response of a transformer coupled amplifier. The gain of the amplifier is constant only for a small range of frequencies. The output voltage is equal to the collector current multiplied by the reactance of primary.



At low frequencies, the reactance of primary begins to fall, resulting in decreased gain. At high frequencies, the capacitance between turns of windings acts as a bypass condenser to reduce the output voltage and hence gain.

So, the amplification of audio signals will not be proportionate and some distortion will also get introduced, which is called as Frequency distortion.

Advantages of Transformer Coupled Amplifier

The following are the advantages of a transformer coupled amplifier –

- An excellent impedance matching is provided.
- Gain achieved is higher.
- There will be no power loss in collector and base resistors.
- Efficient in operation.

Disadvantages of Transformer Coupled Amplifier

The following are the disadvantages of a transformer coupled amplifier –

- Though the gain is high, it varies considerably with frequency. Hence a poor frequency response.
- Frequency distortion is higher.
- Transformers tend to produce hum noise.
- Transformers are bulky and costly.

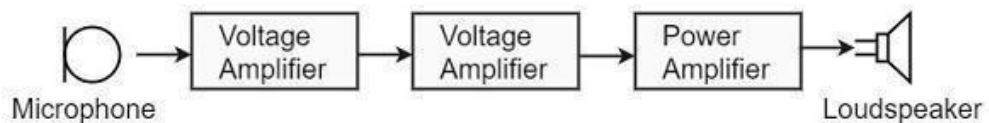
UNIT-2

AUDIO POWER AMPLIFIERS

In practice, any amplifier consists of few stages of amplification. If we consider audio amplification, it has several stages of amplification, depending upon our requirement.

Power Amplifier

After the audio signal is converted into electrical signal, it has several voltage amplifications done, after which the power amplification of the amplified signal is done just before the loud speaker stage. This is clearly shown in the below figure.



While the voltage amplifier raises the voltage level of the signal, the power amplifier raises the power level of the signal. Besides raising the power level, it can also be said that a power amplifier is a device which converts DC power to AC power and whose action is controlled by the input signal.

The DC power is distributed according to the relation, DC

$$\text{power input} = \text{AC power output} + \text{losses}$$

Power Transistor

For such Power amplification, a normal transistor would not do. A transistor that is manufactured to suit the purpose of power amplification is called as a **Power transistor**.

A Power transistor differs from the other transistors, in the following factors.

- It is larger in size, in order to handle large powers.
- The collector region of the transistor is made large and a heat sink is placed at the collector-base junction in order to minimize heat generated.
- The emitter and base regions of a power transistor are heavily doped.
- Due to the low input resistance, it requires low input power.

Hence there is a lot of difference in voltage amplification and power amplification. So, let us now try to get into the details to understand the differences between a voltage amplifier and a power amplifier.

Difference between Voltage and Power Amplifiers:

Let us try to differentiate between voltage and power amplifier.

Voltage Amplifier

The function of a voltage amplifier is to raise the voltage level of the signal. A voltage

amplifier is designed to achieve maximum voltage amplification.

The voltage gain of an amplifier is given by

$$A_v = \beta(R_c/R_{in})$$

The characteristics of a voltage amplifier are as follows –

- The base of the transistor should be thin and hence the value of β should be greater than 100.
- The resistance of the input resistor R_{in} should be low when compared to collector load R_c .
- The collector load R_c should be relatively high. To permit high collector load, the voltage amplifiers are always operated at low collector current.
- The voltage amplifiers are used for small signal voltages.

Power Amplifier

The function of a power amplifier is to raise the power level of input signal. It is required to deliver a large amount of power and has to handle large current.

The characteristics of a power amplifier are as follows –

- The base of transistor is made thicken to handle large currents. The value of β being ($\beta > 100$) high.
- The size of the transistor is made larger, in order to dissipate more heat, which is produced during transistor operation.
- Transformer coupling is used for impedance matching.
- Collector resistance is made low.

The comparison between voltage and power amplifiers is given below in a tabular form.

S.No	Particular	Voltage Amplifier	Power Amplifier
1	β	High (>100)	Low (5 to 20)
2	R_c	High (4-10 K Ω)	Low (5 to 20 Ω)
3	Coupling	Usually R-C coupling	Invariably transformer coupling
4	Input voltage	Low (a few m V)	High (2-4 V)
5	Collector current	Low (≈ 1 mA)	High (> 100 mA)
6	Power output	Low	High
7	Output impendence	High (≈ 12 K Ω)	Low (200 Ω)

The Power amplifiers amplify the power level of the signal. This amplification is done in the last stage in audio applications. The applications related to radio frequencies employ radio power amplifiers. But the **operating point** of a transistor plays a very important role in determining the efficiency of the amplifier. The **main classification** is done based on this mode of operation.

The classification is done based on their frequencies and also based on their mode of operation.

Classification Based on Frequencies

Power amplifiers are divided into two categories, based on the frequencies they handle. They are as follows.

- **Audio Power Amplifiers** – The audio power amplifiers raise the power level of signals that have audio frequency range (20 Hz to 20 KHz). They are also known as **Small signal power amplifiers**.
- **Radio Power Amplifiers** – Radio Power Amplifiers or tuned power amplifiers raise the power level of signals that have radio frequency range (3 KHz to 300 GHz). They are also known as **large signal power amplifiers**.

Classification Based on Mode of Operation

On the basis of the mode of operation, i.e., the portion of the input cycle during which collector current flows, the power amplifiers may be classified as follows.

- **Class A Power amplifier** – When the collector current flows at all times during the full cycle of signal, the power amplifier is known as **class A power amplifier**.
- **Class B Power amplifier** – When the collector current flows only during the positive half cycle of the input signal, the power amplifier is known as **class B power amplifier**.
- **Class C Power amplifier** – When the collector current flows for less than half cycle of the input signal, the power amplifier is known as **class C power amplifier**.

There forms another amplifier called Class AB amplifier, if we combine the class A and class B amplifiers so as to utilize the advantages of both. Before going into the details of these amplifiers, let us have a look at the important terms that have to be considered to determine the efficiency of an amplifier.

Terms Considering Performance

The primary objective of a power amplifier is to obtain maximum output power. In order to

achieve this, the important factors to be considered are collector efficiency, power dissipation capability and distortion. Let us go through them in detail.

Collector Efficiency

This explains how well an amplifier converts DC power to AC power. When the DC supply is given by the battery but no AC signal input is given, the collector output at such a condition is observed as **collector efficiency**.

The collector efficiency is defined as

$$\eta = \text{average a.c power output} / \text{average d.c power input to transistor}$$

The main aim of a power amplifier is to obtain maximum collector efficiency. Hence the higher the value of collector efficiency, the efficient the amplifier will be.

Power Dissipation Capacity

Every transistor gets heated up during its operation. As a power transistor handles large currents, it gets more heated up. This heat increases the temperature of the transistor, which alters the operating point of the transistor. So, in order to maintain the operating point stability, the temperature of the transistor has to be kept in permissible limits. For this, the heat produced has to be dissipated. Such a capacity is called as Power dissipation capability.

Power dissipation capability can be defined as the ability of a power transistor to dissipate the heat developed in it. Metal cases called heat sinks are used in order to dissipate the heat produced in power transistors.

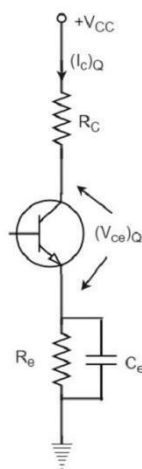
Distortion

A transistor is a non-linear device. When compared with the input, there occur few variations in the output. In voltage amplifiers, this problem is not pre-dominant as small currents are used. But in power amplifiers, as large currents are in use, the problem of distortion certainly arises.

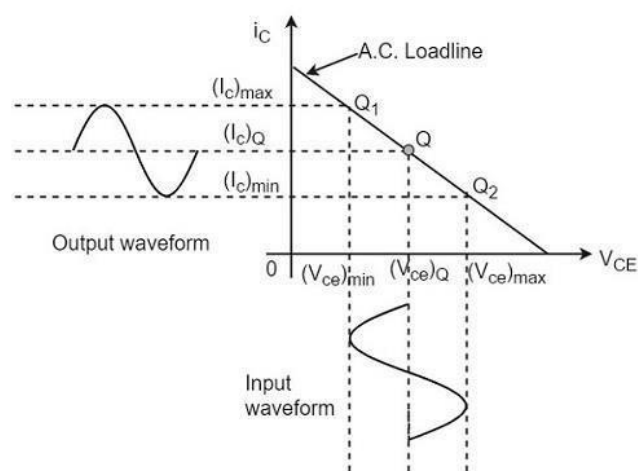
Distortion is defined as the change of output wave shape from the input wave shape of the amplifier. An amplifier that has lesser distortion produces a better output and hence considered efficient.

We have already come across the details of transistor biasing, which is very important for the operation of a transistor as an amplifier. Hence to achieve faithful amplification, the biasing of the transistor has to be done such that the amplifier operates over the linear region.

A Class A power amplifier is one in which the output current flows for the entire cycle of the AC input supply. Hence the complete signal present at the input is amplified at the output. The following figure shows the circuit diagram for Class A Power amplifier.



From the above figure, it can be observed that the transformer is present at the collector as a load. The use of transformer permits the impedance matching, resulting in the transference of maximum power to the load e.g. loud speaker.



The operating point of this amplifier is present in the linear region. It is so selected that the current flows for the entire ac input cycle. The below figure explains the selection of operating point.

The output characteristics with operating point Q is shown in the figure above. Here $(I_c)_Q$ and $(V_{ce})_Q$ represent no signal collector current and voltage between collector and emitter respectively. When signal is applied, the Q-point shifts to Q_1 and Q_2 . The output current increases to $(I_c)_{\max}$ and decreases to $(I_c)_{\min}$. Similarly, the collector-emitter voltage increases to $(V_{ce})_{\max}$ and decreases to $(V_{ce})_{\min}$.

D.C. Power drawn from collector battery V_{cc} is given by

$$P_{in} = \text{voltage} \times \text{current} = V_{cc}(I_c)_Q$$

This power is used in the following two parts –

- Power dissipated in the collector load as heat is given by

$$P_{RC} = (\text{current})^2 \times \text{resistance} = (I_C)^2_Q R_C$$

- Power given to transistor is given by

$$P_{tr} = P_{in} - P_{RC} = V_{cc} - (I_C)_Q R_C$$

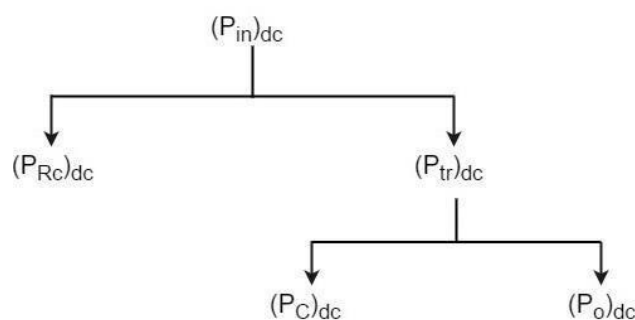
When signal is applied, the power given to transistor is used in the following two parts –

- A.C. Power developed across load resistors R_C which constitutes the a.c. power output.

$$(P_o)_{ac} = I^2 R_C = V^2 / R_C = (V_m / \sqrt{2})^2 / R_C = V_m^2 / 2 R_C$$

- Where I is the R.M.S. value of a.c. output current through load, V is the R.M.S. value of a.c. voltage, and V_m is the maximum value of V .
- The D.C. power dissipated by the transistor (collector region) in the form of heat, i.e., $(P_C)_{dc}$

We have represented the whole power flow in the following diagram.



This class A power amplifier can amplify small signals with least distortion and the output will be an exact replica of the input with increased strength.

Let us now try to draw some expressions to represent efficiencies.

Overall Efficiency

The overall efficiency of the amplifier circuit is given by

$$\begin{aligned}(\eta)_{\text{overall}} &= \frac{\text{a. c power delivered to the load}}{\text{total power delivered by d. c supply}} \\ &= \frac{(P_O)_{ac}}{(P_{in})_{dc}}\end{aligned}$$

Collector Efficiency

The collector efficiency of the transistor is defined as

$$\begin{aligned}(\eta)_{\text{collector}} &= \frac{\text{average a. c power output}}{\text{average d. c power input to transistor}} \\ &= \frac{(P_O)_{ac}}{(P_{tr})_{dc}}\end{aligned}$$

Expression for overall efficiency

$$\begin{aligned}(P_O)_{ac} &= V_{rms} \times I_{rms} \\ &= \frac{1}{\sqrt{2}} \left[\frac{(V_{ce})_{max} - (V_{ce})_{min}}{2} \right] \times \frac{1}{\sqrt{2}} \left[\frac{(I_C)_{max} - (I_C)_{min}}{2} \right] \\ &= \frac{[(V_{ce})_{max} - (V_{ce})_{min}] \times [(I_C)_{max} - (I_C)_{min}]}{8}\end{aligned}$$

Advantages of Class A Amplifiers

The advantages of Class A power amplifier are as follows –

- The current flows for complete input cycle
- It can amplify small signals
- The output is same as input
- No distortion is present

Disadvantages of Class A Amplifiers

The advantages of Class A power amplifier are as follows –

- Low power output
- Low collector efficiency

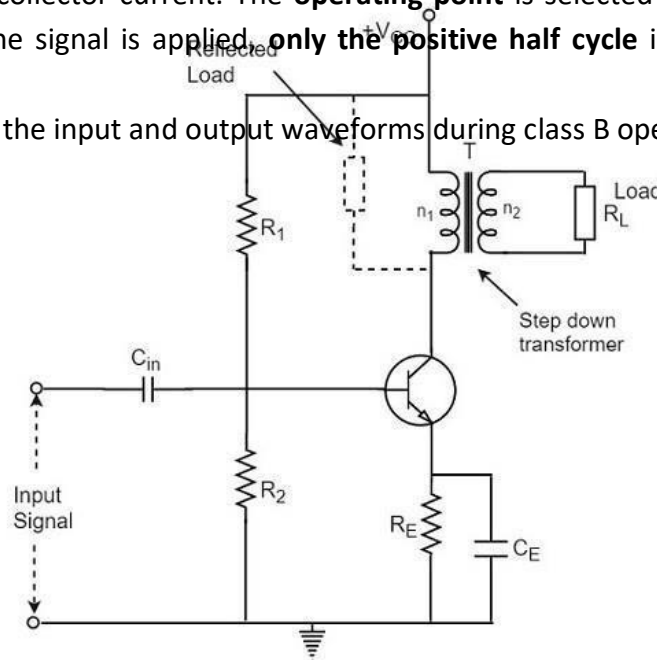
The class A power amplifier as discussed in the previous chapter, is the circuit in which the output current flows for the entire cycle of the AC input supply. We also have learnt about the

disadvantages it has such as low output power and efficiency. In order to minimize those effects, the transformer coupled class A power amplifier has been introduced.

Class B- WORKING PRINCIPLE

The biasing of the transistor in class B operation is in such a way that at zero signal condition, there will be no collector current. The **operating point** is selected to be at collector cut off voltage. So, when the signal is applied, **only the positive half cycle** is amplified at the output.

The figure below shows the input and output waveforms during class B operation.



When the signal is applied, the circuit is forward biased for the positive half cycle of the input and hence the collector current flows. But during the negative half cycle of the input, the circuit is reverse biased and the collector current will be absent. Hence **only the positive half cycle** is amplified at the output.

As the negative half cycle is completely absent, the signal distortion will be high. Also, when the applied signal increases, the power dissipation will be more. But when compared to class A power amplifier, the output efficiency is increased. Well, in order to minimize the disadvantages and achieve low distortion, high efficiency and high output power, the push-pull configuration is used in this class B amplifier.

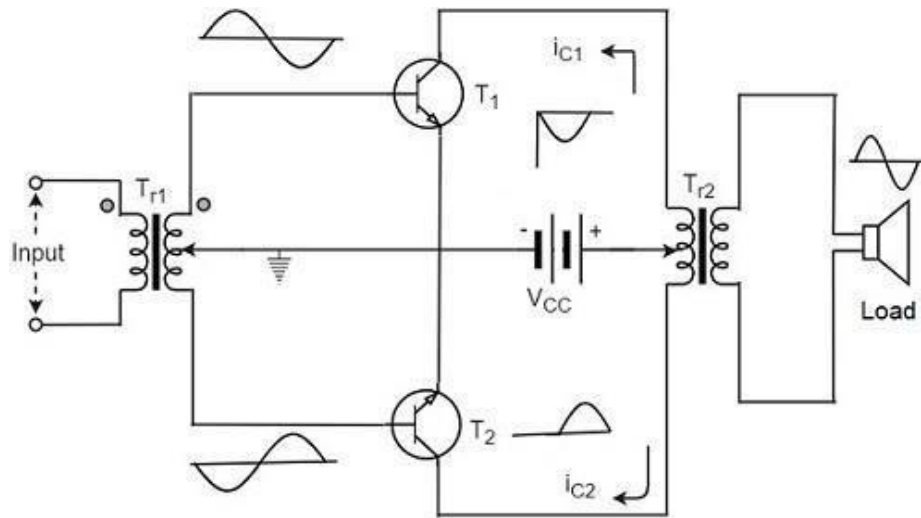
Class B Push-Pull Amplifier

Though the efficiency of class B power amplifier is higher than class A, as only one half cycle of the input is used, the distortion is high. Also, the input power is not completely utilized. In order to compensate these problems, the push-pull configuration is introduced in class B amplifier.

Construction:

The circuit of a push-pull class B power amplifier consists of two identical transistors T_1 and T_2 whose bases are connected to the secondary of the center-tapped input transformer T_{r1} . The emitters are shorted and the collectors are given the V_{CC} supply through the primary of the output transformer T_{r2} .

The circuit arrangement of class B push-pull amplifier, is same as that of class A push-pull amplifier except that the transistors are biased at cut off, instead of using the biasing resistors. The figure below gives the detailing of the construction of a push-pull class B power amplifier.

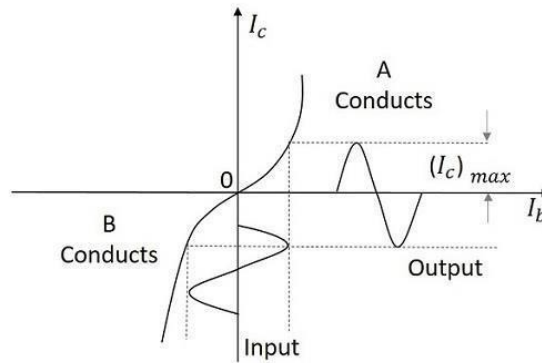


The circuit operation of class B push pull amplifier is detailed below.

Operation

The circuit of class B push-pull amplifier shown in the above figure clears that both the transformers are center-tapped. When no signal is applied at the input, the transistors T_1 and T_2 are in cut off condition and hence no collector currents flow. As no current is drawn from V_{CC} , no power is wasted.

When input signal is given, it is applied to the input transformer T_{r1} which splits the signal into two signals that are 180° out of phase with each other. These two signals are given to the two identical transistors T_1 and T_2 . For the positive half cycle, the base of the transistor T_1 becomes positive and collector current flows. At the same time, the transistor T_2 has negative half cycle, which throws the transistor T_2 into cutoff condition and hence no collector current flows. The waveform is produced as shown in the following figure.



For the next half cycle, the transistor T_1 gets into cut off condition and the transistor T_2 gets into conduction, to contribute the output. Hence for both the cycles, each transistor conducts alternately. The output transformer T_{r3} serves to join the two currents producing an almost undistorted output waveform.

Power Efficiency of Class B Push-Pull Amplifier

The current in each transistor is the average value of half sine loop. For half sine loop, I_{dc} is given by
$$I_{dc} = \frac{(I_c)_{max}}{\pi}$$

Therefore,

$$(P_{in})_{dc} = 2 \times \left[\frac{(I_c)_{max}}{\pi} \times V_{CC} \right]$$

Here factor 2 is introduced as there are two transistors in push-pull amplifier.

R.M.S. value of collector current = $(I_c)_{max} / \sqrt{2}$

R.M.S. value of output voltage = $V_{CC} / \sqrt{2}$

Under ideal conditions of maximum power

Therefore,

$$(P_o)_{ac} = \frac{(I_c)_{max}}{\sqrt{2}} \times \frac{V_{CC}}{\sqrt{2}} = \frac{(I_c)_{max} \times V_{CC}}{2}$$

Now overall maximum efficiency

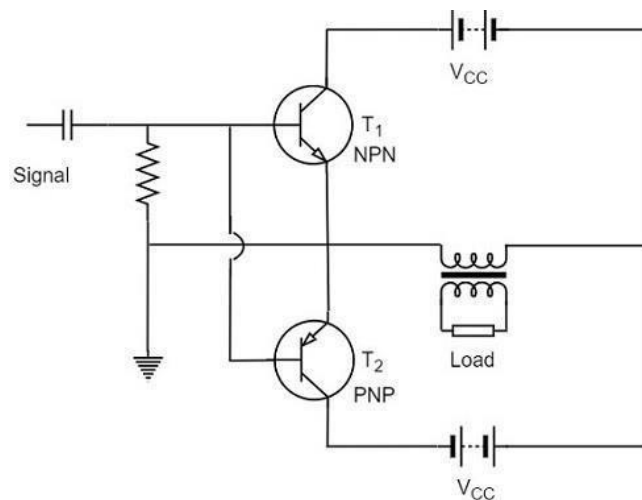
$$\begin{aligned} \eta_{overall} &= \frac{(P_o)_{ac}}{(P_{in})_{dc}} \\ &= \frac{(I_c)_{max} \times V_{CC}}{2} \times \frac{\pi}{2(I_c)_{max} \times V_{CC}} \\ &= \frac{\pi}{4} = 0.785 = 78.5\% \end{aligned}$$

The collector efficiency would be the same.

Hence the class B push-pull amplifier improves the efficiency than the class A push-pull amplifier.

Complementary Symmetry Push-Pull Class B Amplifier

The push pull amplifier which was just discussed improves efficiency but the usage of center-tapped transformers makes the circuit bulky, heavy and costly. To make the circuit simple and to improve the efficiency, the transistors used can be complemented, as shown in the following circuit diagram.



The above circuit employs a NPN transistor and a PNP transistor connected in push pull configuration. When the input signal is applied, during the positive half cycle of the input signal, the NPN transistor conducts and the PNP transistor cuts off. During the negative half cycle, the NPN transistor cuts off and the PNP transistor conducts.

In this way, the NPN transistor amplifies during positive half cycle of the input, while PNP transistor amplifies during negative half cycle of the input. As the transistors are both complement to each other, yet act symmetrically while being connected in push pull configuration of class B, this circuit is termed as **Complementary symmetry push pull class B amplifier**.

Advantages

The advantages of Complementary symmetry push pull class B amplifier are as follows.

- As there is no need of center tapped transformers, the weight and cost are reduced.
- Equal and opposite input signal voltages are not required.

Disadvantages

The disadvantages of Complementary symmetry push pull class B amplifier are as follows.

- It is difficult to get a pair of transistors (NPN and PNP) that have similar characteristics.
- We require both positive and negative supply voltages.

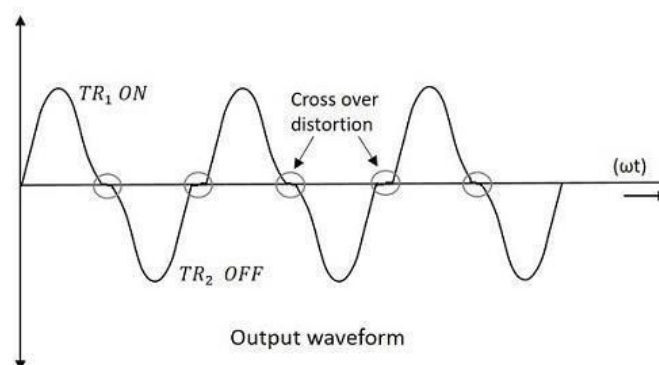
The class A and class B amplifier so far discussed has got few limitations. Let us now try to combine these two to get a new circuit which would have all the advantages of both class A and class B amplifier without their inefficiencies. Before that, let us also go through another important problem, called as **Cross over distortion**, the output of class B encounters with.

Cross-over Distortion:

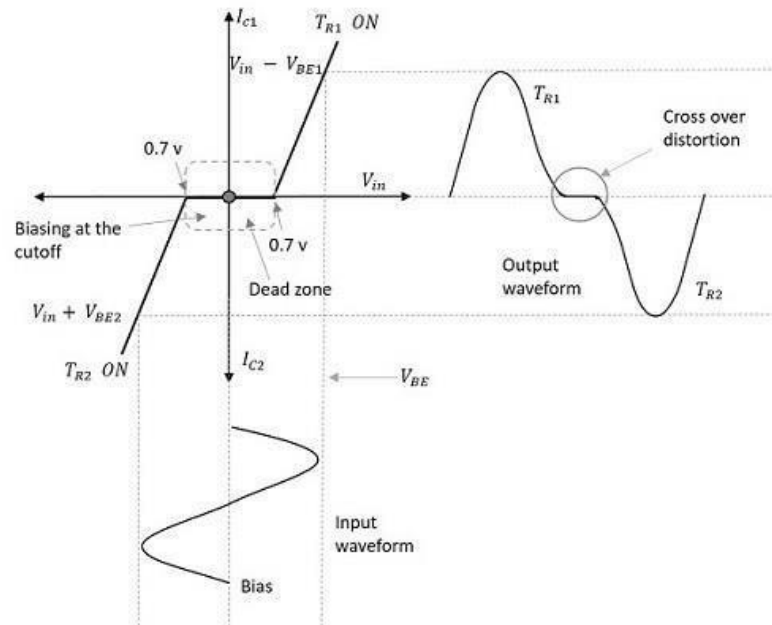
In the push-pull configuration, the two identical transistors get into conduction, one after the other and the output produced will be the combination of both.

When the signal changes or crosses over from one transistor to the other at the zero voltage point, it produces an amount of distortion to the output wave shape. For a transistor in order to conduct, the base emitter junction should cross $0.7V$, the cut off voltage. The time taken for a transistor to get ON from OFF or to get OFF from ON state is called the **transition period**.

At the zero voltage point, the transition period of switching over the transistors from one to the other, has its effect which leads to the instances where both the transistors are OFF at a time. Such instances can be called as **Flat spot** or **Dead band** on the output wave shape.



The above figure clearly shows the cross over distortion which is prominent in the output waveform. This is the main disadvantage. This cross over distortion effect also reduces the overall peak to peak value of the output waveform which in turn reduces the maximum power output. This can be more clearly understood through the non-linear characteristic of the waveform as shown below.



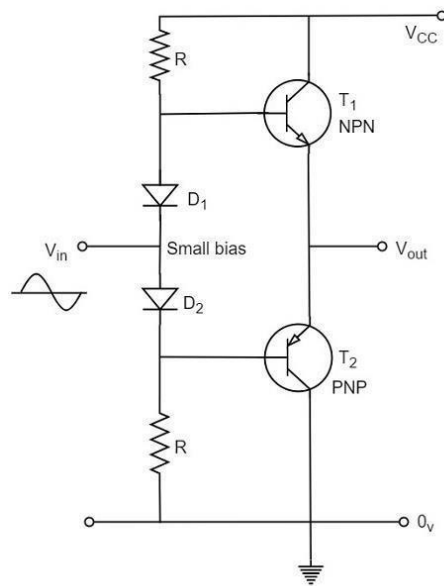
It is understood that this cross-over distortion is less pronounced for large input signals, where as it causes severe disturbance for small input signals. This cross over distortion can be eliminated if the conduction of the amplifier is more than one half cycle, so that both the transistors won't be OFF at the same time.

This idea leads to the invention of class AB amplifier, which is the combination of both class A and class B amplifiers, as discussed below.

Class AB Power Amplifier

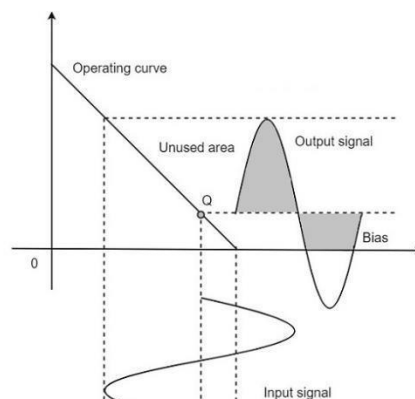
As the name implies, class AB is a combination of class A and class B type of amplifiers. As class A has the problem of low efficiency and class B has distortion problem, this class AB is emerged to eliminate these two problems, by utilizing the advantages of both the classes.

The cross over distortion is the problem that occurs when both the transistors are OFF at the same instant, during the transition period. In order to eliminate this, the condition has to be chosen for more than one half cycle. Hence, the other transistor gets into conduction, before the operating transistor switches to cut off state. This is achieved only by using class AB configuration, as shown in the following circuit diagram.



Therefore, in class AB amplifier design, each of the push-pull transistors is conducting for slightly more than the half cycle of conduction in class B, but much less than the full cycle of conduction of class A.

The conduction angle of class AB amplifier is somewhere between 180° to 360° depending upon the operating point selected. This is understood with the help of below figure.



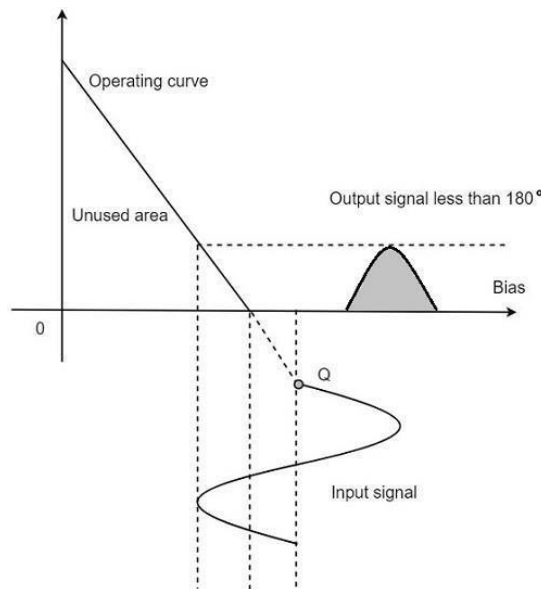
The small bias voltage given using diodes D_1 and D_2 , as shown in the above figure, helps the operating point to be above the cutoff point. Hence the output waveform of class AB results as seen in the above figure. The crossover distortion created by class B is overcome by this class AB, as well the inefficiencies of class A and B don't affect the circuit.

So, the class AB is a good compromise between class A and class B in terms of efficiency and linearity having the efficiency reaching about 50% to 60%. The class A, B and AB amplifiers are called as **linear amplifiers** because the output signal amplitude and phase are linearly related to the input signal amplitude and phase.

Class C Power Amplifier

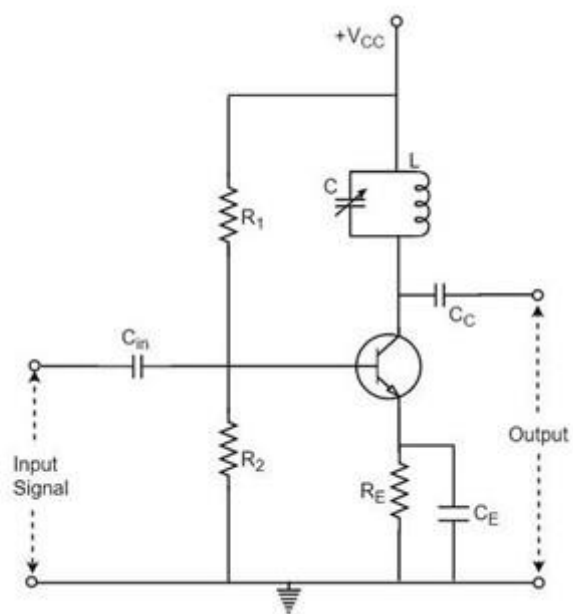
When the collector current flows for less than half cycle of the input signal, the power amplifier is known as **class C power amplifier**. The efficiency of class C amplifier is high while linearity is poor. The conduction angle for class C is less than 180° . It is generally around 90° , which means the transistor remains idle for more than half of the input signal. So, the output current will be delivered for less time compared to the application of input signal.

The following figure shows the operating point and output of a class C amplifier.



This kind of biasing gives a much improved efficiency of around 80% to the amplifier, but introduces heavy distortion in the output signal. Using the class C amplifier, the pulses produced at its output can be converted to complete sine wave of a particular frequency by using LC circuits in its collector circuit.

The types of amplifiers that we have discussed so far cannot work effectively at radio frequencies, even though they are good at audio frequencies. Also, the gain of these amplifiers is such that it will not vary according to the frequency of the signal, over a wide range. This allows the amplification of the signal equally well over a range of frequencies and does not permit the selection of particular desired frequency while rejecting the other frequencies.



UNIT 3

FIELD EFFECT TRANSISTOR

INTRODUCTION

1. The Field effect transistor is abbreviated as FET , it is an another semiconductor device like a BJT which can be used as an amplifier or switch.
2. The Field effect transistor is a voltage operated device. Whereas Bipolar junction transistor is a current controlled device. Unlike BJT a FET requires virtually no input current.
3. This gives it an extremely high input resistance , which is its most important advantage over a bipolar transistor.
4. FET is also a three terminal device, labeled as source, drain and gate.
5. The source can be viewed as BJT's emitter, the drain as collector, and the gate as the counter part of the base.
6. The material that connects the source to drain is referred to as the channel.
7. FET operation depends only on the flow of majority carriers ,therefore they are called uni polar devices. BJT operation depends on both minority and majority carriers.
8. As FET has conduction through only majority carriers it is less noisy than BJT.
9. FETs are much easier to fabricate and are particularly suitable for ICs because they occupy less space than BJTs.
10. FET amplifiers have low gain bandwidth product due to the junction capacitive effects and produce more signal distortion except for small signal operation.
11. The performance of FET is relatively unaffected by ambient temperature changes. As it has a negative temperature coefficient at high current levels, it prevents the FET from thermal breakdown. The BJT has a positive temperature coefficient at high current levels which leads to thermal breakdown.

CLASSIFICATION OF FET:

There are two major categories of field effect transistors:

1. Junction Field Effect Transistors
2. MOSFETs

These are further sub divided in to P- channel and N-channel devices.

MOSFETs are further classified in to two types Depletion MOSFETs and Enhancement . MOSFETs

When the channel is of N-type the JFET is referred to as an N-channel JFET ,when the channel is of P-type the JFET is referred to as P-channel JFET.

The schematic symbols for the P-channel and N-channel JFETs are shown in the figure.



Fig 5.1 schematic symbols for the P-channel and N-channel JFET

CONSTRUCTION AND OPERATION OF N- CHANNEL FET

If the gate is an N-type material, the channel must be a P-type material.

CONSTRUCTION OF N-CHANNEL JFET

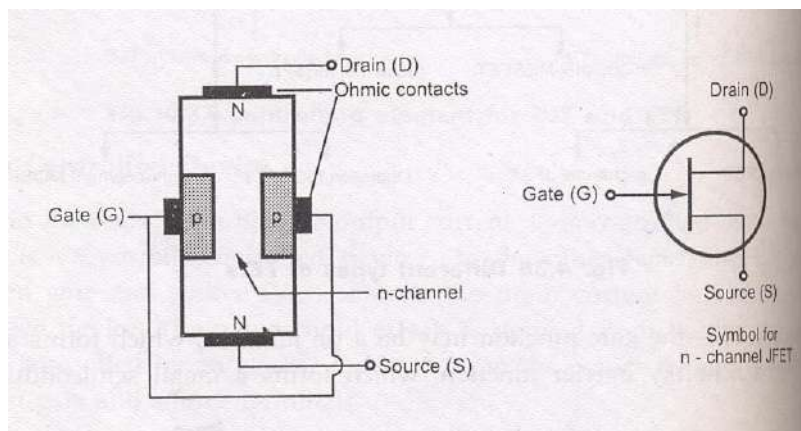


Fig 5.2 Construction of N-Channel JFET

A piece of N- type material, referred to as channel has two smaller pieces of P-type material attached to its sides, forming PN junctions. The channel ends are designated as the drain and source. And the two pieces of P-type material are connected together and their terminal is called the gate. Since this channel is in the N-type bar, the FET is known as N-channel JFET.

OPERATION OF N-CHANNEL JFET:-

The overall operation of the JFET is based on varying the width of the channel to control the drain current.

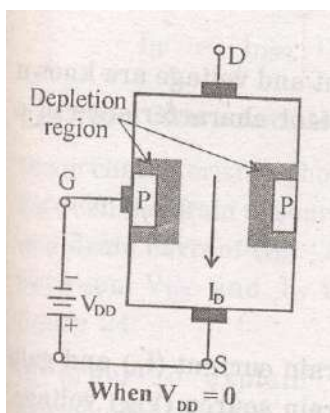
A piece of N type material referred to as the channel, has two smaller pieces of P type material attached to its sides, forming PN –Junctions. The channel's ends are designated the drain and the source. And the two pieces of P type material are connected together and their terminal is called the gate. With the gate terminal not connected and the potential applied positive at the drain negative at the source a drain current I_d flows. When the gate is biased negative with respect to the source the PN junctions are reverse biased and depletion regions are formed. The channel is more lightly doped than the P type gate blocks, so the depletion regions penetrate deeply into the channel. Since depletion region is a region depleted of charge carriers it behaves as an Insulator. The result is that the channel is narrowed. Its resistance is increased and I_d is reduced. When the negative gate bias voltage is further increased, the depletion regions meet at the center and I_d is cut off completely.

There are two ways to control the channel width

1. By varying the value of V_{gs}
2. And by Varying the value of V_{ds} holding V_{gs} constant

1 By varying the value of V_{gs} :-

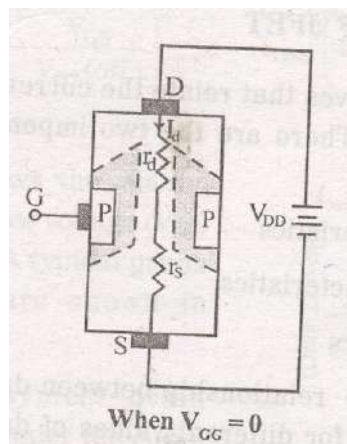
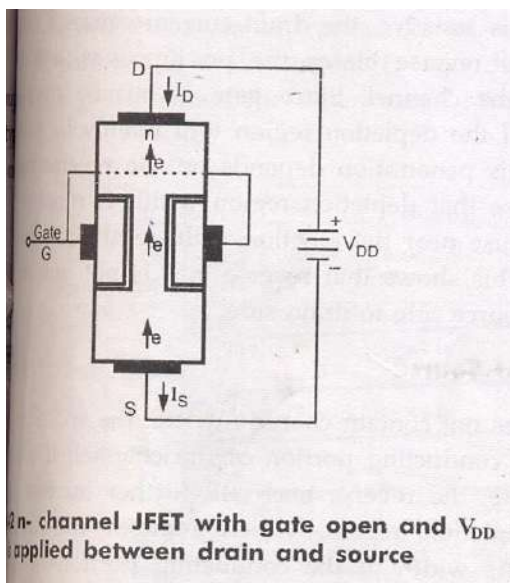
We can vary the width of the channel and in turn vary the amount of drain current. This can be done by varying the value of V_{gs} . This point is illustrated in the fig below. Here we are dealing with N channel FET. So channel is of N type and gate is of P type that constitutes a PN junction. This PN junction is always reverse biased in JFET operation .The reverse bias is applied by a battery voltage V_{gs} connected between the gate and the source terminal i.e positive terminal of the battery is connected to the source and negative terminal to gate.



- 1) When a PN junction is reverse biased the electrons and holes diffuse across junction by leaving immobile ions on the N and P sides , the region containing these immobile ions is known as depletion regions.
- 2) If both P and N regions are heavily doped then the depletion region extends symmetrically on both sides.
- 3) But in N channel FET P region is heavily doped than N type thus depletion region extends more in N region than P region.
- 4) So when no V_{ds} is applied the depletion region is symmetrical and the conductivity becomes Zero. Since there are no mobile carriers in the junction.
- 5) As the reverse bias voltage is increases the thickness of the depletion region also increases. i.e. the effective channel width decreases .
- 6) By varying the value of V_{gs} we can vary the width of the channel.

2 Varying the value of V_{ds} holding V_{gs} constant :-

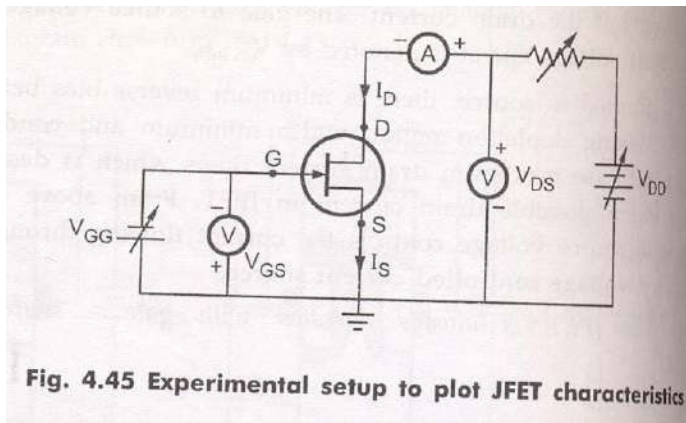
- 1) When no voltage is applied to the gate i.e. $V_{gs}=0$, V_{ds} is applied between source and drain the electrons will flow from source to drain through the channel constituting drain current I_d .
- 2) With $V_{gs}= 0$ for $I_d= 0$ the channel between the gate junctions is entirely open .In response to a small applied voltage V_{ds} , the entire bar acts as a simple semi conductor resistor and the current I_d increases linearly with V_{ds} .
- 3) The channel resistances are represented as r_d and r_s as shown in the fig.



- 4) This increasing drain current I_d produces a voltage drop across r_d which reverse biases the gate to source junction, ($r_d > r_s$) .Thus the depletion region is formed which is not symmetrical .

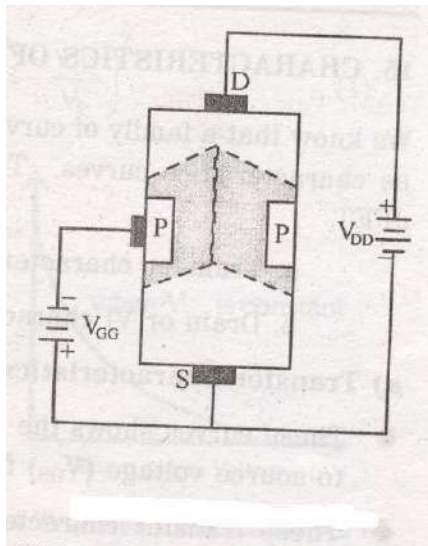
- 5) The depletion region i.e. developed penetrates deeper in to the channel near drain and less towards source because $V_{rd} \gg V_{rs}$. So reverse bias is higher near drain than at source.
- 6) As a result growing depletion region reduces the effective width of the channel. Eventually a voltage V_{ds} is reached at which the channel is pinched off. This is the voltage where the current I_d begins to level off and approach a constant value.
- 7) So, by varying the value of V_{ds} we can vary the width of the channel holding V_{gs} constant.

When both V_{gs} and V_{ds} is applied:-



It is of course in principle not possible for the channel to close Completely and there by reduce the current I_d to Zero for, if such indeed, could be the case the gate voltage V_{gs} is applied in the direction to provide additional reverse bias

- 1) When voltage is applied between the drain and source with a battery V_{dd} , the electrons flow from source to drain through the narrow channel existing between the depletion regions. This constitutes the drain current I_d , its conventional direction is from drain to source.
- 2) The value of drain current is maximum when no external voltage is applied between gate and source and is designated by I_{dss} .



- 3) When V_{gs} is increased beyond Zero the depletion regions are widened. This reduces the effective width of the channel and therefore controls the flow of drain current through the channel.
- 4) When V_{gs} is further increased a stage is reached at which the depletion regions touch each other that means the entire channel is closed with depletion region. This reduces the drain current to Zero.

CHARACTERISTICS OF N-CHANNEL JFET

The family of curves that shows the relation between current and voltage are known as characteristic curves.

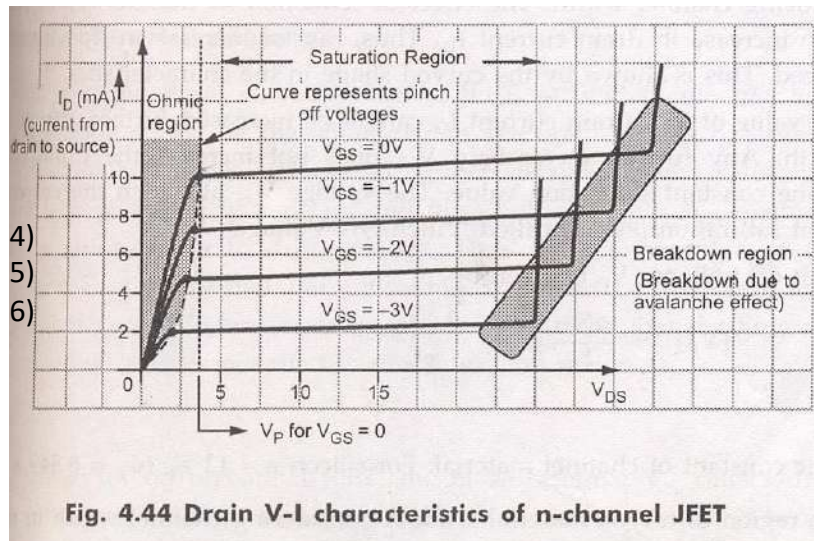
There are two important characteristics of a JFET.

- 1) Drain or V_I Characteristics
- 2) Transfer characteristics

1. Drain Characteristics:-

Drain characteristics show the relation between the drain to source voltage V_{ds} and drain current I_d . In order to explain typical drain characteristics let us consider the curve with $V_{gs} = 0V$.

- 1) When V_{ds} is applied and it is increasing the drain current I_D also increases linearly up to knee point.
- 2) This shows that FET behaves like an ordinary resistor. This region is called as ohmic region.
- 3) I_D increases with increase in drain to source voltage. Here the drain current is increased slowly as compared to ohmic region.



4) It is because of the fact that there is an increase in V_{DS} . This in turn increases the reverse bias voltage across the gate source junction. As a result of this depletion region grows in size thereby reducing the effective width of the channel.

5) All the drain to source voltage corresponding to point the channel width is reduced to a minimum value and is known as pinch off.

5) The drain to source voltage at which channel pinch off occurs is called pinch off voltage (V_P).

PINCH OFF Region:-

- 1) This is the region shown by the curve as saturation region.
- 2) It is also called as saturation region or constant current region. Because of the channel is occupied with depletion region, the depletion region is more towards the drain and less towards the source, so the channel is limited, with this only limited number of carriers are only allowed to cross this channel from source drain causing a current that is constant in this region. To use FET as an amplifier it is operated in this saturation region.
- 3) In this drain current remains constant at its maximum value I_{DSS} .
- 4) The drain current in the pinch off region depends upon the gate to source voltage and is given by the relation

$$I_d = I_{dss} [1 - V_{gs}/V_p]^2$$

This is known as shokley's relation.

BREAKDOWN REGION:-

- 1) The region is shown by the curve. In this region, the drain current increases rapidly as the drain to source voltage is increased.
- 2) It is because of the gate to source junction due to avalanche effect.

- 3) The avalanche break down occurs at progressively lower value of V_{DS} because the reverse bias gate voltage adds to the drain voltage thereby increasing effective voltage across the gate junction

This causes

1. The maximum saturation drain current is smaller
 2. The ohmic region portion decreased.
- 4) It is important to note that the maximum voltage V_{DS} which can be applied to FET is the lowest voltage which causes available break down.

2. TRANSFER CHARACTERISTICS:-

These curves shows the relationship between drain current I_D and gate to source voltage V_{GS} for different values of V_{DS} .

- 1) First adjust the drain to source voltage to some suitable value , then increase the gate to source voltage in small suitable value.
- 2) Plot the graph between gate to source voltage along the horizontal axis and current I_D on the vertical axis. We shall obtain a curve like this.

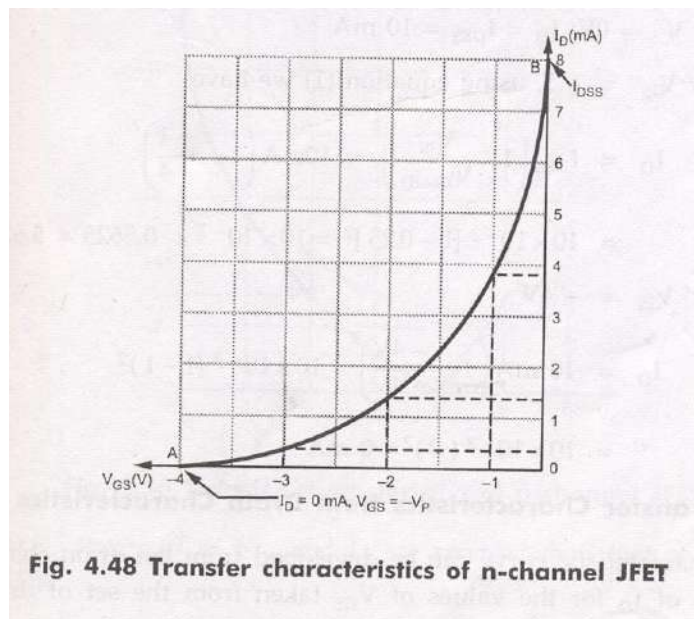


Fig. 4.48 Transfer characteristics of n-channel JFET

- 3) As we know that if V_{GS} is more negative curves drain current to reduce . where V_{GS} is made sufficiently negative, I_d is reduced to zero. This is caused by the widening of the depletion region to a point where it is completely closes the channel. The value of V_{GS} at the cutoff point is designed as V_{GSoff}

- 4) The upper end of the curve as shown by the drain current value is equal to I_{dss} that is when $V_{gs} = 0$ the drain current is maximum.
- 5) While the lower end is indicated by a voltage equal to $V_{gs\text{off}}$
- 6) If V_{gs} continuously increasing, the channel width is reduced, then $I_d = 0$
- 7) It may be noted that curve is part of the parabola; it may be expressed as

$$I_d = I_{dss} [1 - V_{gs}/V_{gs\text{off}}]^2$$

DIFFERENCE BETWEEN V_p AND $V_{gs\text{off}}$ –

V_p is the value of V_{gs} that causes the JFET to become constant current component, It is measured at $V_{gs} = 0V$ and has a constant drain current of $I_d = I_{dss}$. Where $V_{gs\text{off}}$ is the value of V_{gs} that reduces I_d to approximately zero.

Why the gate to source junction of a JFET be always reverse biased ?

The gate to source junction of a JFET is never allowed to become forward biased because the gate material is not designed to handle any significant amount of current. If the junction is allowed to become forward biased, current is generated through the gate material. This current may destroy the component.

There is one more important characteristic of JFET reverse biasing i.e. J FET 's have extremely high characteristic gate input impedance. This impedance is typically in the high mega ohm range. With the advantage of extremely high input impedance it draws no current from the source. The high input impedance of the JFET has led to its extensive use in integrated circuits. The low current requirements of the component makes it perfect for use in ICs. Where thousands of transistors must be etched on to a single piece of silicon. The low current draw helps the IC to remain relatively cool, thus allowing more components to be placed in a smaller physical area.

JFET PARAMETERS

The electrical behavior of JFET may be described in terms of certain parameters. Such parameters are obtained from the characteristic curves.

A C Drain resistance(r_d):

It is also called dynamic drain resistance and is the a.c.resistance between the drain and source terminal,when the JFET is operating in the pinch off or saturation region.It is given by the ratio of small change in drain to source voltage ΔV_{ds} to the corresponding change in drain current ΔI_d for a constant gate to source voltage V_{gs} .

Mathematically it is expressed as $r_d = \Delta V_{ds} / \Delta I_d$ where V_{gs} is held constant.

TRANSCONDUCTANCE (g_m):

It is also called forward transconductance . It is given by the ratio of small change in drain current (ΔI_d) to the corresponding change in gate to source voltage (ΔV_{ds})

Mathematically the transconductance can be written as

$$g_m = \Delta I_d / \Delta V_{ds}$$

AMPLIFICATION FACTOR (μ)

It is given by the ratio of small change in drain to source voltage (ΔV_{ds}) to the corresponding change in gate to source voltage (ΔV_{gs}) for a constant drain current (I_d).

Thus $\mu = \Delta V_{ds} / \Delta V_{gs}$ when I_d held constant

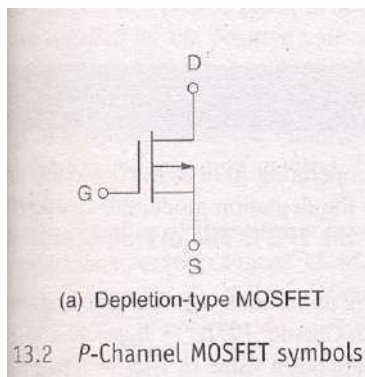
The amplification factor μ may be expressed as a product of transconductance (g_m) and ac drain resistance (r_d)

$$\mu = \Delta V_{ds} / \Delta V_{gs} = g_m r_d$$

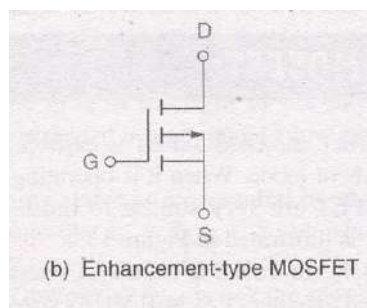
MOSFET

We now turn our attention to the insulated gate FET or metal oxide semi conductor FET which is having the greater commercial importance than the junction FET.

Most MOSFETS however are triodes, with the substrate internally connected to the source. The circuit symbols used by several manufacturers are indicated in the Fig below.



13.2 P-Channel MOSFET symbols.



(a) Depletion type MOSFET

(b) Enhancement type MOSFET

Both of them are P- channel

Here are two basic types of MOSFETS

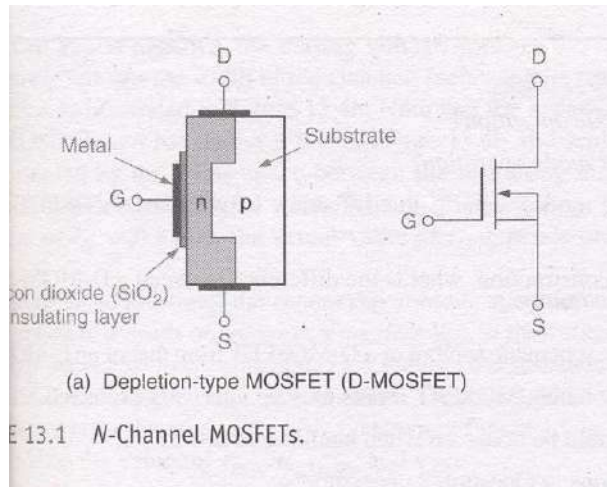
- (1) Depletion type
- (2) Enhancement type MOSFET.

MOSFETS can be operated in both the depletion mode and the enhancement mode.

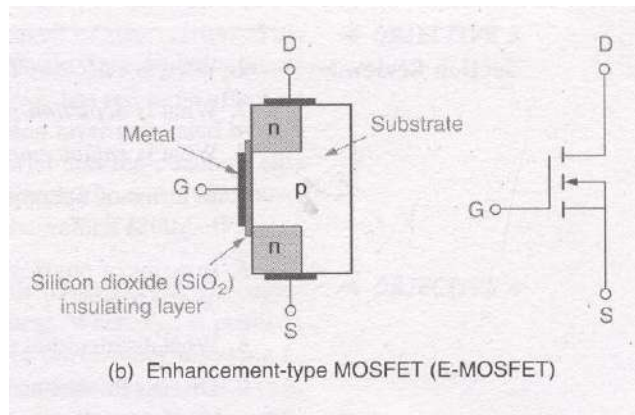
E -MOSFETS

are restricted to operate in enhancement mode. The primary difference between them is their physical construction.

The construction difference between the two is shown in the fig given below.



As we can see the D MOSFET have physical channel between the source and drain terminals(Shaded area)



The E MOSFET on the other hand has no such channel physically. It depends on the gate voltage to form a channel between the source and the drain terminals.

Both MOSFETS have an insulating layer between the gate and the rest of the component. This insulating layer is made up of SiO_2 a glass like insulating material. The gate material is made up of

metal conductor. Thus going from gate to substrate, we can have metal oxide semiconductor which is where the term MOSFET comes from.

Since the gate is insulated from the rest of the component, the MOSFET is sometimes referred to as an insulated gate FET or IGFET.

The foundation of the MOSFET is called the substrate. This material is represented in the schematic symbol by the center line that is connected to the source.

In the symbol for the MOSFET, the arrow is placed on the substrate. As with JFET an arrow pointing in represents an N-channel device, while an arrow pointing out represents p-channel device.

CONSTRUCTION OF AN N-CHANNEL MOSFET:-

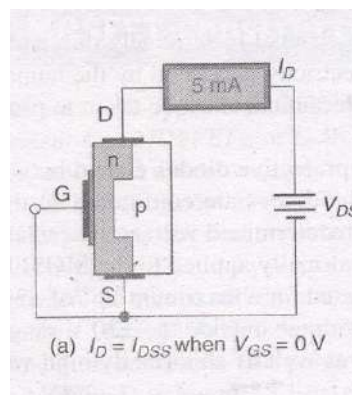
The N-channel MOSFET consists of a lightly doped p type substance into which two heavily doped n+ regions are diffused as shown in the Fig. These n+ sections, which will act as source and drain.

A thin layer of insulation silicon dioxide (SiO_2) is grown over the surface of the structure, and holes are cut into oxide layer, allowing contact with the source and drain. Then the gate metal area is overlaid on the oxide, covering the entire channel region. Metal contacts are made to drain and source and the contact to the metal over the channel area is the gate terminal. The metal area of the gate, in conjunction with the insulating dielectric oxide layer and the semiconductor channel, forms a parallel plate capacitor. The insulating layer of SiO_2

is the reason why this device is called the insulated gate field effect transistor. This layer results in an extremely high input resistance (10^{10} to 10^{15} ohms) for MOSFET.

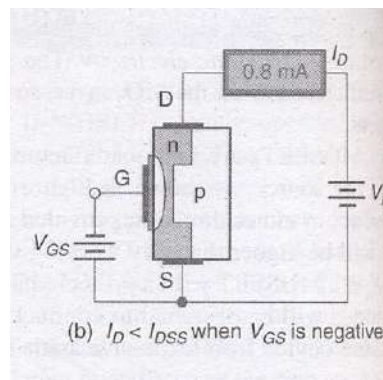
DEPLETION MOSFET

The basic structure of D-MOSFET is shown in the fig. An N-channel is diffused between source and drain with the device an appreciable drain current I_{DSS} flows for zero gate to source voltage, $V_{GS}=0$.



Depletion mode operation:-

- 1) The above fig shows the D-MOSFET operating conditions with gate and source terminals shorted together ($V_{GS}=0V$)
- 2) At this stage $I_D = I_{DSS}$ where $V_{GS}=0V$, with this voltage V_{DS} , an appreciable drain current I_{DSS} flows.
- 3) If the gate to source voltage is made negative i.e. V_{GS} is negative. Positive charges are induced in the channel through the SiO_2 of the gate capacitor.
- 4) Since the current in a FET is due to majority carriers (electrons for an N-type material), the induced positive charges make the channel less conductive and the drain current drops as V_{GS} made more negative.
- 5) The re distribution of charge in the channel causes an effective depletion of majority carriers, which accounts for the designation depletion MOSFET.
- 6) That means biasing voltage V_{GS} depletes the channel of free carriers. This effectively reduces the width of the channel, increasing its resistance.
- 7) Note that negative V_{GS} has the same effect on the MOSFET as it has on the JFET.

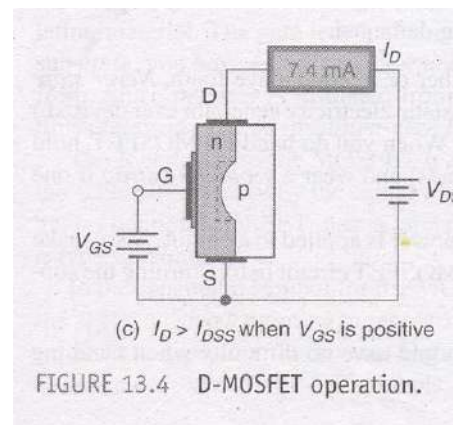


- 8) As shown in the fig above, the depletion layer generated by V_{GS} (represented by the white space between the insulating material and the channel) cuts into the channel, reducing its width. As a result, $I_D < I_{DSS}$. The actual value of I_D depends on the value of I_{DSS} , $V_{GS}(\text{off})$ and V_{GS} .

Enhancement mode operation of the D-MOSFET:-

- 1) This operating mode is a result of applying a positive gate to source voltage V_{GS} to the device.
- 2) When V_{GS} is positive the channel is effectively widened. This reduces the resistance of the channel allowing I_D to exceed the value of I_{DSS}
- 3) When V_{GS} is given positive the majority carriers in the p-type are holes. The holes in the p type substrate are repelled by the +ve gate voltage.

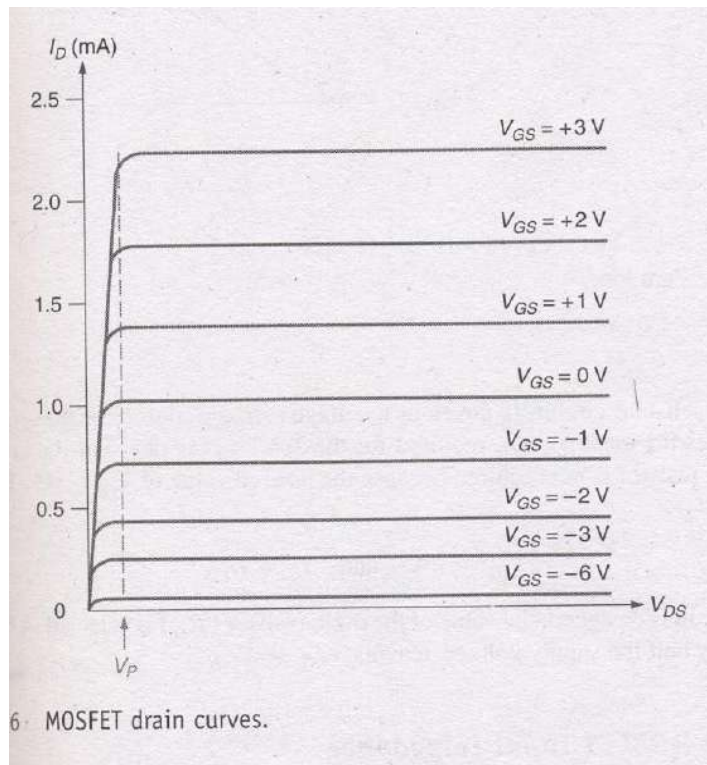
- 4) At the same time, the conduction band electrons (minority carriers) in the p type material are attracted towards the channel by the +gate voltage.
- 5) With the build up of electrons near the channel, the area to the right of the physical channel effectively becomes an N type material.
- 6) The extended n type channel now allows more current, $I_D > I_{DSS}$



Characteristics of Depletion MOSFET:-

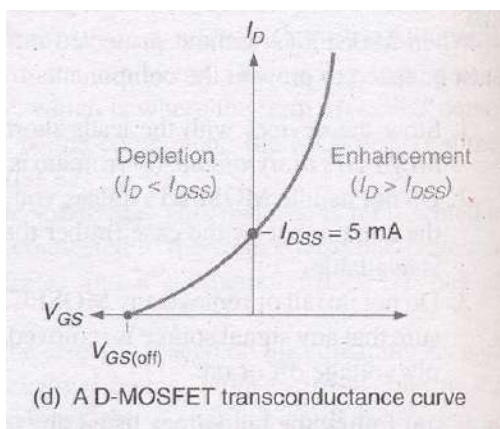
The fig. shows the drain characteristics for the N channel depletion type MOSFET

- 1) The curves are plotted for both V_{GS} positive and V_{GS} negative voltages
- 2) When $V_{GS}=0$ and negative the MOSFET operates in depletion mode when V_{GS} is positive, the MOSFET operates in the enhancement mode.
- 3) The difference between JFET and D MOSFET is that JFET does not operate for positive values of V_{GS} .
- 4) When $V_{DS}=0$, there is no conduction takes place between source to drain, if $V_{GS}<0$ and $V_{DS}>0$ then I_D increases linearly.
- 5) But as $V_{GS}, 0$ induces positive charges holes in the channel, and controls the channel width. Thus the conduction between source to drain is maintained as constant, i.e. I_D is constant.
- 6) If $V_{GS}>0$ the gate induces more electrons in channel side, it is added with the free electrons generated by source. again the potential applied to gate determines the channel width and maintains constant current flow through it as shown in Fig



TRANSFER CHARACTERISTICS:-

The combination of 3 operating states i.e. $V_{GS}=0\text{V}$, $V_{GS}<0\text{V}$, $V_{GS}>0\text{V}$ is represented by the D MOSFET transconductance curve shown in Fig.

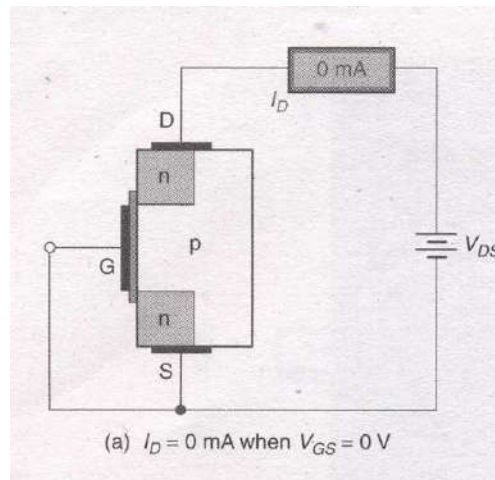


- 1) Here in this curve it may be noted that the region AB of the characteristics similar to that of JFET.
- 2) This curve extends for the positive values of V_{GS}

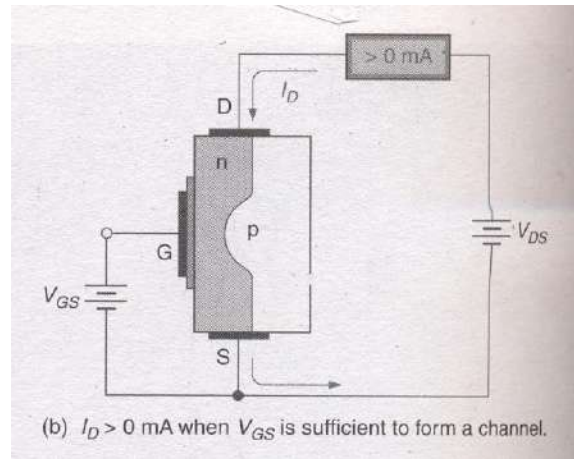
- 3) Note that $I_D = I_{DSS}$ for $V_{GS} = 0V$ when V_{GS} is negative, $I_D < I_{DSS}$ when $V_{GS} = V_{GS(off)}$, I_D is reduced to approximately $0mA$. Where V_{GS} is positive $I_D > I_{DSS}$. So obviously I_{DSS} is not the maximum possible value of I_D for a MOSFET.
- 4) The curves are similar to JFET so that the D MOSFET have the same transconductance equation.

E-MOSFETS:-

The E MOSFET is capable of operating only in the enhancement mode. The gate potential must be positive w.r.t to source.



- 1) when the value of $V_{GS} = 0V$, there is no channel connecting the source and drain materials.
- 2) As a result, there can be no significant amount of drain current.
- 3) When $V_{GS} = 0$, the V_{DD} supply tries to force free electrons from source to drain but the presence of p-region does not permit the electrons to pass through it. Thus there is no drain current at $V_{GS} = 0$,
- 4) If V_{GS} is positive, it induces a negative charge in the p type substrate just adjacent to the SiO_2 layer.
- 5) As the holes are repelled by the positive gate voltage, the minority carrier electrons attracted toward this voltage. This forms an effective N type bridge between source and drain providing a path for drain current.
- 6) This +ve gate voltage forms a channel between the source and drain.
- 7) This produces a thin layer of N type channel in the P type substrate. This layer of free electrons is called N type inversion layer.



- 8) The minimum V_{GS} which produces this inversion layer is called threshold voltage and is designated by $V_{GS(th)}$. This is the point at which the device turns on is called the threshold voltage $V_{GS(th)}$
- 9) When the voltage V_{GS} is $< V_{GS(th)}$ no current flows from drain to source.
- 10) However when the voltage $V_{GS} > V_{GS(th)}$ the inversion layer connects the drain to source and we get significant values of current.

CHARACTERISTICS OF E- MOSFET:-

1. DRAIN CHARACTERISTICS

The volt ampere drain characteristics of an N-channel enhancement mode MOSFET are given in the

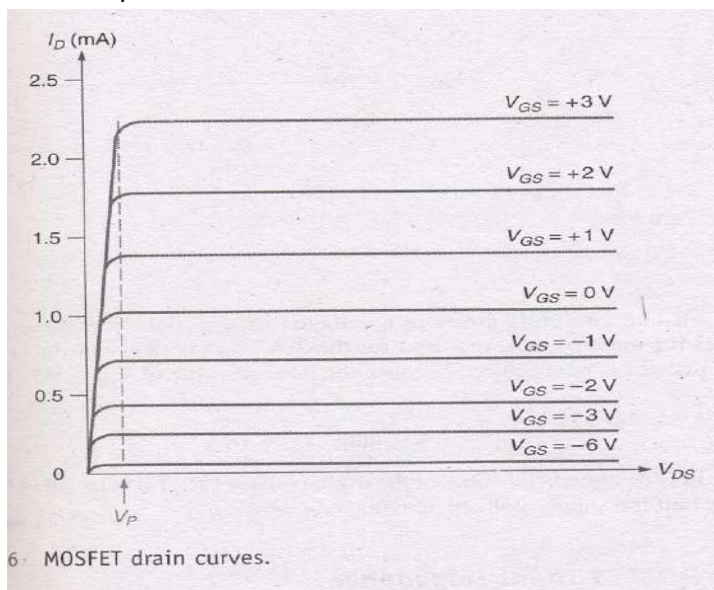


fig.

2. TRANSFER CHARACTERISTICS:-

- 1) The current I_{DSS} at $V_{GS} \leq 0$ is very small being of the order of a few nano amps.
- 2) As V_{GS} is made +ve, the current I_D increases slowly at first, and then much more rapidly with an increase in V_{GS} .
- 3) The standard transconductance formula will not work for the E MOSFET.
- 4) To determine the value of I_D at a given value of V_{GS} we must use the following relation

$$I_D = K[V_{GS} - V_{GS(Th)}]^2$$

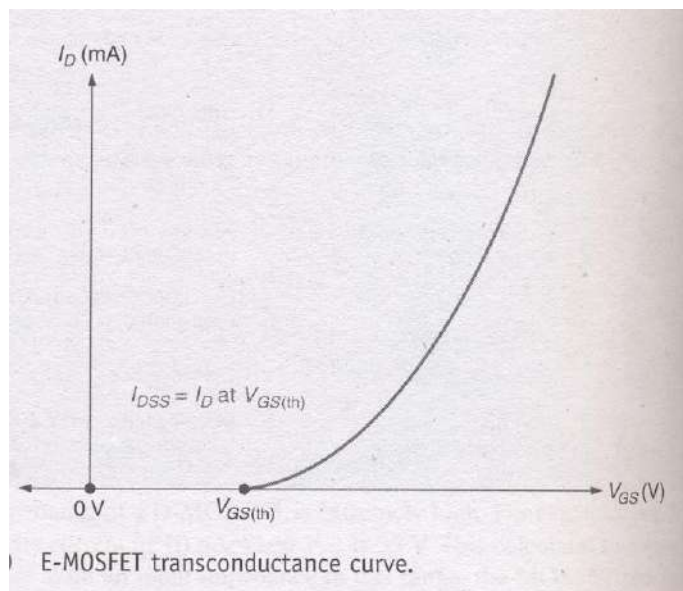
Where K is constant for the MOSFET . found as

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

From the data specification sheets, the 2N7000 has the following ratings.

$I_{D(on)} = 75\text{mA (minimum)}$.

And $V_{GS(th)} = 0.8\text{(minimum)}$



APPLICATION OF MOSFET

One of the primary contributions to electronics made by MOSFETs can be found in the area of digital (computer electronics). The signals in digital circuits are made up of rapidly switching dc levels. This signal is called as a rectangular wave, made up of two dc levels (or logic levels). These logic levels are 0V and +5V.

A group of circuits with similar circuitry and operating characteristics is referred to as a logic family. All the circuits in a given logic family respond to the same logic levels, have similar speed and power-handling capabilities , and can be directly connected together. One such logic family is complementary MOS (or CMOS) logic. This logic family is made up entirely of MOSFETs.

UNIT-4

Feedback Amplifier & Oscillator

Feedback Amplifier

A practical amplifier has a gain of nearly one million *i.e.* its output is one million times the input. Consequently, even a casual disturbance at the input will appear in the amplified form in the output. There is a strong tendency in amplifiers to introduce *hum* due to sudden temperature changes or stray electric and magnetic fields. Therefore, every high gain amplifier tends to give noise along with signal in its output. The noise in the output of an amplifier is undesirable and must be kept to as small a level as possible. The noise level in amplifiers can be reduced considerably by the use of *negative feedback i.e.* by injecting a fraction of output in phase opposition to the input signal. The object of this chapter is to consider the effects and methods of providing negative feedback in transistor amplifiers.

Ideally an amplifier should reproduce the input signal, with change in magnitude and with or without change in phase. But some of the short comings of the amplifier circuit are

1. Change in the value of the gain due to variation in supplying voltage, temperature or due to components.
2. Distortion in wave-form due to non linearities in the operating characters of the Amplifying device.
3. The amplifier may introduce noise (undesired signals)

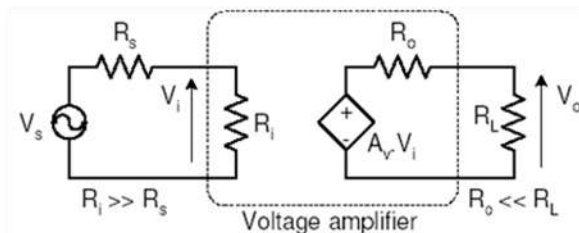
The above drawbacks can be minimizing if we introduce feedback.

CLASSIFICATION OF AMPLIFIERS

Amplifiers can be classified broadly as:

1. Voltage amplifiers.
2. Current amplifiers.
3. Tran conductance amplifiers.
4. Tran resistance amplifiers.

1.1 Voltage amplifier



if $R_i \gg R_s$

then $V_i \approx V_s$

and if $R_o \ll R_L$

then

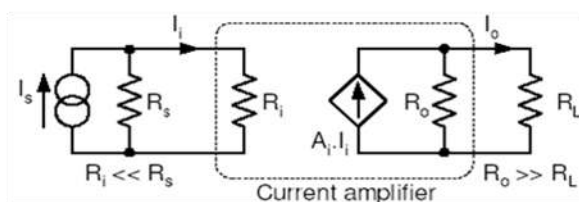
$$V_o \approx A_v V_i \approx A_v V_s$$

hence $A_v \equiv \frac{V_o}{V_i}$

with $R_L = \infty$

represent the open circuit voltage gain.

1.2 Current amplifier



if $R_i \ll R_s$

then $I_i \approx I_s$

and if $R_o \gg R_L$

then

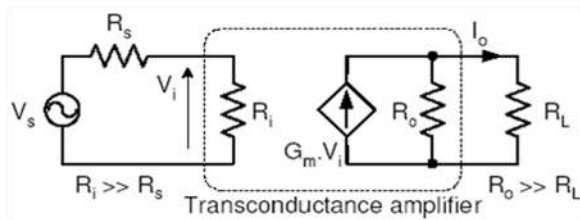
$$I_o \approx A_i I_i \approx A_i I_s$$

hence $A_i \equiv \frac{I_o}{I_i}$

with $R_L = 0$

represent the short circuit current gain.

1.3 Transconductance amplifier



if $R_i \gg R_s$

then $V_i \approx V_s$

and if $R_o \gg R_L$

then

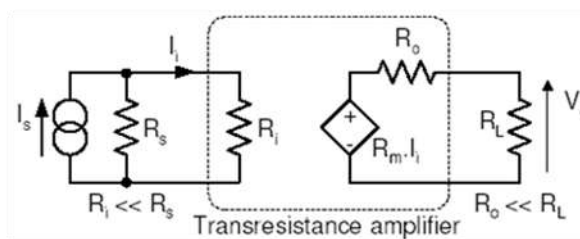
$$I_o \approx G_m V_i \approx G_m V_s$$

hence $G_m \equiv \frac{I_o}{V_i}$

with $R_L = 0$

represent the short circuit mutual or transfer conductance

1.4 Transresistance amplifier



if $R_i \ll R_s$

then $I_i \approx I_s$

and if $R_o \ll R_L$

then

$$V_o \approx R_m I_i \approx R_m I_s$$

hence $R_m \equiv \frac{V_o}{I_i}$

with $R_L = \infty$

represent the open circuit mutual or transfer resistance.

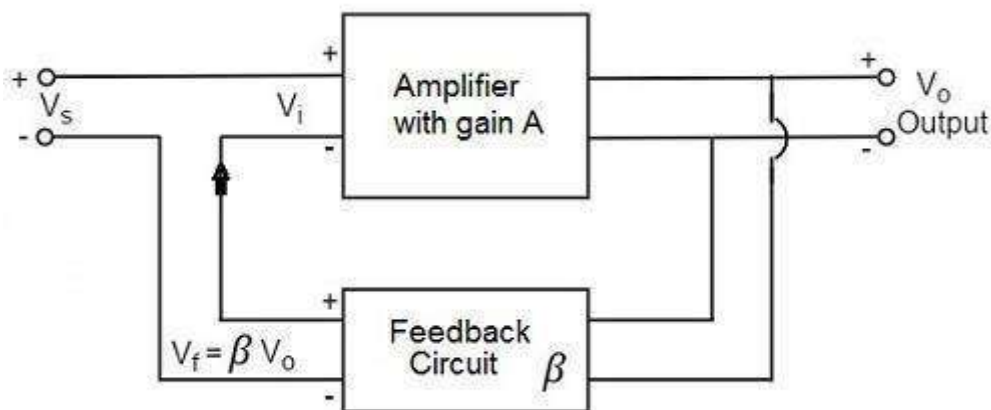
Concept of Feedback

An amplifier circuit simply increases the signal strength. But while amplifying, it just increases the strength of its input signal whether it contains information or some noise along with information. This noise or some disturbance is introduced in the amplifiers because of their strong tendency to introduce hum due to sudden temperature changes or stray electric and magnetic fields. Therefore, every high gain amplifier tends to give noise along with signal in its output, which is very undesirable.

The noise level in the amplifier circuits can be considerably reduced by using negative feedback done by injecting a fraction of output in phase opposition to the input signal.

Principle of Feedback Amplifier

A feedback amplifier generally consists of two parts. They are the amplifier and the feedback circuit. The feedback circuit usually consists of resistors. The concept of feedback amplifier can be understood from the following figure.



From the above figure, the gain of the amplifier is represented as A . the gain of the amplifier is the ratio of output voltage V_o to the input voltage V_i . the feedback network extracts a voltage $V_f = \beta V_o$ from the output V_o of the amplifier.

This voltage is added for positive feedback and subtracted for negative feedback, from the signal voltage V_s . Now,

$$V_i = V_s + V_f = V_s + \beta V_o$$

$$V_i = V_s - V_f = V_s - \beta V_o$$

The quantity $\beta = V_f/V_o$ is called as feedback ratio or feedback fraction.

Let us consider the case of negative feedback. The output V_o must be equal to the input voltage $(V_s - \beta V_o)$ multiplied by the gain A of the amplifier.

Hence,

$$(V_s - \beta V_o)A = V_o$$

Or

$$AV_s - A\beta V_o = V_o$$

Or

$$AV_s = V_o(1 + A\beta)$$

Therefore,

$$\frac{V_o}{V_s} = \frac{A}{1 + A\beta}$$

Let A_f be the overall gain (gain with the feedback) of the amplifier. This is defined as the ratio of output voltage V_o to the applied signal voltage V_s , i.e.,

$$A_f = \frac{A}{1 + A\beta}$$

The equation of gain of the feedback amplifier, with positive feedback is given by

$$A_f = \frac{A}{1 - A\beta}$$

These are the standard equations to calculate the gain of feedback amplifiers.

Types of Feedbacks

The process of injecting a fraction of output energy of some device back to the input is known as Feedback. It has been found that feedback is very useful in reducing noise and making the amplifier operation stable.

Depending upon whether the feedback signal aids or opposes the input signal, there are two types of feedbacks used.

Positive Feedback

The feedback in which the feedback energy i.e., either voltage or current is in phase with the input signal and thus aids it is called as Positive feedback.

Both the input signal and feedback signal introduces a phase shift of 180° thus making a 360° resultant phase shift around the loop, to be finally in phase with the input signal.

Though the positive feedback increases the gain of the amplifier, it has the disadvantages such as

- Increasing distortion
- Instability

It is because of these disadvantages the positive feedback is not recommended for the amplifiers. If the positive feedback is sufficiently large, it leads to oscillations, by which oscillator circuits are formed.

Negative Feedback

The feedback in which the feedback energy i.e., either voltage or current is out of phase with the input and thus opposes it, is called as negative feedback.

In negative feedback, the amplifier introduces a phase shift of 180° into the circuit while the feedback network is so designed that it produces no phase shift or zero phase shift. Thus the resultant feedback voltage V_f is 180° out of phase with the input signal V_{in} .

Though the gain of negative feedback amplifier is reduced, there are many advantages of negative feedback such as

- Stability of gain is improved
- Reduction in distortion
- Reduction in noise
- Increase in input impedance
- Decrease in output impedance
- Increase in the range of uniform application

It is because of these advantages negative feedback is frequently employed in amplifiers.

Negative feedback in an amplifier is the method of feeding a portion of the amplified output to the input but in opposite phase. The phase opposition occurs as the amplifier provides 180° phase shift whereas the feedback network doesn't.

While the output energy is being applied to the input, for the voltage energy to be taken as feedback, the output is taken in shunt connection and for the current energy to be taken as feedback, the output is taken in series connection.

There are two main types of negative feedback circuits. They are –

- Negative Voltage Feedback
- Negative Current Feedback

Negative Voltage Feedback

In this method, the voltage feedback to the input of amplifier is proportional to the output voltage. This is further classified into two types –

- Voltage-series feedback
- Voltage-shunt feedback

Negative Current Feedback

In this method, the voltage feedback to the input of amplifier is proportional to the output current. This is further classified into two types.

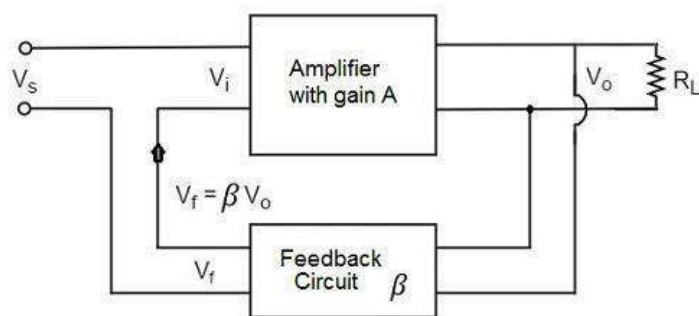
- Current-series feedback
- Current-shunt feedback

Let us have a brief idea on all of them.

Voltage-Series Feedback

In the voltage series feedback circuit, a fraction of the output voltage is applied in series with the input voltage through the feedback circuit. This is also known as shunt-driven series-fed feedback, i.e., a parallel-series circuit.

The following figure shows the block diagram of voltage series feedback, by which it is evident that the feedback circuit is placed in shunt with the output but in series with the input.

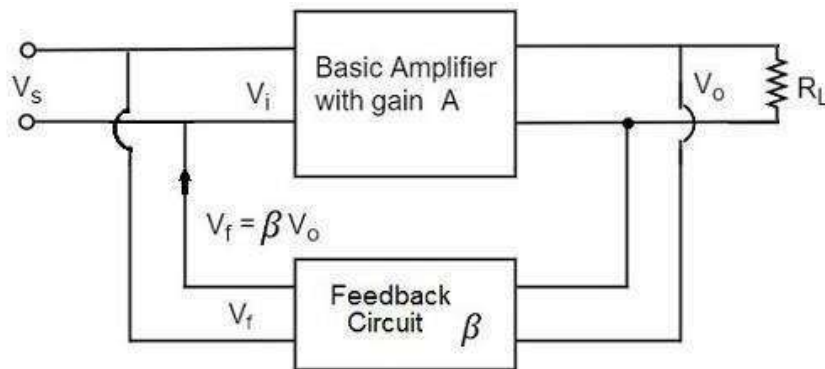


As the feedback circuit is connected in shunt with the output, the output impedance is decreased and due to the series connection with the input, the input impedance is increased.

Voltage-Shunt Feedback

In the voltage shunt feedback circuit, a fraction of the output voltage is applied in parallel with the input voltage through the feedback network. This is also known as shunt-driven shunt-fed feedback i.e., a parallel-parallel proto type.

The below figure shows the block diagram of voltage shunt feedback, by which it is evident that the feedback circuit is placed in shunt with the output and also with the input.

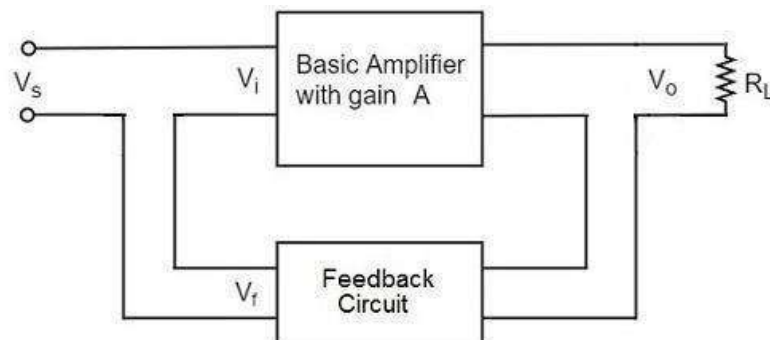


As the feedback circuit is connected in shunt with the output and the input as well, both the output impedance and the input impedance are decreased.

Current-Series Feedback

In the current series feedback circuit, a fraction of the output voltage is applied in series with the input voltage through the feedback circuit. This is also known as series-driven series-fed feedback i.e., a series-series circuit.

The following figure shows the block diagram of current series feedback, by which it is evident that the feedback circuit is placed in series with the output and also with the input.

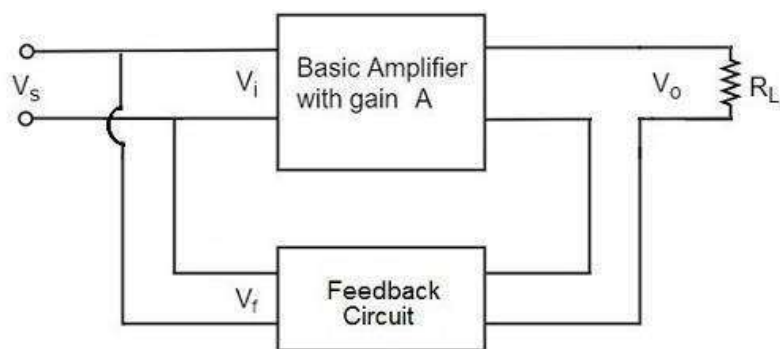


As the feedback circuit is connected in series with the output and the input as well, both the output impedance and the input impedance are increased.

Current-Shunt Feedback

In the current shunt feedback circuit, a fraction of the output voltage is applied in series with the input voltage through the feedback circuit. This is also known as series-driven shunt-fed feedback i.e., a series-parallel circuit.

The below figure shows the block diagram of current shunt feedback, by which it is evident that the feedback circuit is placed in series with the output but in parallel with the input.



As the feedback circuit is connected in series with the output, the output impedance is increased and due to the parallel connection with the input, the input impedance is decreased.

Let us now tabulate the amplifier characteristics that get affected by different types of negative feedbacks.

Characteristics	Types of Feedback			
	Voltage-Series	Voltage-Shunt	Current-Series	Current-Shunt
Voltage Gain	Decreases	Decreases	Decreases	Decreases
Bandwidth	Increases	Increases	Increases	Increases
Input resistance	Increases	Decreases	Increases	Decreases
Output resistance	Decreases	Decreases	Increases	Increases
Harmonic distortion	Decreases	Decreases	Decreases	Decreases
Noise	Decreases	Decreases	Decreases	Decreases

Advantages of Negative Feedback

1. Stabilization of gain
 - Make the gain less sensitive to changes in circuit components e.g. due to changes in temperature.
2. Reduce non-linear distortion
 - Make the output proportional to the input, keeping the gain constant, independent of signal level.
3. Reduce the effect of noise
 - Minimize the contribution to the output of unwanted signals generated in circuit components or extraneous interference.
4. Extend the bandwidth of the amplifier
 - Reduce the gain and increase the bandwidth
5. Modification the input and output impedances
 - Raise or lower the input and output impedances by selection of the appropriate feedback topology.

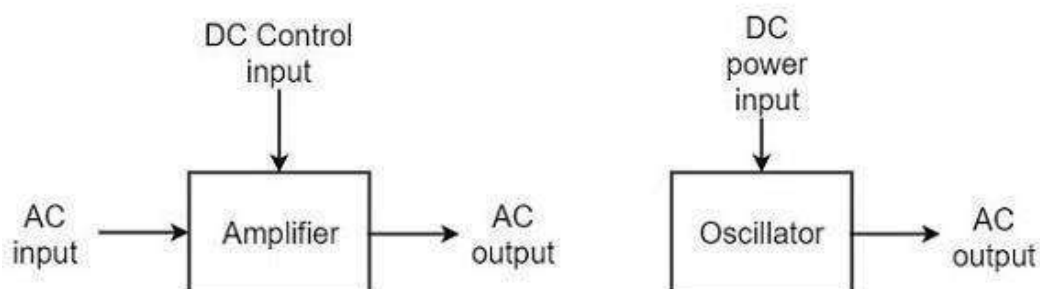
Oscillators

An **oscillator** generates output without any ac input signal. An electronic oscillator is a circuit which converts dc energy into ac at a very high frequency. An amplifier with a positive feedback can be understood as an oscillator.

Amplifier vs. Oscillator

An **amplifier** increases the signal strength of the input signal applied, whereas an **oscillator** generates a signal without that input signal, but it requires dc for its operation. This is the main difference between an amplifier and an oscillator.

Take a look at the following illustration. It clearly shows how an amplifier takes energy from d.c. power source and converts it into a.c. energy at signal frequency. An oscillator produces an oscillating a.c. signal on its own.



The frequency, waveform, and magnitude of a.c. power generated by an amplifier, is controlled by the a.c. signal voltage applied at the input, whereas those for an oscillator are controlled by the components in the circuit itself, which means no external controlling voltage is required.

Alternator vs. Oscillator

An **alternator** is a mechanical device that produces sinusoidal waves without any input. This a.c. generating machine is used to generate frequencies up to 1000Hz. The output frequency depends on the number of poles and the speed of rotation of the armature.

The following points highlight the differences between an alternator and an oscillator –

- An alternator converts mechanical energy to a.c. energy, whereas the oscillator converts d.c. energy into a.c. energy.
- An oscillator can produce higher frequencies of several MHz whereas an alternator cannot.
- An alternator has rotating parts, whereas an electronic oscillator doesn't.
- It is easy to change the frequency of oscillations in an oscillator than in an alternator.

Oscillators can also be considered as opposite to rectifiers that convert a.c. to d.c. as these convert d.c. to a.c.

Classification of Oscillators

Electronic oscillators are classified mainly into the following two categories –

- **Sinusoidal Oscillators** – The oscillators that produce an output having a sine waveform are called **sinusoidal** or **harmonic oscillators**. Such oscillators can provide output at frequencies ranging from 20 Hz to 1 GHz.
- **Non-sinusoidal Oscillators** – The oscillators that produce an output having a square, rectangular or saw-tooth waveform are called **non-sinusoidal** or **relaxation oscillators**. Such oscillators can provide output at frequencies ranging from 0 Hz to 20 MHz.

Sinusoidal Oscillators

Sinusoidal oscillators can be classified in the following categories –

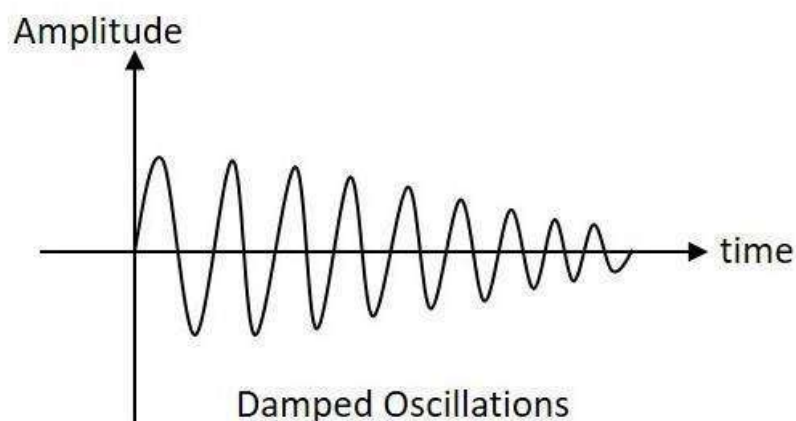
- **Tuned Circuit Oscillators** – These oscillators use a tuned-circuit consisting of inductors (L) and capacitors (C) and are used to generate high-frequency signals. Thus they are also known as radio frequency R.F. oscillators. Such oscillators are Hartley, Colpitts, Clapp-oscillators etc.
- **RC Oscillators** – These oscillators use resistors and capacitors and are used to generate low or audio-frequency signals. Thus they are also known as audio-frequency (A.F.) oscillators. Such oscillators are Phase –shift and Wein-bridge oscillators.
- **Crystal Oscillators** – These oscillators use quartz crystals and are used to generate highly stabilized output signal with frequencies up to 10 MHz. The Piezo oscillator is an example of a crystal oscillator.
- **Negative-resistance Oscillator** – These oscillators use negative-resistance characteristic of the devices such as tunnel devices. A tuned diode oscillator is an example of a negative-resistance oscillator.

Nature of Sinusoidal Oscillations

The nature of oscillations in a sinusoidal wave is generally of two types. They are **damped** and **undamped oscillations**.

Damped Oscillations

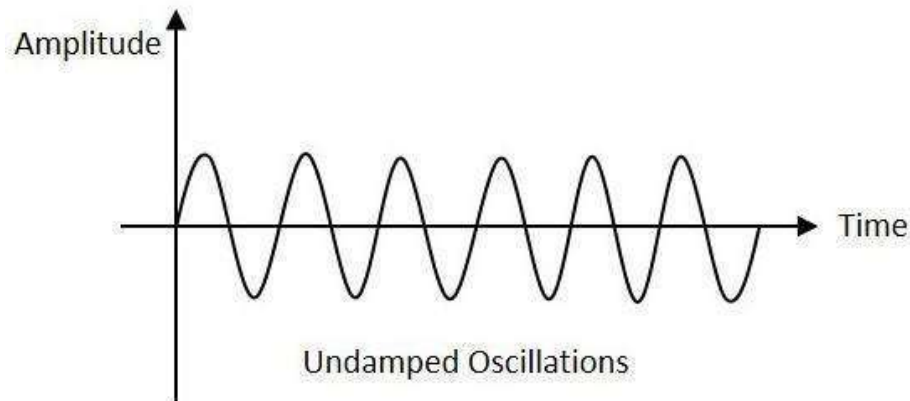
The electrical oscillations whose amplitude goes on decreasing with time are called as **Damped Oscillations**. The frequency of the damped oscillations may remain constant depending upon the circuit parameters.



Damped oscillations are generally produced by the oscillatory circuits that produce power losses and doesn't compensate if required.

Undamped Oscillations

The electrical oscillations whose amplitude remains constant with time are called as **Undamped Oscillations**. The frequency of the undamped oscillations remains constant.



Undamped oscillations are generally produced by the oscillatory circuits that produce no power losses and follow compensation techniques if any power losses occur.

An amplifier with positive feedback produces its output to be in phase with the input and increases the strength of the signal. Positive feedback is also called as **degenerative feedback** or **direct feedback**. This kind of feedback makes a feedback amplifier, an oscillator.

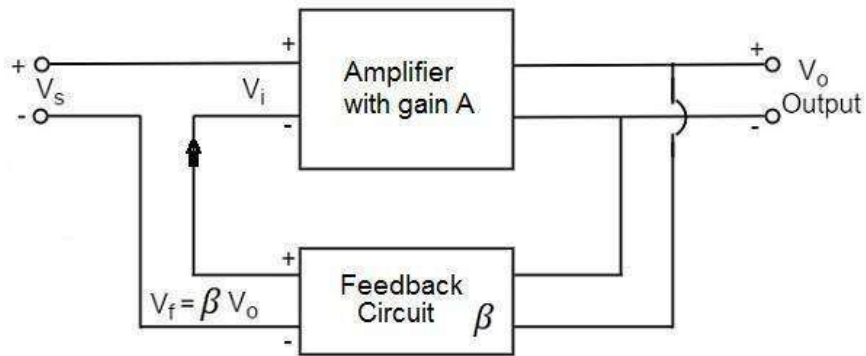
The use of positive feedback results in a feedback amplifier having closed-loop gain greater than the open-loop gain. It results in **instability** and operates as an oscillatory circuit. An oscillatory circuit provides a constantly varying amplified output signal of any desired frequency.

The Barkhausen Criterion

With the knowledge we have till now, we understood that a practical oscillator circuit consists of a tank circuit, a transistor amplifier circuit and a feedback circuit. so, let us now try to brush up the concept of feedback amplifiers, to derive the gain of the feedback amplifiers.

Principle of Feedback Amplifier

A feedback amplifier generally consists of two parts. They are the **amplifier** and the **feedback circuit**. The feedback circuit usually consists of resistors. The concept of feedback amplifier can be understood from the following figure below.



From the above figure, the gain of the amplifier is represented as A . The gain of the amplifier is the ratio of output voltage V_o to the input voltage V_i . The feedback network extracts a voltage $V_f = \beta V_o$ from the output V_o of the amplifier.

This voltage is added for positive feedback and subtracted for negative feedback, from the signal voltage V_s .

So, for a positive feedback,

$$V_i = V_s + V_f = V_s + \beta V_o$$

The quantity $\beta = V_f/V_o$ is called as feedback ratio or feedback fraction.

The output V_o must be equal to the input voltage $(V_s + \beta V_o)$ multiplied by the gain A of the amplifier.

Hence,

$$(V_s + \beta V_o)A = V_o$$

Or

$$AV_s + A\beta V_o = V_o$$

Or

$$AV_s = V_o(1 - A\beta)$$

Therefore

$$\frac{V_o}{V_s} = \frac{A}{(1 - A\beta)}$$

Let A_f be the overall gain (gain with the feedback) of the amplifier. This is defined as the ratio of output voltage V_o to the applied signal voltage V_s , i.e.,

—

from the above two equations, we can understand that, the equation of gain of the feedback amplifier with positive feedback is given by

$$A_f = \frac{A}{1 - A\beta}$$

Where $A\beta$ is the **feedback factor** or the **loop gain**.

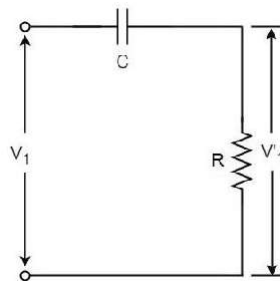
If $A\beta = 1$, $A_f = \infty$. Thus the gain becomes infinity, i.e., there is output without any input. In another words, the amplifier works as an Oscillator.

The condition $A\beta = 1$ is called as **Barkhausen Criterion of oscillations**. This is a very important factor to be always kept in mind, in the concept of Oscillators

RC-Phase-shift Oscillators

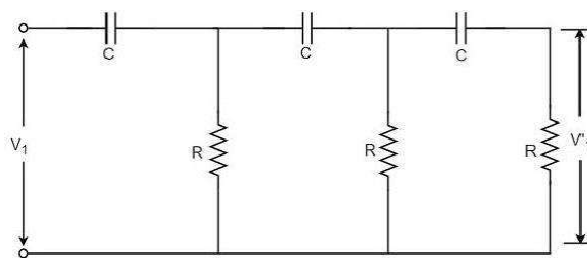
Principle of Phase-shift oscillators

We know that the output voltage of an RC circuit for a sinewave input leads the input voltage. The phase angle by which it leads is determined by the value of RC components used in the circuit. The following circuit diagram shows a single section of an RC network.



The output voltage V_1' across the resistor R leads the input voltage applied input V_1 by some phase angle ϕ° . If R were reduced to zero, V_1' will lead the V_1 by 90° i.e., $\phi^\circ = 90^\circ$.

However, adjusting R to zero would be impracticable, because it would lead to no voltage across R. Therefore, in practice, R is varied to such a value that makes V_1' to lead V_1 by 60° . The following circuit diagram shows the three sections of the RC network.



Each section produces a phase shift of 60° . Consequently, a total phase shift of 180° is produced, i.e., voltage V_2 leads the voltage V_1 by 180° .

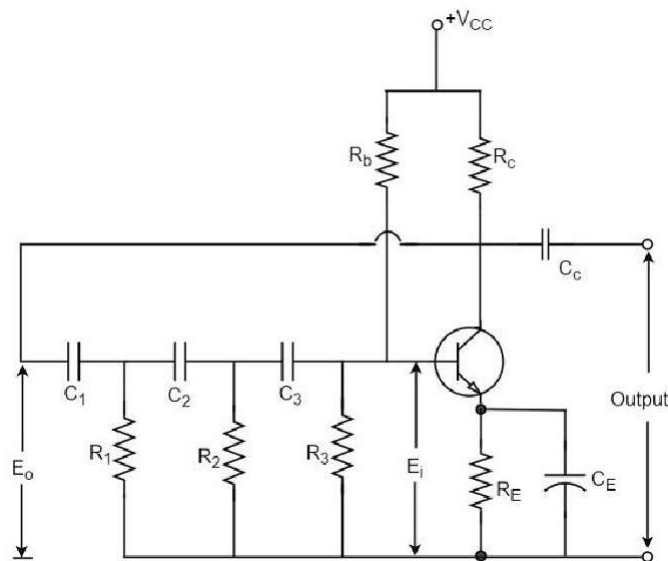
Phase-shift Oscillator Circuit

The oscillator circuit that produces a sine wave using a phase-shift network is called as a Phase-shift oscillator circuit. The constructional details and operation of a phase-shift oscillator circuit are as given below.

Construction

The phase-shift oscillator circuit consists of a single transistor amplifier section and a RC phase-shift network. The phase shift network in this circuit, consists of three RC sections. At the resonant frequency f_o , the phase shift in each RC section is 60° so that the total phase shift produced by RC network is 180° .

The following circuit diagram shows the arrangement of an RC phase-shift oscillator.



The frequency of oscillations is given by

$$f_o = \frac{1}{2\pi RC\sqrt{6}}$$

Where

$$R_1 = R_2 = R_3 = R$$

$$C_1 = C_2 = C_3 = C$$

Operation

The circuit when switched ON oscillates at the resonant frequency f_o . The output E_o of the amplifier is fed back to RC feedback network. This network produces a phase shift of 180° and a voltage E_i appears at its output. This voltage is applied to the transistor amplifier.

The feedback applied will be

$$m = E_i/E_o$$

The feedback is in correct phase, whereas the transistor amplifier, which is in CE configuration, produces a 180° phase shift. The phase shift produced by network and the transistor add to form a phase shift around the entire loop which is 360° .

Advantages

The advantages of RC phase shift oscillator are as follows –

- It does not require transformers or inductors.
- It can be used to produce very low frequencies.
- The circuit provides good frequency stability.

Disadvantages

The disadvantages of RC phase shift oscillator are as follows –

- Starting the oscillations is difficult as the feedback is small.
- The output produced is small.

Another type of popular audio frequency oscillator is the Wien bridge oscillator circuit. This is mostly used because of its important features. This circuit is free from the **circuit fluctuations** and the **ambient temperature**.

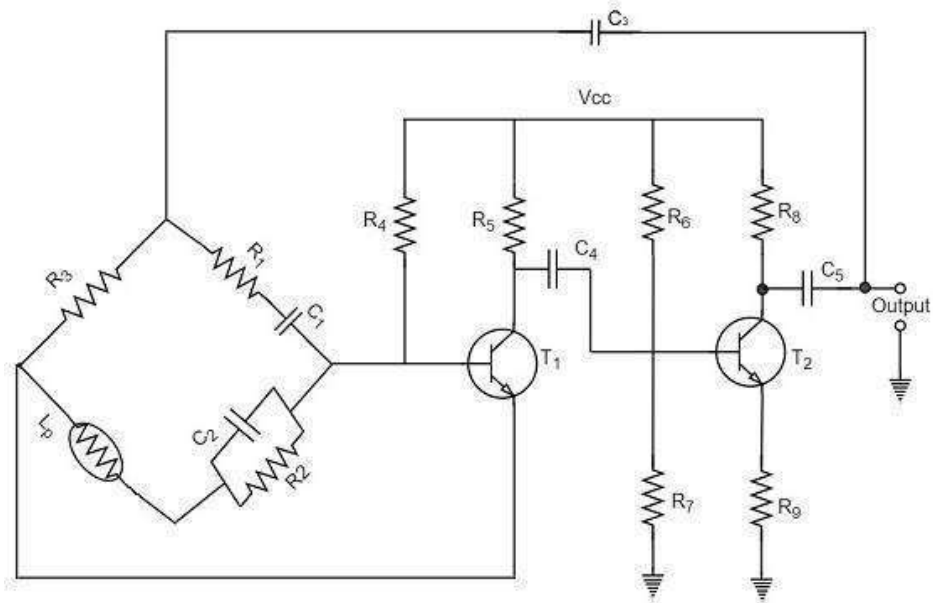
The main advantage of this oscillator is that the frequency can be varied in the range of 10Hz to about 1MHz whereas in RC oscillators, the frequency is not varied.

Wien bridge oscillator

Construction

The circuit construction of Wien bridge oscillator can be explained as below. It is a two-stage amplifier with RC bridge circuit. The bridge circuit has the arms R_1C_1 , R_3 , R_2C_2 and the tungsten lamp L_p . Resistance R_3 and the lamp L_p are used to stabilize the amplitude of the output.

The following circuit diagram shows the arrangement of a Wien bridge oscillator.



The transistor T_1 serves as an oscillator and an amplifier while the other transistor T_2 serves as an inverter. The inverter operation provides a phase shift of 180° . This circuit provides positive feedback through R_1C_1 , C_2R_2 to the transistor T_1 and negative feedback through the voltage divider to the input of transistor T_2 .

The frequency of oscillations is determined by the series element R_1C_1 and parallel element R_2C_2 of the bridge.

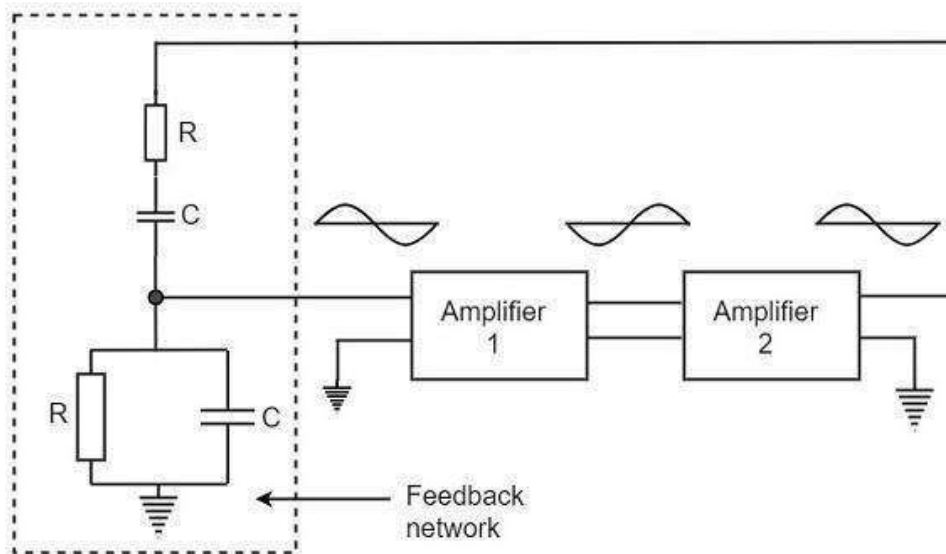
$$f = \frac{1}{2\pi\sqrt{R_1C_1R_2C_2}}$$

If $R_1 = R_2$ and $C_1 = C_2 = C$

Then,

$$f = \frac{1}{2\pi RC}$$

Now, we can simplify the above circuit as follows –



The oscillator consists of two stages of RC coupled amplifier and a feedback network. The voltage across the parallel combination of R and C is fed to the input of amplifier 1. The net phase shift through the two amplifiers is zero.

The usual idea of connecting the output of amplifier 2 to amplifier 1 to provide signal regeneration for oscillator is not applicable here as the amplifier 1 will amplify signals over a wide range of frequencies and hence direct coupling would result in poor frequency stability. By adding Wien bridge feedback network, the oscillator becomes sensitive to a particular frequency and hence frequency stability is achieved.

Operation

When the circuit is switched ON, the bridge circuit produces oscillations of the frequency stated above. The two transistors produce a total phase shift of 360° so that proper positive feedback is ensured. The negative feedback in the circuit ensures constant output. This is achieved by temperature sensitive tungsten lamp L_p . Its resistance increases with current.

If the amplitude of the output increases, more current is produced and more negative feedback is achieved. Due to this, the output would return to the original value. Whereas, if the output tends to decrease, reverse action would take place.

Advantages

The advantages of Wien bridge oscillator are as follows –

- The circuit provides good frequency stability.

- It provides constant output.
- The operation of circuit is quite easy.
- The overall gain is high because of two transistors.
- The frequency of oscillations can be changed easily.
- The amplitude stability of the output voltage can be maintained more accurately, by replacing R_2 with a thermistor.

Disadvantages

The disadvantages of Wien bridge oscillator are as follows –

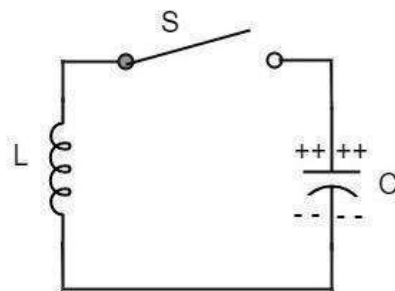
- The circuit cannot generate very high frequencies.
- Two transistors and number of components are required for the circuit construction.

LC Oscillators

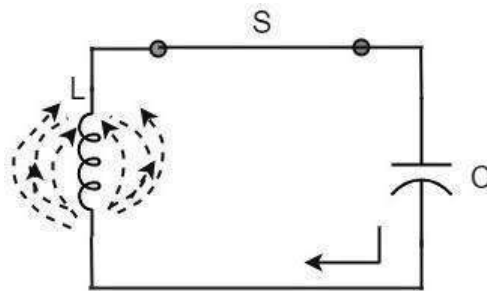
An oscillatory circuit produces electrical oscillations of a desired frequency. They are also known as **tank circuits**.

A simple tank circuit comprises of an inductor L and a capacitor C both of which together determine the oscillatory frequency of the circuit.

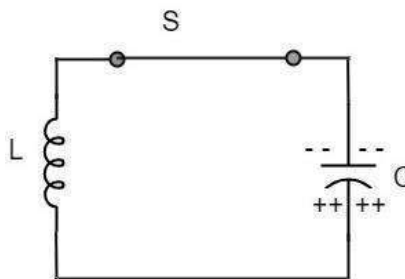
To understand the concept of oscillatory circuit, let us consider the following circuit. The capacitor in this circuit is already charged using a dc source. In this situation, the upper plate of the capacitor has excess of electrons whereas the lower plate has deficit of electrons. The capacitor holds some electrostatic energy and there is a voltage across the capacitor.



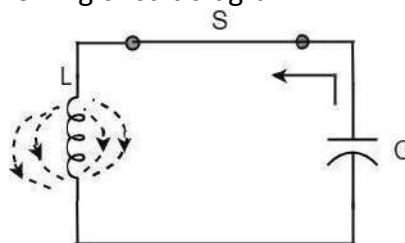
When the switch **S** is closed, the capacitor discharges and the current flows through the inductor. Due to the inductive effect, the current builds up slowly towards a maximum value. Once the capacitor discharges completely, the magnetic field around the coil is maximum.



Now, let us move on to the next stage. Once the capacitor is discharged completely, the magnetic field begins to collapse and produces a counter EMF according to Lenz's law. The capacitor is now charged with positive charge on the upper plate and negative charge on the lower plate.



Once the capacitor is fully charged, it starts to discharge to build up a magnetic field around the coil, as shown in the following circuit diagram.



This continuation of charging and discharging results in alternating motion of electrons or an **oscillatory current**. The interchange of energy between L and C produce continuous **oscillations**.

In an ideal circuit, where there are no losses, the oscillations would continue indefinitely. In a practical tank circuit, there occur losses such as **resistive** and **radiation losses** in the coil and **dielectric losses** in the capacitor. These losses result in damped oscillations.

Frequency of Oscillations

The frequency of the oscillations produced by the tank circuit are determined by the components of the tank circuit, **the L** and **the C**. The actual frequency of oscillations is the **resonant frequency** (or natural frequency) of the tank circuit which is given by

$$f_r = \frac{1}{2\pi\sqrt{LC}}$$

Capacitance of the capacitor

The frequency of oscillation f_o is inversely proportional to the square root of the capacitance of a capacitor. So, if the value of the capacitor used is large, the charge and discharge time periods will be large. Hence the frequency will be lower.

Mathematically, the frequency,

$$f_o \propto \frac{1}{\sqrt{C}}$$

Self-Inductance of the coil

The frequency of the oscillation f_o is proportional to the square root of the self-inductance of the coil. If the value of the inductance is large, the opposition to change of current flow is greater and hence the time required to complete each cycle will be longer, which means time period will be longer and frequency will be lower.

Mathematically, the frequency,

$$f_o \propto \frac{1}{\sqrt{L}}$$

Combining both the above equations,

$$f_o \propto \frac{1}{\sqrt{LC}}$$

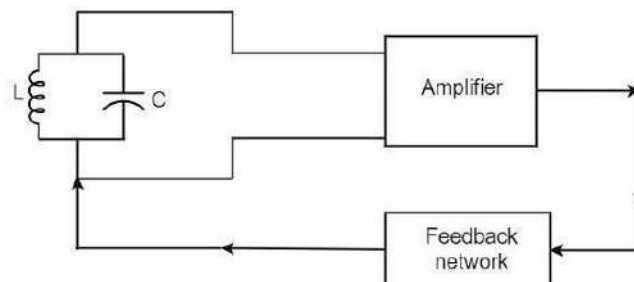
$$f_o = \frac{1}{2\pi\sqrt{LC}}$$

The above equation, though indicates the output frequency, matches the **natural frequency** or **resonance frequency** of the tank circuit.

An Oscillator circuit is a complete set of all the parts of circuit which helps to produce the oscillations. These oscillations should sustain and should be Undamped as just discussed before. Let us try to analyze a practical Oscillator circuit to have a better understanding on how an Oscillator circuit works.

Practical Oscillator Circuit

A Practical Oscillator circuit consists of a tank circuit, a transistor amplifier, and a feedback circuit. The following circuit diagram shows the arrangement of a practical oscillator.



Let us now discuss the parts of this practical oscillator circuit.

Tank Circuit – The tank circuit consists of an inductance L connected in parallel with capacitor C . The values of these two components determine the frequency of the oscillator circuit and hence this is called as **Frequency determining circuit**.

- **Transistor Amplifier** – The output of the tank circuit is connected to the amplifier circuit so that the oscillations produced by the tank circuit are amplified here. Hence the output of these oscillations are increased by the amplifier.
- **Feedback Circuit** – The function of feedback circuit is to transfer a part of the output energy to LC circuit in proper phase. This feedback is positive in oscillators while negative in amplifiers.

Frequency Stability of an Oscillator

The frequency stability of an oscillator is a measure of its ability to maintain a constant frequency, over a long time interval. When operated over a longer period of time, the oscillator frequency may have a drift from the previously set value either by increasing or by decreasing.

The change in oscillator frequency may arise due to the following factors –

- Operating point of the active device such as BJT or FET used should lie in the linear region of the amplifier. Its deviation will affect the oscillator frequency.
- The temperature dependency of the performance of circuit components affect the oscillator frequency.
- The changes in d.c. supply voltage applied to the active device, shift the oscillator frequency. This can be avoided if a regulated power supply is used.
- A change in output load may cause a change in the Q-factor of the tank circuit, thereby causing a change in oscillator output frequency.
- The presence of inter element capacitances and stray capacitances affect the oscillator output frequency and thus frequency stability.

Tuned circuit oscillators are the circuits that produce oscillations with the help of tuning circuits. The tuning circuits consists of an inductance L and a capacitor C . These are also known as **LC oscillators, resonant circuit oscillators or tank circuit oscillators**.

The tuned circuit oscillators are used to produce an output with frequencies ranging from 1 MHz to 500 MHz. Hence these are also known as **R.F. Oscillators**. A BJT or a FET is used as an amplifier with tuned circuit oscillators. With an amplifier and an LC tank circuit, we can feedback a signal with right amplitude and phase to maintain oscillations.

Types of Tuned Circuit Oscillators

Most of the oscillators used in radio transmitters and receivers are of LC oscillators type. Depending upon the way the feedback is used in the circuit, the LC oscillators are divided as the following types.

- **Hartley Oscillator** – It uses inductive feedback.
- **Colpitts Oscillator** – It uses capacitive feedback.

Hartley Oscillator

A very popular **local oscillator** circuit that is mostly used in **radio receivers** is the **Hartley Oscillator** circuit. The constructional details and operation of a Hartley oscillator are as discussed below.

Construction

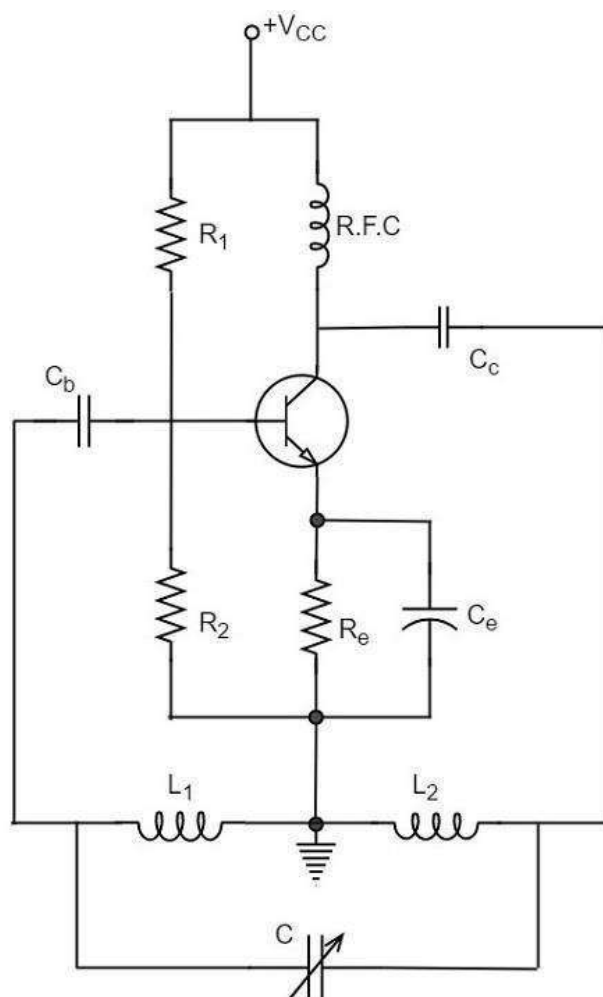
In the circuit diagram of a Hartley oscillator shown below, the resistors R_1 , R_2 and R_e provide necessary bias condition for the circuit. The capacitor C_e provides a.c. ground thereby providing any signal degeneration. This also provides temperature stabilization.

The capacitors C_c and C_b are employed to block d.c. and to provide an a.c. path. The radio frequency choke (R.F.C) offers very high impedance to high frequency currents which means it shorts for d.c. and opens for a.c. Hence it provides d.c. load for collector and keeps a.c. currents out of d.c. supply source.

Tank Circuit

The frequency determining network is a parallel resonant circuit which consists of the inductors L_1 and L_2 along with a variable capacitor C . The junction of L_1 and L_2 are earthed. The coil L_1 has its one end connected to base via C_c and the other to emitter via C_e . So, L_2 is in the output circuit. Both the coils L_1 and L_2 are inductively coupled and together form an **Auto-transformer**.

The following circuit diagram shows the arrangement of a Hartley oscillator. The tank circuit is **shunt fed** in this circuit. It can also be a **series-fed**.



Operation

When the collector supply is given, a transient current is produced in the oscillatory or tank circuit. The oscillatory current in the tank circuit produces a.c. voltage across L_1 .

The **auto-transformer** made by the inductive coupling of L_1 and L_2 helps in determining the frequency and establishes the feedback. As the CE configured transistor provides 180° phase shift, another 180° phase shift is provided by the transformer, which makes 360° phase shift between the input and output voltages.

This makes the feedback positive which is essential for the condition of oscillations. When the **loop gain $|\beta A|$ of the amplifier is greater than one**, oscillations are sustained in the circuit.

Frequency

The equation for **frequency of Hartley oscillator** is given as

$$f = \frac{1}{2\pi\sqrt{L_T C}}$$

$$L_T = L_1 + L_2 + 2M$$

Here, L_T is the total cumulatively coupled inductance; L_1 and L_2 represent inductances of 1st and 2nd coils; and M represents mutual inductance.

Mutual inductance is calculated when two windings are considered.

Advantages

The advantages of Hartley oscillator are

- Instead of using a large transformer, a single coil can be used as an auto-transformer.
- Frequency can be varied by employing either a variable capacitor or a variable inductor.
- Less number of components are sufficient.
- The amplitude of the output remains constant over a fixed frequency range.

Disadvantages

The disadvantages of Hartley oscillator are

- It cannot be a low frequency oscillator.
- Harmonic distortions are present.

Applications

The applications of Hartley oscillator are

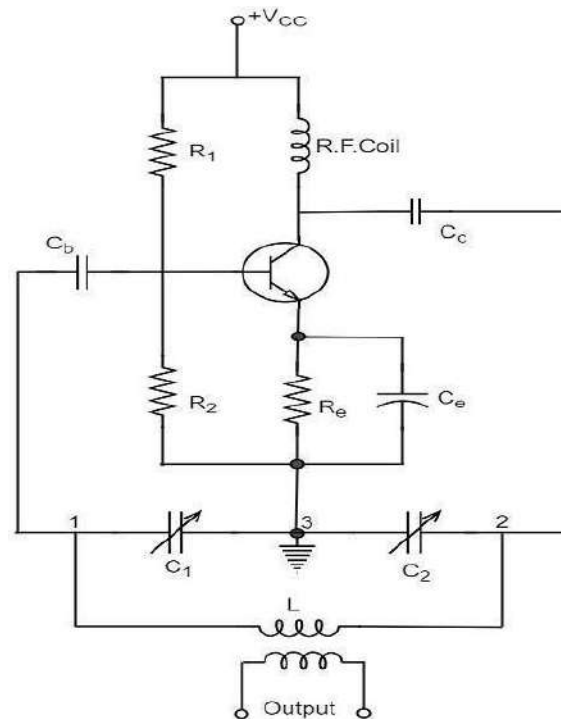
- It is used to produce a sine wave of desired frequency.
- Mostly used as a local oscillator in radio receivers.
- It is also used as R.F. Oscillator.

Colpitts oscillator

A Colpitts oscillator looks just like the Hartley oscillator but the inductors and capacitors are replaced with each other in the tank circuit. The constructional details and operation of a Colpitts oscillator are as discussed below.

Construction

Let us first take a look at the circuit diagram of a Colpitts oscillator.



The resistors R_1 , R_2 and R_e provide necessary bias condition for the circuit. The capacitor C_e provides a.c. ground thereby providing any signal degeneration. This also provides temperature stabilization.

The capacitors C_c and C_b are employed to block d.c. and to provide an a.c. path. The radio frequency choke (R.F.C) offers very high impedance to high frequency currents which means it shorts for d.c. and opens for a.c. Hence it provides d.c. load for collector and keeps a.c. currents out of d.c. supply source.

Tank Circuit

The frequency determining network is a parallel resonant circuit which consists of variable capacitors C_1 and C_2 along with an inductor L . The junction of C_1 and C_2 are earthed. The capacitor C_1 has its one end connected to base via C_c and the other to emitter via C_e . The voltage developed across C_1 provides the regenerative feedback required for the sustained oscillations.

Operation

When the collector supply is given, a transient current is produced in the oscillatory or tank circuit. The oscillatory current in the tank circuit produces a.c. voltage across C_1 which are

applied to the base emitter junction and appear in the amplified form in the collector circuit and supply losses to the tank circuit.

If terminal 1 is at positive potential with respect to terminal 3 at any instant, then terminal 2 will be at negative potential with respect to 3 at that instant because terminal 3 is grounded. Therefore, points 1 and 2 are out of phase by 180°.

As the CE configured transistor provides 180° phase shift, it makes 360° phase shift between the input and output voltages. Hence, feedback is properly phased to produce continuous Undamped oscillations. When the **loop gain $|\beta A|$ of the amplifier is greater than one, oscillations are sustained** in the circuit.

Frequency

The equation for **frequency of Colpitts oscillator** is given as

$$f = \frac{1}{2\pi\sqrt{LC_T}}$$

C_T is the total capacitance of C_1 and C_2 connected in series.

$$\frac{1}{C_T} = \frac{1}{C_1} + \frac{1}{C_2}$$

$$C_T = \frac{C_1 \times C_2}{C_1 + C_2}$$

Advantages

The advantages of Colpitts oscillator are as follows –

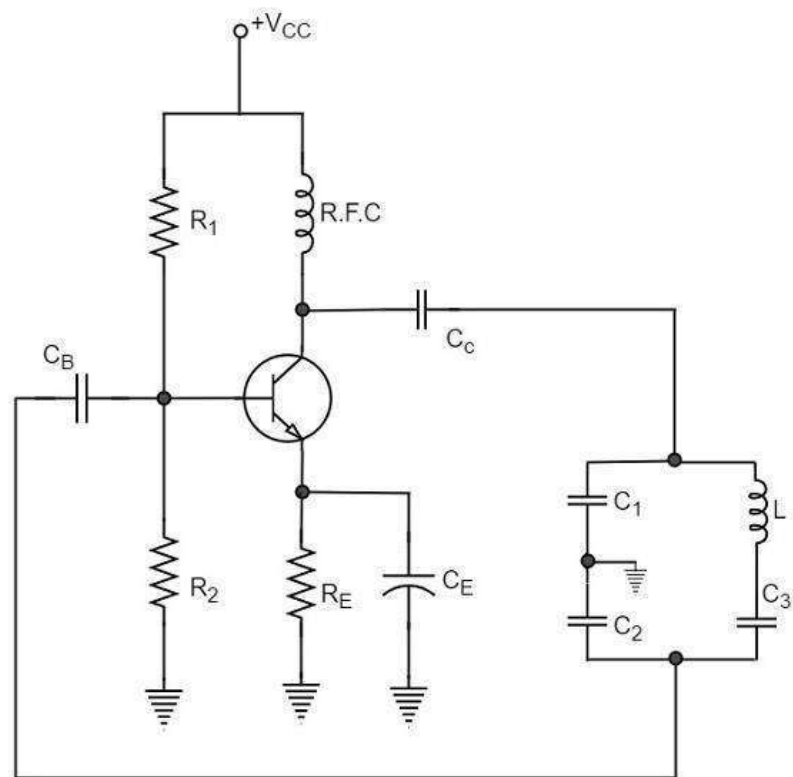
- Colpitts oscillator can generate sinusoidal signals of very high frequencies.
- It can withstand high and low temperatures.
- The frequency stability is high.
- Frequency can be varied by using both the variable capacitors.
- Less number of components are sufficient.
- The amplitude of the output remains constant over a fixed frequency range.

The Colpitts oscillator is designed to eliminate the disadvantages of Hartley oscillator and is known to have no specific disadvantages. Hence there are many applications of a colpitts oscillator.

Applications

The applications of Colpitts oscillator are as follows –

- Colpitts oscillator can be used as High frequency sinewave generator.
- This can be used as a temperature sensor with some associated circuitry.
- Mostly used as a local oscillator in radio receivers.
- It is also used as R.F. Oscillator.
- It is also used in Mobile applications.
- It has got many other commercial applications.



$$f_o = \frac{1}{2\pi\sqrt{L.C}}$$

Drawbacks of LC circuits

Though they have few applications, the **LC** circuits have few **drawbacks** such as

$$C = \frac{1}{\frac{1}{C_1} + \frac{1}{C_2} + \frac{1}{C_3}}$$

- Frequency instability
- Waveform is poor
- Cannot be used for low frequencies

$$f_o = \frac{1}{2\pi\sqrt{L.C_3}}$$

Inductors are bulky and expensive

Whenever an oscillator is under continuous operation, its **frequency stability** gets affected. There occur changes in its frequency. The main factors that affect the frequency of an oscillator are

- Power supply variations
- Changes in temperature
- Changes in load or output resistance

In RC and LC oscillators the values of resistance, capacitance and inductance vary with temperature and hence the frequency gets affected. In order to avoid this problem, the piezo electric crystals are being used in oscillators.

Crystal Oscillators

The use of piezo electric crystals in parallel resonant circuits provide high frequency stability in oscillators. Such oscillators are called as **Crystal Oscillators**.

Crystal Oscillators

The principle of crystal oscillators depends upon the **Piezo electric effect**. The natural shape of a crystal is hexagonal. When a crystal wafer is cut perpendicular to X-axis, it is called as X-cut and when it is cut along Y-axis, it is called as Y-cut.

The crystal used in crystal oscillator exhibits a property called as Piezo electric property. So, let us have an idea on piezo electric effect.

Piezo Electric Effect

The crystal exhibits the property that when a mechanical stress is applied across one of the faces of the crystal, a potential difference is developed across the opposite faces of the crystal. Conversely, when a potential difference is applied across one of the faces, a mechanical stress is produced along the other faces. This is known as **Piezo electric effect**.

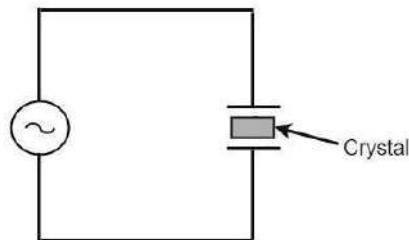
Certain crystalline materials like Rochelle salt, quartz and tourmaline exhibit piezo electric effect and such materials are called as **Piezo electric crystals**. Quartz is the most commonly used piezo electric crystal because it is inexpensive and readily available in nature.

When a piezo electric crystal is subjected to a proper alternating potential, it vibrates mechanically. The amplitude of mechanical vibrations becomes maximum when the frequency of alternating voltage is equal to the natural frequency of the crystal.

Working of a Quartz Crystal

In order to make a crystal work in an electronic circuit, the crystal is placed between two metal plates in the form of a capacitor. **Quartz** is the mostly used type of crystal because of its availability and strong nature while being inexpensive. The ac voltage is applied in parallel to the crystal.

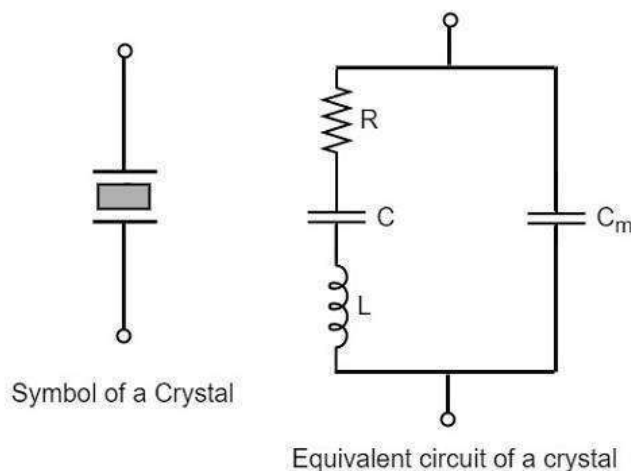
The circuit arrangement of a Quartz Crystal will be as shown below –



If an AC voltage is applied, the crystal starts vibrating at the frequency of the applied voltage. However, if the frequency of the applied voltage is made equal to the natural frequency of the crystal, **resonance** takes place and crystal vibrations reach a maximum value. This natural frequency is almost constant.

Equivalent circuit of a Crystal

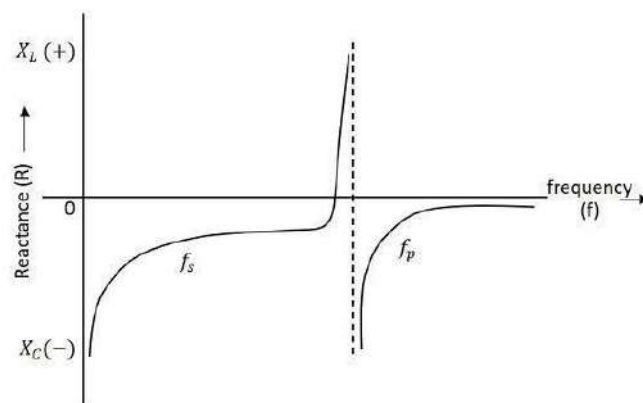
If we try to represent the crystal with an equivalent electric circuit, we have to consider two cases, i.e., when it vibrates and when it doesn't. The figures below represent the symbol and electrical equivalent circuit of a crystal respectively.



The above equivalent circuit consists of a series R-L-C circuit in parallel with a capacitance C_m . When the crystal mounted across the AC source is not vibrating, it is equivalent to the capacitance C_m . When the crystal vibrates, it acts like a tuned R-L-C circuit.

Frequency response

The frequency response of a crystal is as shown below. The graph shows the reactance (X_L or X_C) versus frequency (f). It is evident that the crystal has two closely spaced resonant frequencies.



The first one is the series resonant frequency (f_s), which occurs when reactance of the inductance (L) is equal to the reactance of the capacitance C . In that case, the impedance of the equivalent circuit is equal to the resistance R and the frequency of oscillation is given by the relation,

$$f = \frac{1}{2\pi\sqrt{L \cdot C}}$$

The second one is the parallel resonant frequency (f_p), which occurs when the reactance of R - L - C branch is equal to the reactance of capacitor C_m . At this frequency, the crystal offers a very high impedance to the external circuit and the frequency of oscillation is given by the relation.

$$f_p = \frac{1}{2\pi\sqrt{L \cdot C_T}}$$

Where

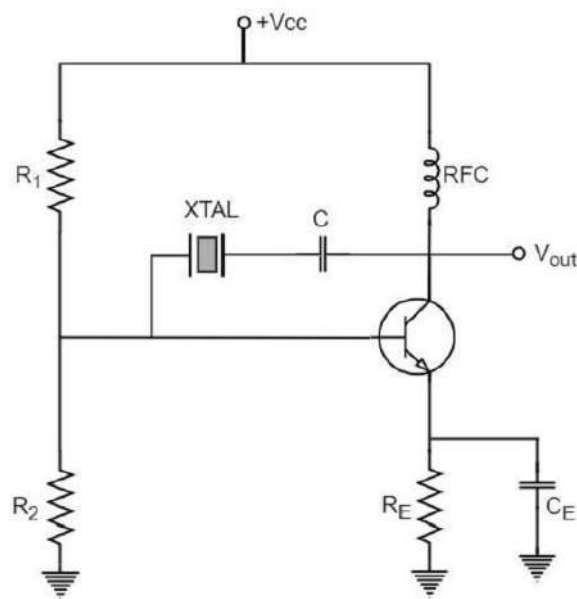
$$C_T = \frac{CC_m}{(C + C_m)}$$

The value of C_m is usually very large as compared to C . Therefore, the value of C_T is approximately equal to C and hence the series resonant frequency is approximately equal to the parallel resonant frequency (i.e., $f_s = f_p$).

Crystal Oscillator Circuit pierce crystal oscillator

A crystal oscillator circuit can be constructed in a number of ways like a Crystal controlled tuned collector oscillator, a Colpitts crystal oscillator, a Clap crystal oscillator etc. But the **transistor pierce crystal oscillator** is the most commonly used one. This is the circuit which is normally referred as a crystal oscillator circuit.

The following circuit diagram shows the arrangement of a transistor pierce crystal oscillator.



In this circuit, the crystal is connected as a series element in the feedback path from collector to the base. The resistors R_1 , R_2 and R_E provide a voltage-divider stabilized d.c. bias circuit. The capacitor C_E provides a.c. bypass of the emitter resistor and RFC (radio frequency choke) coil provides for d.c. bias while decoupling any a.c. signal on the power lines from affecting the output signal. The coupling capacitor C has negligible impedance at the circuit operating frequency. But it blocks any d.c. between collector and base.

The circuit frequency of oscillation is set by the series resonant frequency of the crystal and its value is given by the relation,

$$f = \frac{1}{2\pi\sqrt{L.C}}$$

It may be noted that the changes in supply voltage, transistor device parameters etc. have no effect on the circuit operating frequency, which is held stabilized by the crystal.

Advantages

The advantages of crystal oscillator are as follows –

- They have a high order of frequency stability.
- The quality factor (Q) of the crystal is very high.

Disadvantages

The disadvantages of crystal oscillator are as follows –

- They are fragile and can be used in low power circuits.
- The frequency of oscillations cannot be changed appreciably.

Frequency Stability of an Oscillator

An Oscillator is expected to maintain its frequency for a longer duration without any variations, so as to have a smoother clear sinewave output for the circuit operation. Hence the term frequency stability really matters a lot, when it comes to oscillators, whether sinusoidal or non-sinusoidal.

The frequency stability of an oscillator is defined as the ability of the oscillator to maintain the required frequency constant over a long time interval as possible. Let us try to discuss the factors that affect this frequency stability.

Change in operating point

We have already come across the transistor parameters and learnt how important an operating point is. The stability of this operating point for the transistor being used in the circuit for amplification (BJT or FET), is of higher consideration.

The operating of the active device used is adjusted to be in the linear portion of its characteristics. This point is shifted due to temperature variations and hence the stability is affected.

Variation in temperature

The tank circuit in the oscillator circuit, contains various frequency determining components such as resistors, capacitors and inductors. All of their parameters are temperature dependent. Due to the change in temperature, their values get affected. This brings the change in frequency of the oscillator circuit.

Due to power supply

The variations in the supplied power will also affect the frequency. The power supply variations lead to the variations in V_{cc} . This will affect the frequency of the oscillations produced.

In order to avoid this, the regulated power supply system is implemented. This is in short called as RPS.

Change in output load

The variations in output resistance or output load also affects the frequency of the oscillator. When a load is connected, the effective resistance of the tank circuit is changed. As a result, the Q-factor of LC tuned circuit is changed. This results a change in output frequency of oscillator.

Changes in inter-element capacitances

Inter-element capacitances are the capacitances that develop in PN junction materials such as diodes and transistors. These are developed due to the charge present in them during their operation.

The inter element capacitors undergo change due to various reasons as temperature, voltage etc. This problem can be solved by connecting swamping capacitor across offending inter-element capacitor.

UNIT-5

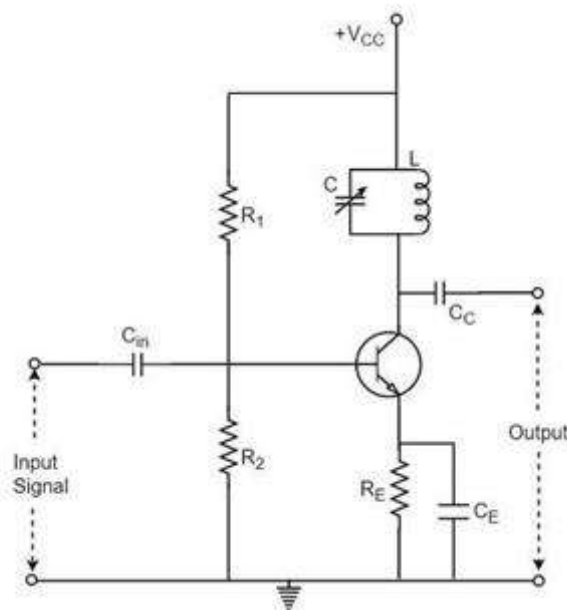
Single Tuned Amplifier

An amplifier circuit with a single tuner section being at the collector of the amplifier circuit is called as Single tuner amplifier circuit.

Construction

A simple transistor amplifier circuit consisting of a parallel tuned circuit in its collector load, makes a single tuned amplifier circuit. The values of capacitance and inductance of the tuned circuit are selected such that its resonant frequency is equal to the frequency to be amplified.

The following circuit diagram shows a single tuned amplifier circuit.



The output can be obtained from the coupling capacitor C_C as shown above or from a secondary winding placed at L .

Operation

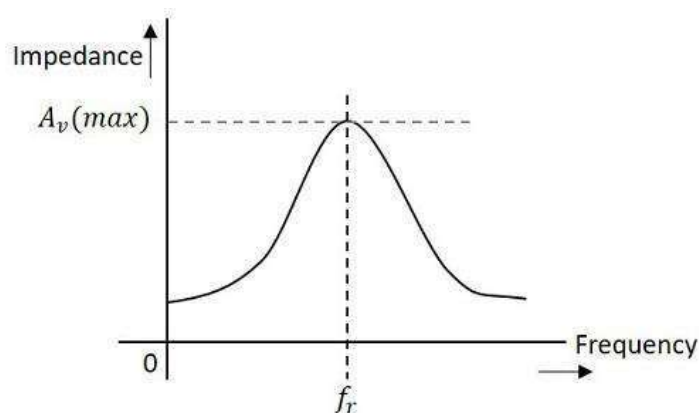
The high frequency signal that has to be amplified is applied at the input of the amplifier. The resonant frequency of the parallel tuned circuit is made equal to the frequency of the signal applied by altering the capacitance value of the capacitor C , in the tuned circuit. At this stage, the tuned circuit offers high impedance to the signal frequency, which helps to offer high output across the tuned circuit. As high impedance is offered only for the tuned frequency, all the other frequencies which get lower impedance are rejected by the tuned circuit. Hence the tuned amplifier selects and amplifies the desired frequency signal.

Frequency Response

The parallel resonance occurs at resonant frequency f_r when the circuit has a high Q. the resonant frequency f_r is given by

$$f_r = 1/\sqrt{2\pi LC}$$

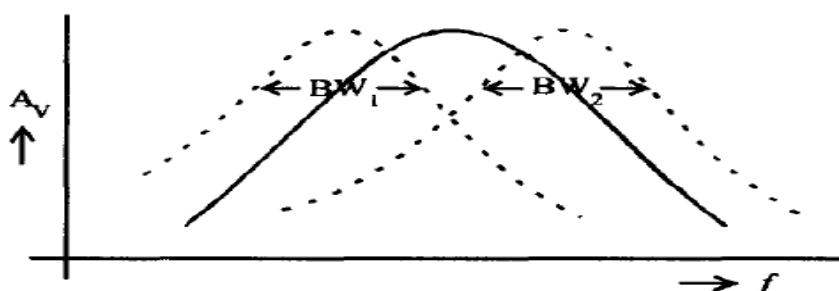
The following graph shows the frequency response of a single tuned amplifier circuit.



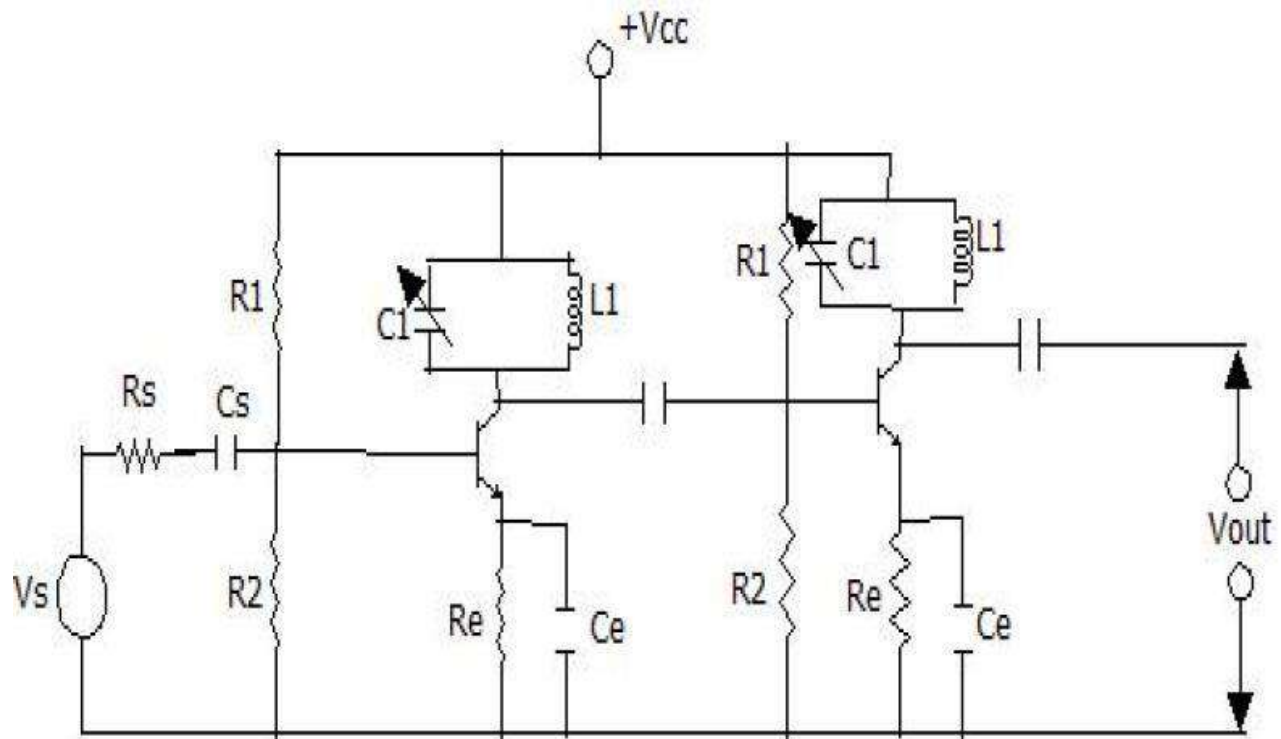
At resonant frequency f_r the impedance of parallel tuned circuit is very high and is purely resistive. The voltage across R_L is therefore maximum, when the circuit is tuned to resonant frequency. Hence the voltage gain is maximum at resonant frequency and drops off above and below it. The higher the Q, the narrower will the curve be.

Stagger Tuning

Tuned amplifiers have large gain, since at resonance, Z is maximum. So A_v is maximum. To get this large A_v over a wide range of frequencies, stagger tuned amplifiers are employed. This is done by taking two single tuned circuits of a certain Bandwidth, and displacing or staggering their resonance peaks by an amount equal to their Bandwidth. The resultant staggered pair will have a Bandwidth, $\sqrt{2}$ times as great as that of each of individual pairs.



The circuit of stagger tuned amplifier is as shown below:



Stagger Tuned Amplifiers are used to improve the overall frequency response of tuned Amplifiers. Stagger tuned Amplifiers are usually designed so that the overall response exhibits maximal flatness around the centre frequency.

It needs a number of tuned circuit operating in union. The overall frequency response of a Stagger tuned amplifier is obtained by adding the individual response together. Since the resonant Frequencies of different tuned circuits are displaced or staggered, they are referred as Stagger Tuned Amplifier.

The main advantage of stagger tuned amplifier is increased bandwidth. Its Drawback is Reduced Selectivity and critical tuning of many tank circuits. They are used in RF amplifier stage in Radio Receivers.

The stagger tuning in this circuit is achieved by resonating the tuned circuits L1 C1, L2 C2 to slightly different Frequencies

MULTIVIBRATORS

Multi means many; vibrator means oscillator. A circuit which can oscillate at a number of frequencies is called a multivibrator. Basically there are three types of multivibrators:

1. Bistable multivibrator
2. Monostable multivibrator
3. Astable multivibrator

Each of these multivibrators has two states. As the names indicate, a bistable multivibrator has got two stable states, a monostable multivibrator has got only one stable state (the other state being quasi stable) and the astable multivibrator has got no stable state (both the

states being quasi stable). The stable state of a multivibrator is the state in which the device can stay permanently. Only when a proper external triggering signal is applied, it will change its state. Quasi stable state means temporarily stable state. The device cannot stay permanently in this state. After a predetermined time, the device will automatically come out of the quasi stable state.

Multivibrators find applications in a variety of systems where square waves or timed intervals are required. For example, before the advent of low-cost integrated circuits, chains of multivibrators found use as frequency dividers. A free-running multivibrator with a frequency of one-half to one-tenth of the reference frequency would accurately lock to the reference frequency. This technique was used in early electronic organs, to keep notes of different octaves accurately in tune. Other applications included early television systems, where the various line and frame frequencies were kept synchronized by pulses included in the video signal.

BISTABLE MULTIVIBRATOR

A bistable multivibrator is a multivibrator which can exist indefinitely in either of its two stable states and which can be induced to make an abrupt transition from one state to the other by means of external excitation. In a bistable multivibrator both the coupling elements are resistors (dc coupling). The bistable multivibrator is also called a multi, Eccles-Jordan circuit (after its inventors), trigger circuit, scale-of-two toggle circuit, flip-flop, and binary. There are two types of bistable multivibrators:

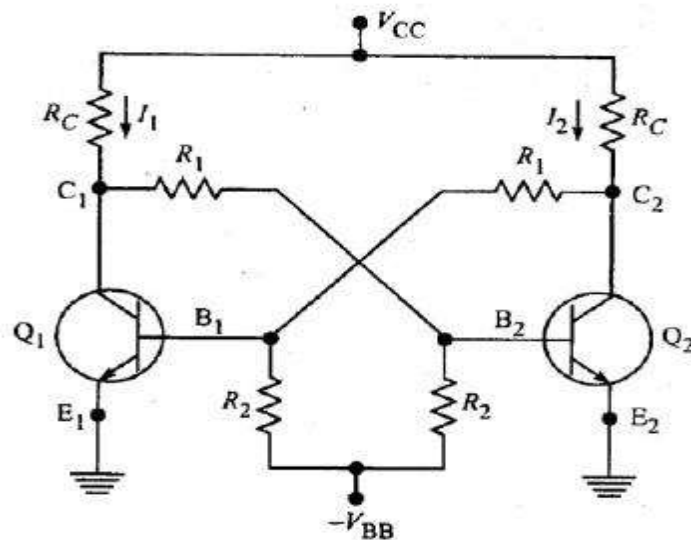
1. Collector coupled bistable multivibrator
2. Emitter coupled bistable multivibrator

There are two types of collector-coupled bistable multivibrators:

1. Fixed-bias bistable multivibrator
2. Self-bias bistable multivibrator

A FIXED-BIAS BISTABLE MULTIVIBRATOR

The Figure below shows the circuit diagram of a fixed-bias bistable multivibrator using transistors (inverters). Note, that the output of each amplifier is direct coupled to the input of the other amplifier.



In one of the stable states, transistor Q_1 is ON (i.e. in saturation) and Q_2 is OFF (i.e. in cut-off), and in the other stable state Q_1 is OFF and Q_2 is ON. Even though the circuit is symmetrical, it is not possible for the circuit to remain in a stable state with both the transistors conducting (i.e. both operating in the active region) simultaneously and carrying equal currents. The reason is that if we assume that both the transistors are biased equally and are carrying equal currents I_1 and I_2 and suppose there is a minute fluctuation in the current I_1 —let us say it increases by a small amount—then the voltage at the collector of Q_1 decreases. This will result in a decrease in voltage at the base of Q_2 . So Q_2 conducts less and I_2 decreases and hence the potential at the collector of Q_2 increases. This results in an increase in the base potential of Q_1 . So, Q_1 conducts still more and I_1 is further increased and the potential at the collector of Q_1 is further reduced, and so on. So, the current I_1 keeps on increasing and the current I_2 keeps on decreasing till Q_1 goes into saturation and Q_2 goes into cut-off. This action takes place because of the regenerative feedback incorporated into the circuit and will occur only if the loop gain is greater than one. *A stable state of a binary is one in which the voltages and currents satisfy the Kirchhoff's laws and are consistent with the device characteristics and in which, in addition, the condition of the loop gain being less than unity is satisfied.*

The condition with respect to loop gain will certainly be satisfied, if either of the two devices is below cut-off or if either device is in saturation. But normally the circuit is designed such that in a stable state one transistor is in saturation and the other one is in cut-off, because if one transistor is biased to be in cut-off and the other one to be in active region, as the temperature changes or the devices age and the device parameters vary, the quiescent point changes and the quiescent output voltage may also change appreciably. Sometimes the drift may

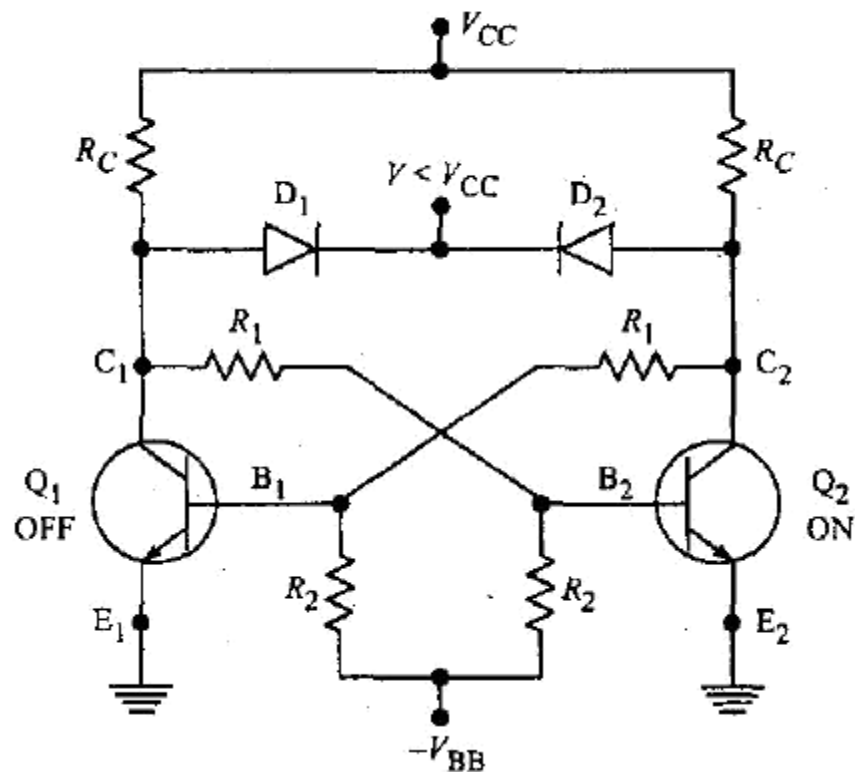
be so much that the device operating in the active region may go into cut-off, and with both the devices in cut-off the circuit will be useless.

Selection of components in the fixed-bias bistable multivibrator

In the fixed-bias binary shown in Figure 4.1., nearly the full supply voltage V_{CC} will appear across the transistor that is OFF. Since this supply voltage V_{CC} is to be reasonably smaller than the collector breakdown voltage SV_{ce} . V_{CC} restricted to a maximum of a few tens of volts. Under saturation conditions the collector current I_c is maximum. Hence RC must be chosen so that this value of $C (= V_{CC}/\Delta G)$ does not exceed the maximum permissible limit. The values of R_1 , R_2 and V_{BB} must be selected such that in one stable state the base current is large enough to drive the transistor into saturation whereas in the second stable state the emitter junction must be below cut-off. The signal at a collector called the output swing V_w is the change in collector voltage resulting from a transistor going from one state to the other, i.e. $V_w = V_{C1} - V_{C2}$. If the loading caused by R_1 can be neglected, then the collector voltage of the OFF transistor is V_{CC} . Since the collector saturation voltage is few tenths of a volt, then the swing $V_w = V_{CC}$, independently of RQ . The component values, the supply voltages and the values of β_{CBO} , h^{β} , $V_{BE}(sat)$, and $V_{CE}(sat)$ are sufficient for the analysis of transistor binary circuits.

Loading

The bistable multivibrator may be used to drive other circuits and hence at one or both the collectors there are shunting loads, which are not shown in Figure 4.1. These loads reduce the magnitude of the collector voltage V_{C1} of the OFF transistor. This will result in reduction of the output voltage swing. A reduced V_{C1} will decrease V_{B2} and it is possible that Q_2 may not be driven into saturation. Hence the flip-flop circuit components must be chosen such that under the heaviest load, which the binary drives, one transistor remains in saturation while the other is in cut-off. Since the resistor R_1 also loads the OFF transistor, to reduce loading, the value of R_1 should be as large as possible compared to the value of R_c . But to ensure a loop gain in excess of unity during the transition between the states, R^{β} should be selected such that For some applications, the loading varies with the operation being performed. In such cases, the extent to which a transistor is driven into saturation is variable. A constant output swing $V_w = V$, and a constant base saturation current I_{B2} can be obtained by clamping the collectors to an auxiliary voltage $V < V_{CC}$ through the diodes D_1 and D_2 as indicated in Figure 4.2. As Q_1 cuts OFF, its collector voltage rises and when it reaches V , the "collector catching diode" D_1 conducts and clamps the output to V .



Standard specifications

In the cut-off region, i.e. for the OFF state

$$V_{BE} \text{ (cut-off)} : \leq 0 \text{ V for silicon transistor} \\ \leq -0.1 \text{ V for germanium transistor}$$

In the saturation region, i.e. for the ON state

$$V_{BE} \text{ (sat)} : 0.7 \text{ V for silicon transistor} \\ 0.3 \text{ V for germanium transistor} \\ V_{CE} \text{ (sat)} : 0.3 \text{ V for silicon transistor} \\ 0.1 \text{ V for germanium transistor}$$

The above values hold good for n-p-n transistors. For p-n-p transistors the above values with opposite sign are to be taken.

Test for saturation

To test whether a transistor is really in saturation or not evaluate the collector current i_C and the base current i_B independently.

If $i_B > i_B \text{ (min)}$, where $i_B \text{ (min)} = i_C / h_{FE} \text{ (min)}$ the transistor is really in saturation.

If $i_B \leq i_B \text{ (min)}$, the transistor is not in saturation.

Test for cut-off

To test whether a transistor is really cut-off or not, find its base-to-emitter voltage. If V_{BE} is negative for an n-p-n transistor or positive for a p-n-p transistor, the transistor is really cut-off.

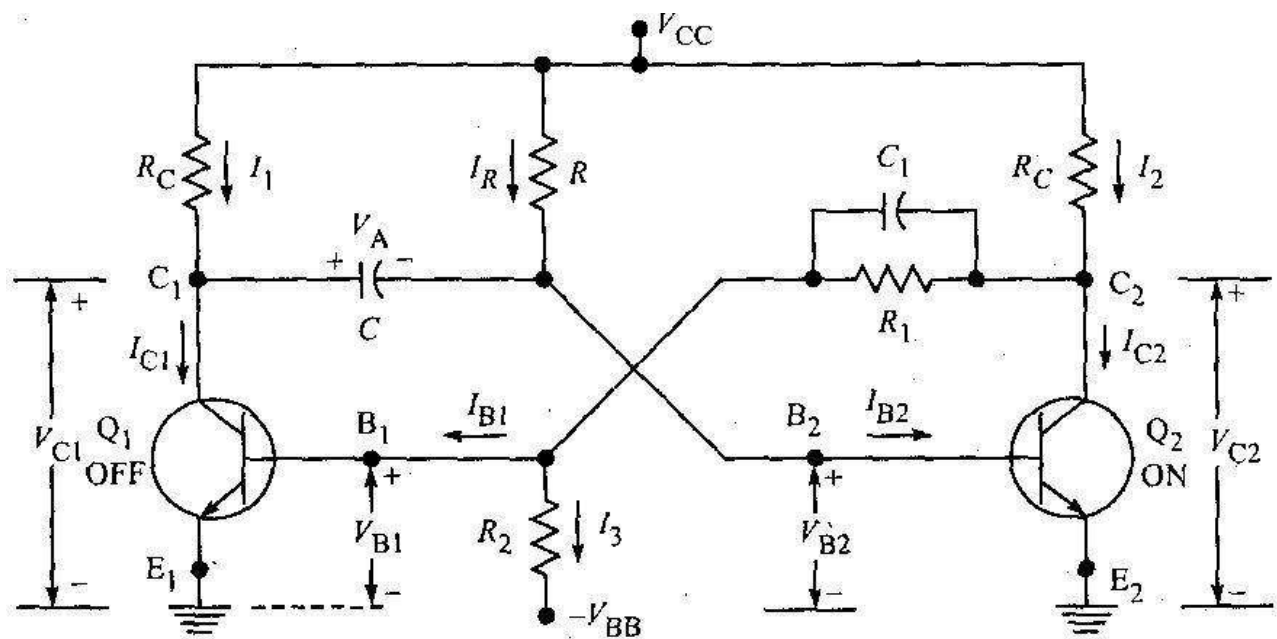
MONOSTABLE MULTIVIBRATOR

Monostable Multivibrators have only **one** stable state (hence their name: "Mono"), and produce a single output pulse when it is triggered externally. Monostable multivibrators only return back to their first original and stable state after a period of time determined by the time constant of the RC coupled circuit.

Monostable multivibrators or "One-Shot Multivibrators" as they are also called, are used to generate a single output pulse of a specified width, either "HIGH" or "LOW" when a suitable external trigger signal or pulse T is applied. This trigger signal initiates a timing cycle which causes the output of the monostable to change its state at the start of the timing cycle and will remain in this second state, which is determined by the time constant of the timing capacitor, C_T and the resistor, R_T until it resets or returns itself back to its original (stable) state. It will then remain in this original stable state indefinitely until another input pulse or trigger signal is received. Then, **Monostable Multivibrators** have only **ONE** stable state and go through a full cycle in response to a single triggering input pulse.

THE COLLECTOR COUPLED MONOSTABLE MULTIVIBRATOR

The below Figure shows the circuit diagram of a collector-to-base coupled (simply called collectorcoupled) monostable multivibrator using n-p-n transistors. The collector of Q_2 is coupled to the base of Q_1 by a resistor R_f (dc coupling) and the collector of Q_1 is coupled to the base of Q_2 by a capacitor C (ac coupling). C_i is the commutating capacitor introduced to increase the speed of operation. The base of Q_1 is connected to $-V_{BB}$ through a resistor R_2 , to ensure that Q_1 is cut off under quiescent conditions. The base of Q_2 is connected to V_{CC} through R to ensure that Q_2 is ON under quiescent conditions. In fact, R may be returned to even a small positive voltage but connecting it to V_{CC} is advantageous. The circuit parameters are selected such that under quiescent conditions, the monostable multivibrator finds itself in its permanent stable state with Q_2 ON (i.e. in saturation) and Q_1 OFF (i.e. in cut-off)- The multivibrator may be induced to make a transition out of its stable state by the application of a negative trigger at the base of Q_2 or at the collector of Q_1 . Since the triggering signal is applied to only one device and not to both the devices simultaneously, unsymmetrical triggering is employed. When a negative signal is applied at the base of Q_2 at $t \sim 0$, due to regenerative action Q_2 goes to OFF state and Q_1 goes to ON state. When Q_1 is ON, a current I_1 flows through its R_c and hence its collector voltage drops suddenly by $I_1 R_c$. This drop will be instantaneously



transmitted through the coupling capacitor C to the base of Q_2 . So at $t = 0+$, the base voltage of Q_2 is

$$V_{BE}(\text{sat}) - I_1 R_C$$

The circuit cannot remain in this state for a long time (it stays in this state only for a finite time T) because when Q_1 conducts, the coupling capacitor C charges from V_{CC} through the conducting transistor Q_1 and hence the potential at the base of Q_2 rises exponentially with a time constant

$(R + R_o)C \approx RC$, where R_o is the conducting transistor output impedance including the resistance R_c . When it passes the cut-in voltage V_y of Q_2 (at a time $t = T$), a regenerative action takes place turning Q_1 OFF and eventually returning the multivibrator to its initial stable state. The transition from the stable state to the quasi-stable state takes place at $t = 0$, and the reverse transition from the quasi-stable state to the stable state takes place at $t = T$. The time T for which the circuit is in its quasi-stable state is also referred to as the delay time, and also as the gate width, pulse width, or pulse duration. The delay time may be varied by varying the time constant $t(= RC)$.

Expression for the gate width T of a monostable multivibrator neglecting the reverse saturation current /CBO

The below Figure (a) shows the waveform at the base of transistor Q_2 of the monostable multivibrator

For $t < 0$, Q2 is ON and so $v_{B2} = V_{BE}(\text{sat})$. At $t = 0$, a negative signal applied brings Q2 to OFF state and Q1 into saturation. A current I flows through R_C of Q1 and hence v_{C1} drops abruptly by $I R_C$ volts and so v_{B2} also drops by $I R_C$ instantaneously. So at $t = 0^+$, $v_{B2} = V_{BE}(\text{sat}) - I R_C$. For $t > 0$, the capacitor charges with a time constant τ , and hence the base voltage of Q2 rises exponentially towards V_{CC} with the same time constant. At $t = T$, when this base voltage rises to the cut-in voltage level V_γ of the transistor, Q2 goes to ON state, and Q1 to OFF state and the pulse ends. In the interval $0 < t < T$, the base voltage of Q2, i.e. v_{B2} is given by

$$v_{B2} = V_{CC} - (V_{CC} - \{V_{BE}(\text{sat}) - I R_C\})e^{-t/\tau}$$

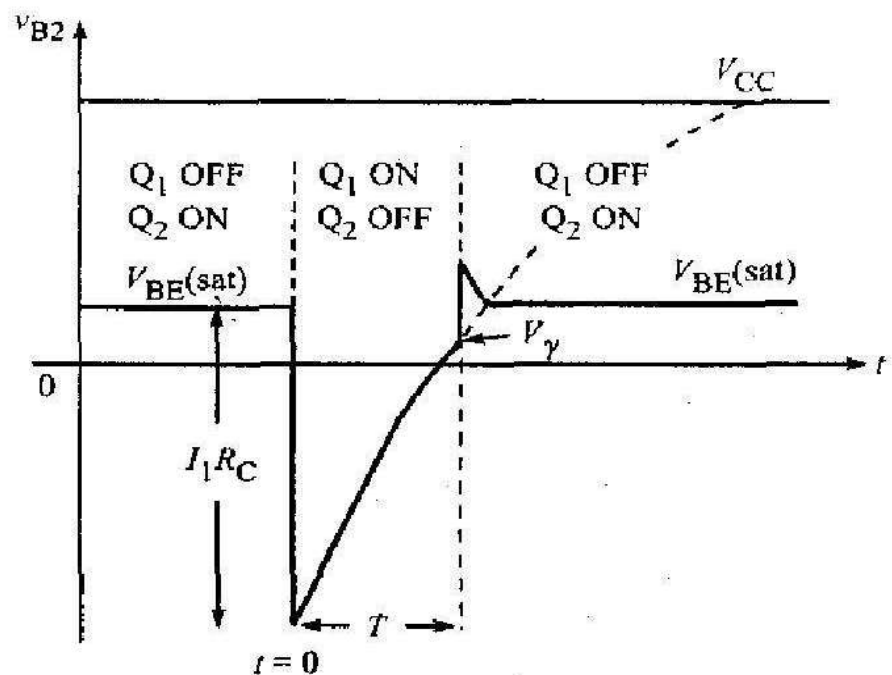


Fig a) Voltage variation at the base of Q2 during the quasi-stable state

But $I R_C = V_{CC} - V_{CE}(\text{sat})$ (because at $t = 0^-$, $v_{C1} = V_{CC}$ and at $t = 0^+$, $v_{C1} = V_{CE}(\text{sat})$)

$$\begin{aligned} \therefore v_{B2} &= V_{CC} - [V_{CC} - \{V_{BE}(\text{sat}) - (V_{CC} - V_{CE}(\text{sat}))\}]e^{-t/\tau} \\ &= V_{CC} - [2V_{CC} - \{V_{BE}(\text{sat}) + V_{CE}(\text{sat})\}]e^{-t/\tau} \end{aligned}$$

At $t = T$, $v_{B2} = V_\gamma$

$$\therefore V_\gamma = V_{CC} - [2V_{CC} - \{V_{CE}(\text{sat}) + V_{BE}(\text{sat})\}]e^{-T/\tau}$$

$$\text{i.e. } e^{T/\tau} = \frac{2V_{CC} - [V_{CE}(\text{sat}) + V_{BE}(\text{sat})]}{V_{CC} - V_\gamma}$$

$$\therefore \frac{T}{\tau} = \frac{\ln \left[2 \left(V_{CC} - \frac{V_{CE}(\text{sat}) + V_{BE}(\text{sat})}{2} \right) \right]}{V_{CC} - V_\gamma}$$

$$\text{i.e. } T = \tau \ln 2 + \tau \ln \frac{V_{CC} - \frac{V_{CE}(\text{sat}) + V_{BE}(\text{sat})}{2}}{V_{CC} - V_\gamma}$$

Normally for a transistor, at room temperature, the cut-in voltage is the average of the saturation junction

$$V_{\gamma} = \frac{V_{CE}(\text{sat}) + V_{BE}(\text{sat})}{2}$$

voltages for either Ge or Si transistors, i.e.

Neglecting the second term in the expression for T

$$T = \tau \ln 2$$

$$T = (R + R_o)C \ln 2 = 0.693(R + R_o)C$$

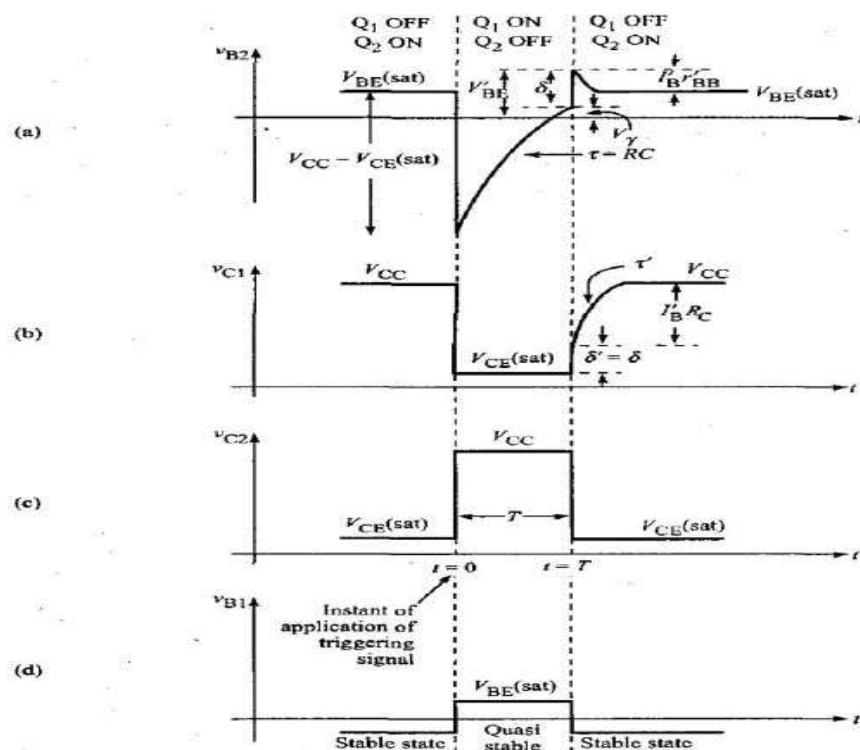
but for a transistor in saturation $R_a \ll R$.

$$\text{Gate width, } T = 0.693RC$$

The larger the V_{CC} is, compared to the saturation junction voltages, the more accurate the result is. The gate width can be made very stable (almost independent of transistor characteristic supply voltages, and resistance values) if Q1 is driven into saturation during the quasi-stable state.

Waveforms of the collector-coupled monostable multivibrator

The waveforms at the collectors and bases of both the transistors Q1 and Q2 are shown below



(a) at the base of Q2, (b) at the collector of Q1, (c) at the collector of Q2, and (d) at the base of Q1

ASTABLE MULTIVIBRATOR

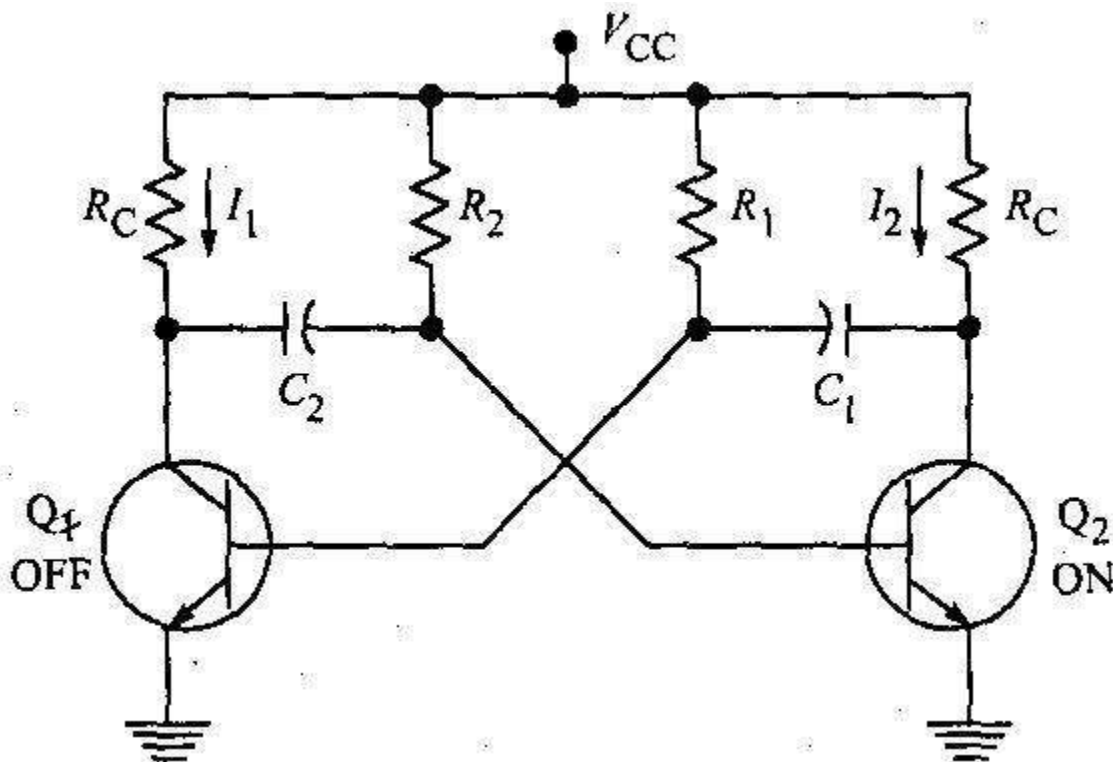
As the name indicates an astable multivibrator is a multivibrator with no permanent stable state. Both of its states are quasi stable only. It cannot remain in any one of its states indefinitely and keeps on oscillating between its two quasi stable states the moment it is connected to the supply. It remains in each of its two quasi stable states for only a short designed interval of time and then goes to the other quasi stable state. No triggering signal is required. Both the coupling elements are capacitors (ac coupling) and hence both the states are quasi stable. It is a free running multivibrator. It generates square waves. It is used as a master oscillator.

There are two types of astable multivibrators:

1. Collector-coupled astable multivibrator
2. Emitter-coupled astable multivibrator

THE COLLECTOR-COUPLED ASTABLE MULTIVIBRATOR

The below Figure shows the circuit diagram of a collector-coupled astable multivibrator using n-p-n transistors. The collectors of both the transistors Q_1 and Q_2 are connected to the bases



of the other transistors through the coupling capacitors C_1 and C_2 . Since both are ac couplings, neither transistor can remain permanently at cut-off. Instead, the circuit has two quasi-stable states, and it makes periodic transitions between these states. Hence it is used as a master oscillator. No triggering signal is required for this multivibrator. The component values are selected such that, the moment it is connected to the supply, due to supply transients one

transistor will go into saturation and the other into cut-off, and also due to capacitive couplings it keeps on oscillating between its two quasi stable states.

The waveforms at the bases and collectors for the astable multivibrator, are shown in below Figure. Let us say at $t = 0$, Q2 goes to ON state and Q1 to OFF state. So, for $t < 0$, Q2 was OFF and Q1 was ON

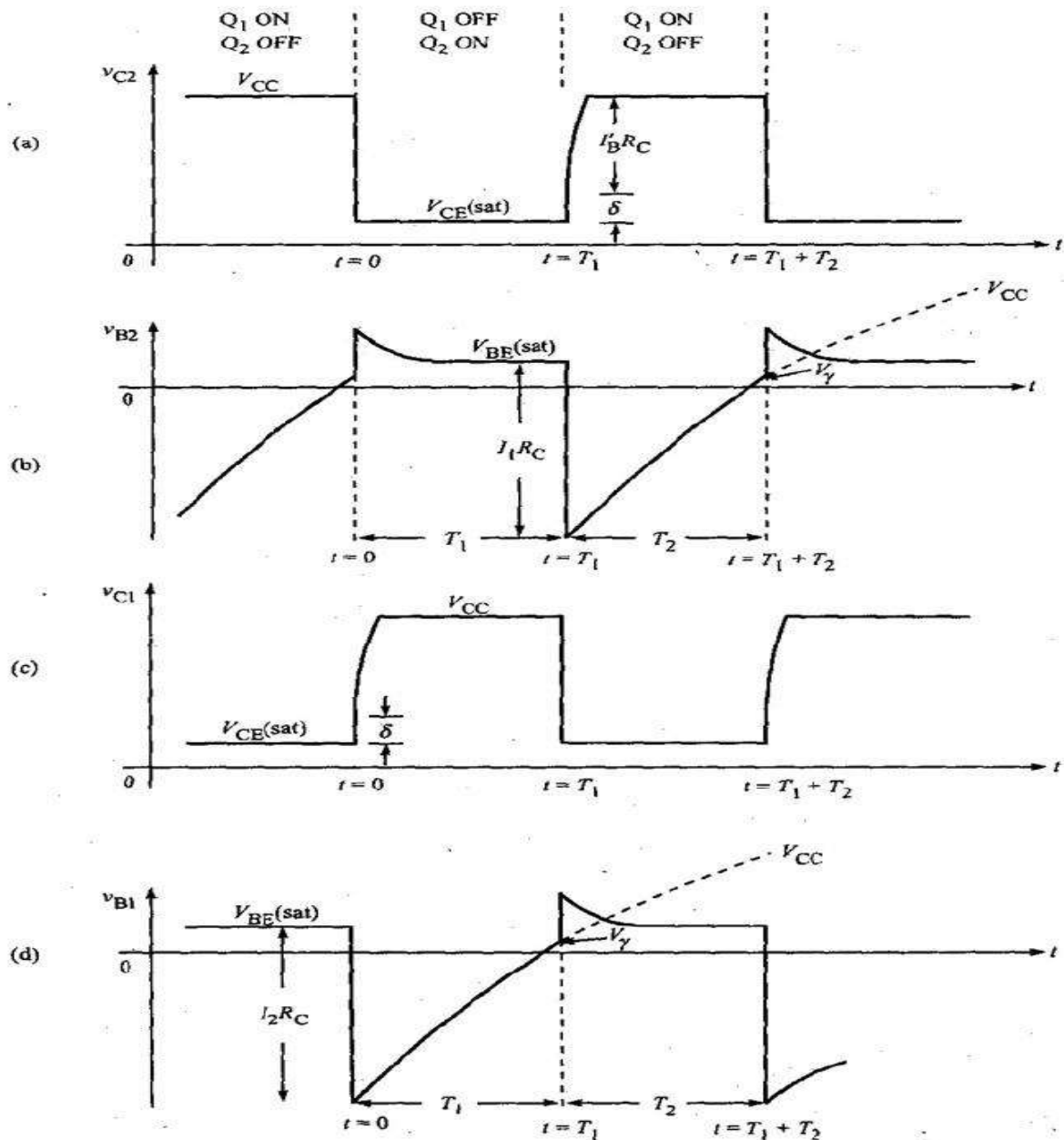


Fig: waveforms at the bases and collectors of a collector-coupled astable multivibrator

Hence for $t < 0$, v_{B2} is negative, $v_{C2} = V_{CC}$, $v_{B1} = V_{BE(sat)}$ and $v_{C1} = V_{CE(sat)}$. The capacitor C_2 charges from V_{CC} through R_2 and v_{B2} rises exponentially towards V_{CC} . At $t = 0$, v_{B2} reaches the cut-in voltage V_γ and Q2 conducts. As Q2 conducts, its collector voltage v_{C2} drops by $V_{CC} - V_{CE(sat)}$. This drop in v_{C2} is transmitted to the base of Q1 through the coupling capacitor C_2 .

and hence v_{B1} also falls by $V_{CE(sat)}$. Q_1 goes to OFF state. So, $V_{B1} = V_{BE(sat)} - I_2 R_C$, and its collector voltage v_{C1} rises towards V_{CC} . This rise in v_{C1} is coupled through the coupling capacitor C_2 to the base of Q_2 , causing an overshoot δ in v_{B2} and the abrupt rise by the same amount δ in V_{C1} as shown in Figure 4.51(c). Now since Q_2 is ON, C_2 charges from V_{CC} through R_1 and hence V_{B1} rises exponentially. At $t = T_1$, when V_{B1} rises to V_γ , Q_1 conducts and due to regenerative action Q_1 goes into saturation and Q_2 to cut-off. Now, for $t > T_1$, the coupling capacitor C_2 charges from V_{CC} through R_2 and at $t = T_1 + T_2$, when v_{B2} rises to the cut-in voltage V_r , Q_2 conducts and due to regenerative feedback Q_2 goes to ON state and Q_1 to OFF state. The cycle of events repeats and the circuit keeps on oscillating between its two quasi-stable states. Hence the output is a square wave. It is called a square wave generator or square wave oscillator or relaxation oscillator. It is a free running oscillator.

Expression for the frequency of oscillation of an astable multivibrator

Consider the waveform at the base of Q_1 shown in Figure 4.54(d). At $t = 0$,

$$v_{B1} = V_{BE(sat)} - I_2 R_C$$

But

$$I_2 R_C = V_{CC} - V_{CE(sat)}$$

$\therefore$

$$\text{At } t = 0, v_{B1} = V_{BE(sat)} - V_{CC} + V_{CE(sat)}$$

For $0 < t < T_1$, v_{B1} rises exponentially towards V_{CC} given by the equation,

$$v_o = v_f - (v_f - v_i)e^{-t/\tau}$$

$$\therefore v_{B1} = V_{CC} - [V_{CC} - (V_{BE(sat)} - V_{CC} + V_{CE(sat)})]e^{-t/\tau_1}, \text{ where } \tau_1 = R_1 C_1$$

At $t = T_1$, when v_{B1} rises to V_γ , Q_1 conducts

$\therefore$

$$V_\gamma = V_{CC} - [2V_{CC} - (V_{BE(sat)} + V_{CE(sat)})]e^{-T_1/R_1 C_1}$$

or

$$e^{T_1/R_1 C_1} = \frac{2 \left[V_{CC} - \frac{V_{BE(sat)} + V_{CE(sat)}}{2} \right]}{V_{CC} - V_\gamma}$$

$$T_1 = R_1 C_1 \ln \frac{2 \left[V_{CC} - \frac{V_{CE(sat)} + V_{BE(sat)}}{2} \right]}{V_{CC} - V_\gamma}$$

$$T_1 = R_1 C_1 \ln 2 + R_1 C_1 \ln \frac{\left[V_{CC} - \frac{V_{CE(sat)} + V_{BE(sat)}}{2} \right]}{V_{CC} - V_\gamma}$$

At room temperature for a transistor,

$$V_\gamma \approx \frac{V_{CE(sat)} + V_{BE(sat)}}{2}$$

$\therefore$

$$T_1 = R_1 C_1 \ln 2 = 0.693 R_1 C_1$$

On similar lines considering the waveform of above Figure , we can show that the time T_2 for which Q2 is OFF and Q1 is ON is given by The period of the waveform, The frequency of oscillation, If $R_1 = R_2 = R$, and $C_1 = C_2 = C$, then $T_1 = T_2 = T$.

$$T = 2 \times 0.693RC = 1.386RC \quad \text{and} \quad f = \frac{1}{1.386RC}$$

The frequency of oscillation may be varied over the range from cycles to mega cycles by varying RC . It is also possible to vary the frequency electrically by connecting R_1 and R_2 to an auxiliary voltage source V (the collector supply remains $+V_{CC}$) and then varying this voltage V .

THE EMITTER-COUPLED ASTABLE MULTIVIBRATOR

An emitter-coupled astable multivibrator may be obtained by using three power supplies or a single power supply. The below Figure (a) shows the circuit diagram of a free-running emitter coupled multivibrator using n-p-n transistors. Figure 4.64 shows its waveforms. Three power supplies are indicated for the sake of simplifying the analysis. A more practical circuit using a single supply is indicated in below Figure (b). Let us assume that the circuit operates in such a manner that Q1 switches between cut-off and saturation and Q2 switches between cut-off and its active region.

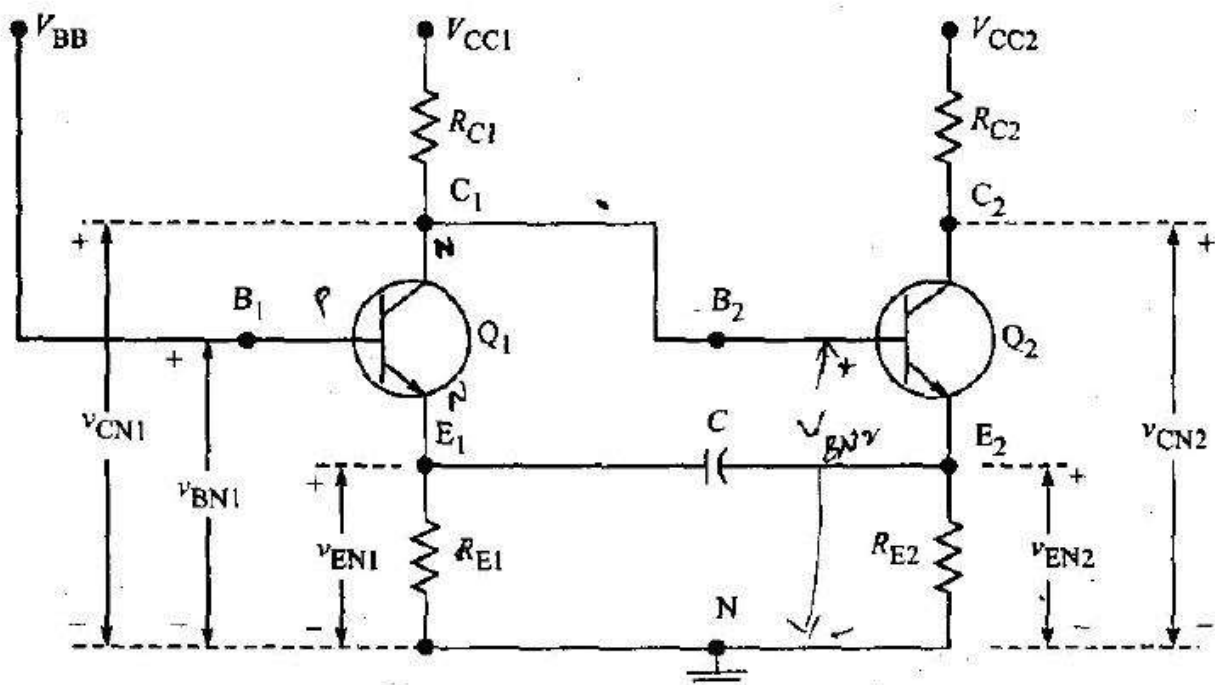


Fig a) Astable Emitter-Coupled Multivibrator

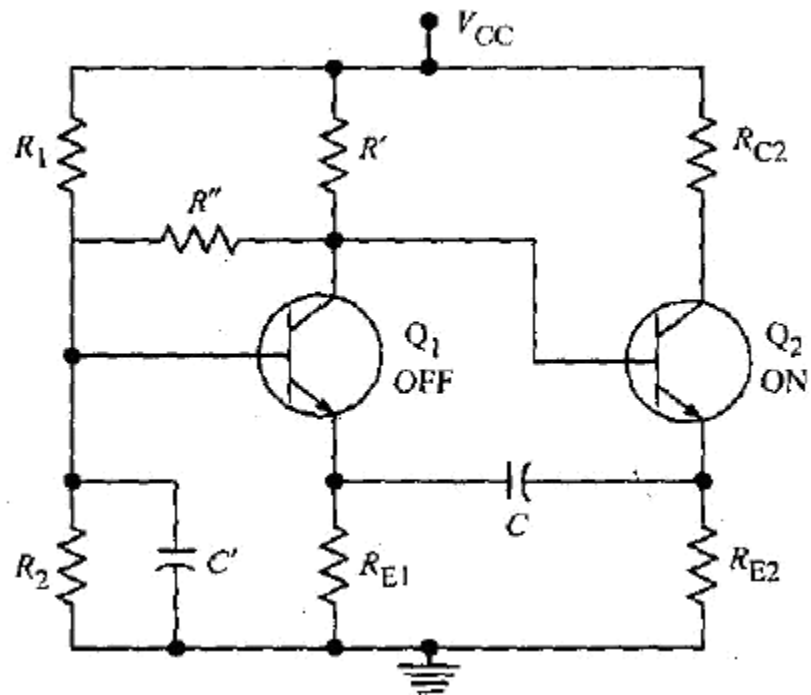


Fig b) Emitter Coupled multivibrator

The waveforms at the base and collector are as shown below:

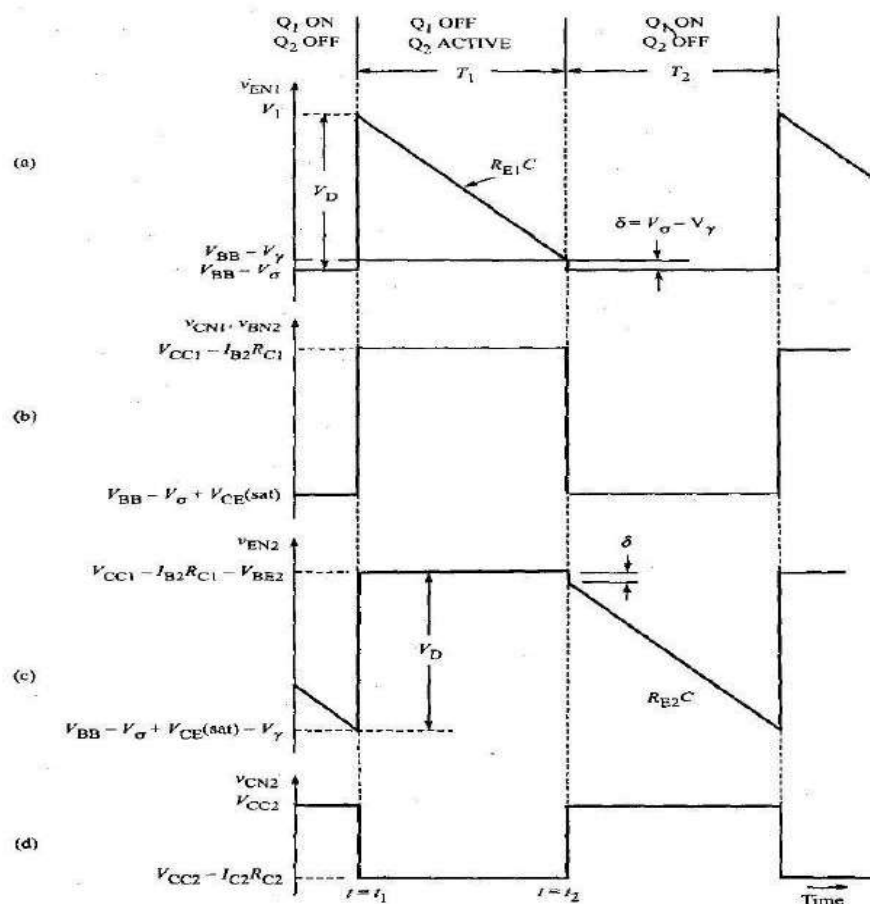


Fig waveforms of the emitter-coupled astable multivibrator

Advantages

1. It is inherently self-starting.
2. The collector of Q2 where the output is taken may be loaded heavily even capacitively.
3. The output is free of recovery transients.
4. Because it has an isolated input at the base of Q1, synchronization is convenient.
5. Frequency adjustment is convenient because only one capacitor is used.

Disadvantages

1. This circuit is more difficult to adjust for proper operating conditions.
2. This circuit cannot be operated with T_1 and T_2 widely different.
3. This circuit uses more components than does the collector-coupled circuit.

P-N DIODE CLIPPER CIRCUIT:

Clipper circuits are the circuits that clip off or remove a portion of an input signal, without causing any distortion to the remaining part of the waveform. These are also known as clippers, clipping circuits, limiters, slicers etc.

Clippers are basically wave shaping circuits that control the shape of an output waveform. It consists of linear and non-linear elements but does not contain energy storing elements.

The basic operation of a diode clipping circuits is such that, in forward biased condition, the diode allows current to pass through it, clamping the voltage. But in reverse biased condition, no any current flows through the diode, and thus voltage remains unaffected across its terminals.

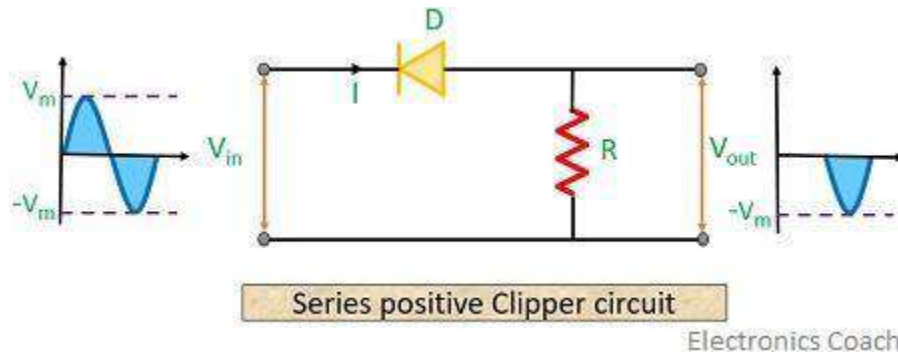
Clipper circuits are basically termed as protection devices. As electronic devices are voltage sensitive and voltage of large amplitude can permanently destroy the device. So, in order to protect the device clipper circuits are used.

Usually, clippers employ resistor–diode combination in its circuitry.

Classification of Clipper circuits

Clippers are basically classified in the following categories:

- **Series positive Clipper circuit**



Let's have a look at the circuit diagram of a series positive clipper. Here, the diode is connected in series with the output thus it is named so.

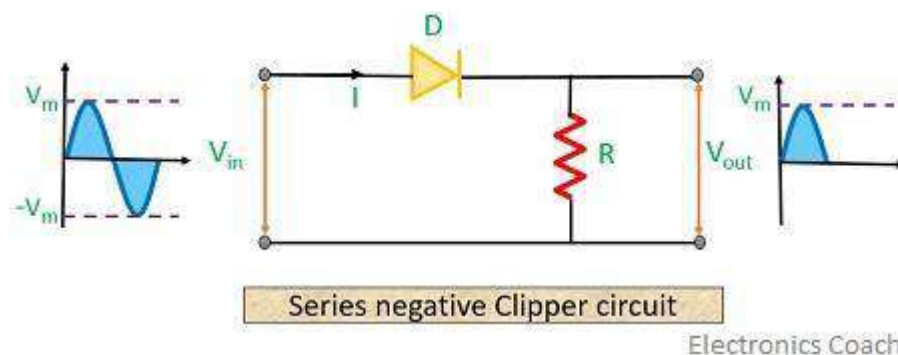
The positive half of the input waveform reverse biases the diode. Thus it acts as an open switch and all the applied input voltage drops across the diode. Resultantly providing no output voltage for positive half of the input waveform.

For the negative half of the input waveform, the diode is in the forward biased state. Thus it acts as a closed switch causing no any voltage drop at the diode.

Hence input voltage will appear across the resistor, ultimately at the output of the circuit.

- **Series negative clipper circuit**

The circuit below shows the figure of a series negative clipping circuit.



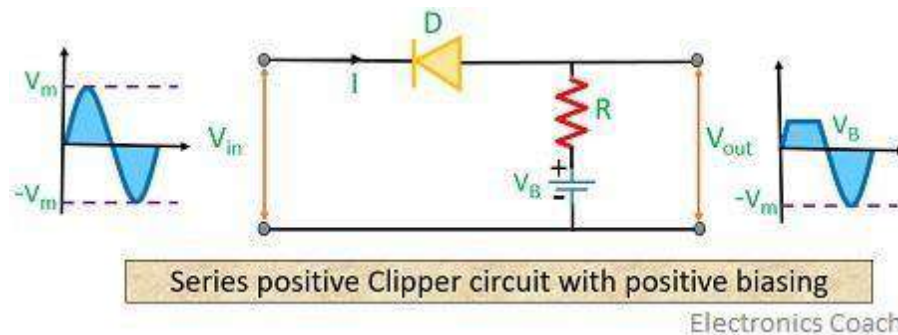
Here, during the positive half cycle of input waveform, the diode becomes forward biased, thus ensuring a closed circuit. Due to which current appears across the resistor of the circuit.

For negative half of the input waveform, the diode now becomes reverse biased acting as an open switch. This causes no current to flow through the circuit. Resultantly providing no output for negative half of the input waveform.

- **Series positive clipper circuit with bias**

Whenever there is a need to clip or remove a certain portion of positive half of input waveform, series positive clippers with biasing are needed.

1. The case of a Positively biased circuit:

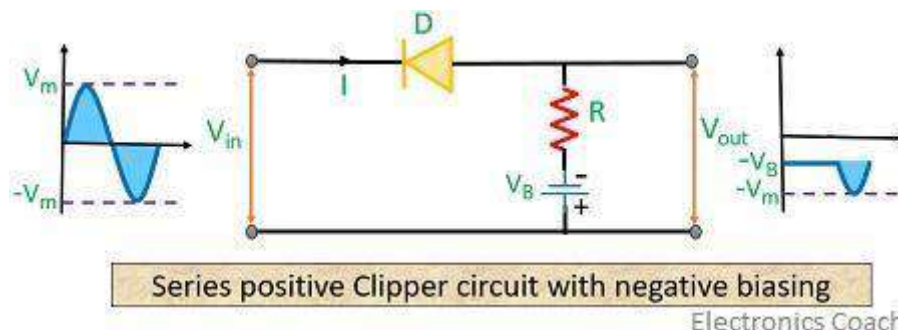


Here in the circuit shown above, we can see that the diode is in forward bias condition concerning the battery. But positive half of the input waveform puts the diode in reverse biased condition.

The diode will conduct until the supply voltage is less than the battery potential. As battery potential dominates the supply voltage, the signal appears at the positive half of output waveform. But as the supply voltage exceeds the battery potential, the diode is now reverse biased. Resultantly no further current will flow through the diode.

For the negative half cycle of the input waveform, the diode is forward biased concerning both supply voltage and battery potential. Hence, we achieve a complete negative half cycle at the output waveform.

2. The case of a negatively biased circuit:



As we can see in the circuit shown above, the diode is reverse bias due to both supply voltage and battery potential. This cuts off the complete positive half of the input waveform.

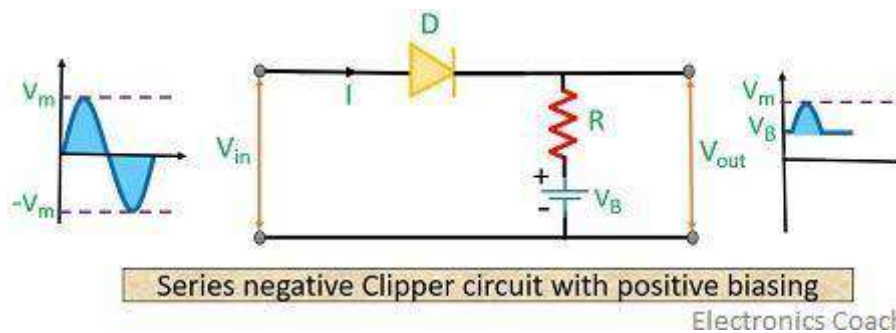
But during the negative half cycle of the input waveform, the diode is in forward biased condition due to supply voltage but is reverse biased by the battery potential.

Here also initially when battery dominates the supply voltage, the diode is in reverse biased condition. But, as the supply voltage becomes greater than the battery potential, the diode will automatically come in forward biased condition. Thus, the signal starts to appear at the output.

- **Series negative clipper circuits with bias**

Whenever there is a need to clip or remove a certain portion of the negative half of the input waveform, then series negative clipper circuits with biasing is needed.

1. The case of a positively biased circuit:

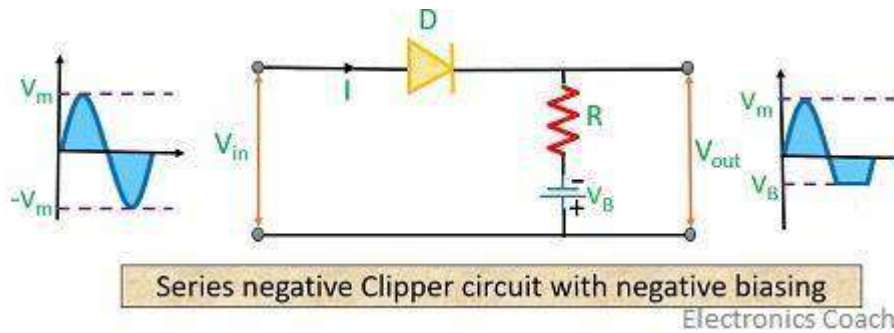


Here in the circuit, we have noticed that for the positive half of input waveform, the diode becomes forward bias. However, the battery potential causes the diode to be at the reverse biased condition.

In the beginning, the battery potential is higher than the supply voltage. Thus causing no current to flow through the circuit at that time. But when the positive half supply voltage exceeds the battery potential then diode becomes forward biased and starts conduction. Hence signal appears at the output.

The negative half cycle of the input waveform causes the diode to be reverse biased. At the same time applied battery potential reverse biases the diode. Thus no any signal is achieved at the output in such a condition.

2. The case of negatively biased circuits:



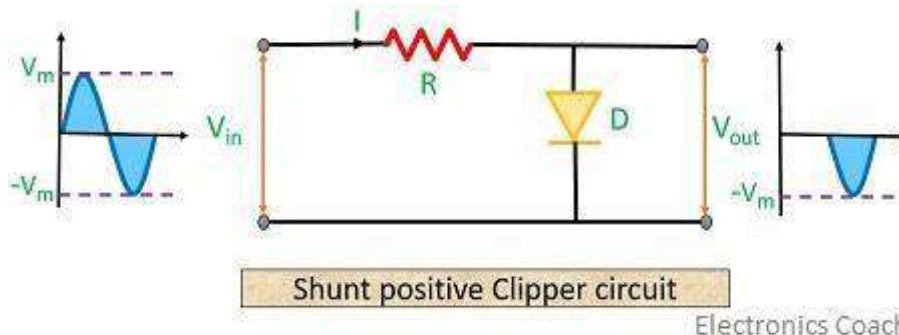
For the positive half of the input cycle, supply voltage and battery potential both cause the diode to be in forward biased condition. Thus complete output for the positive half is achieved in this case.

But in the case of the negative half cycle of the input waveform, the diode is now in reverse biased condition. At the same time due to battery potential, the diode is forward biased. So, until the battery potential is greater than the supply voltage, current flows through the circuit. Hence output is achieved.

But as the supply voltage becomes greater than the battery potential, the diode will now become reverse biased and hence conduction stops. Ultimately no any signal appears for that portion of the input waveform at the output.

- Shunt positive Clipper circuits

Now, let's have a look at the circuit of shunt clippers shown below



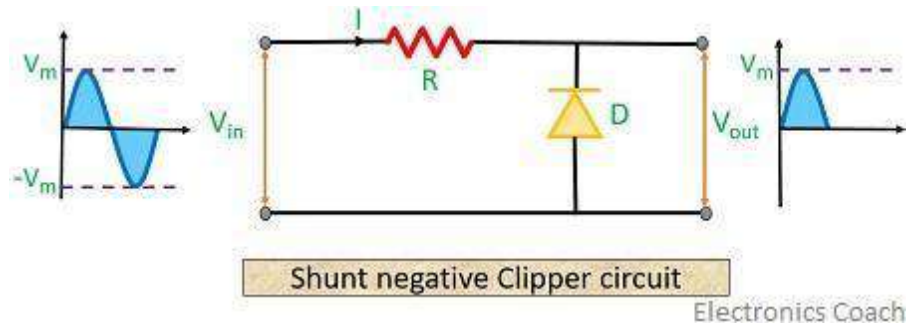
Here, the diode is connected in parallel with the load. Hence its working principle is exactly opposite to that of the shunt clippers. In shunt positive clippers, the output signal is observed only when the diode is reverse biased.

During the positive half of the input signal, the diode gets forward biased due to which the current flows through the diode. Hence, due to the parallel combination of diode and load, no current is observed at the load. Resultantly no output waveform for the positive half is achieved.

On the contrary, during the negative half of the input signal, the diode gets reverse biased. Thus no current flows through it, and the output current is observed at the load. So, for the negative half of input, the entire negative half appears at the output.

- **Shunt negative Clippers circuits**

Let's have a look at the figure shown below of shunt negative clippers.



For negative shunt clippers, during the positive half of input, the diode gets reverse biased. Thus no current flows through it, and the output current is observed at the load.

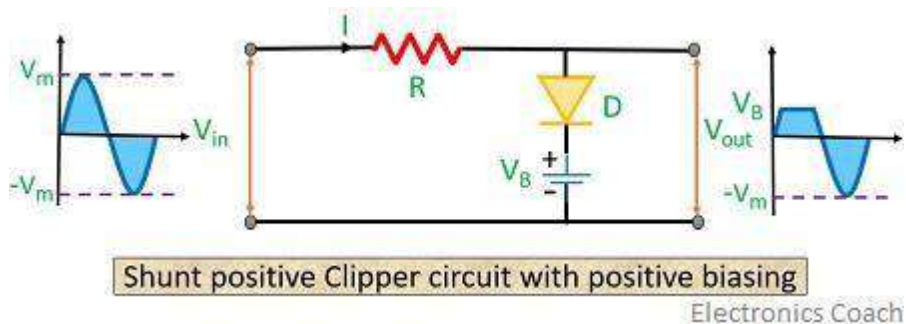
Hence output signal is achieved for positive half of the input signal.

During the negative half of the input signal, the diode gets forward biased and hence no load current is achieved. Ultimately no output is observed for negative half of the input signal.

- **Shunt positive clipper circuits with bias**

Here, we will discuss both the cases of positive and negative biasing separately as we have done in series clipper with bias.

1. **The case of positively biased circuits:**

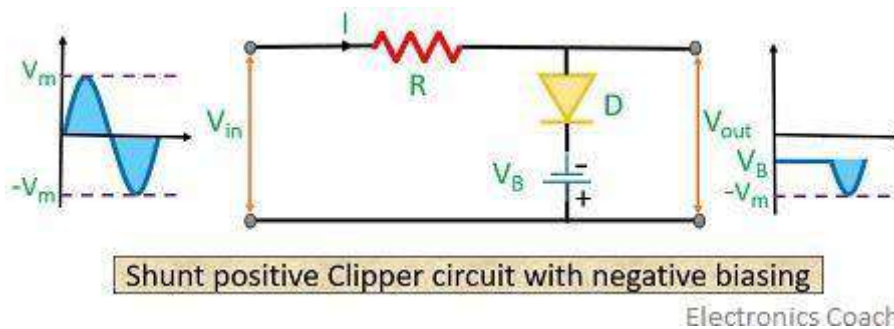


During the positive half of the input cycle, the diode gets forward biased but at the same time battery potential causes the diode to be at the reverse biased condition.

Until the battery potential is greater than the supply voltage, the reverse biased diode by the battery potential does not allow the flow of current through it. Hence current flows through load and signal is observed at the output. But when input voltage increases more than the battery potential, the diode becomes forward biased due to the supply voltage. Hence, allowing current to flow through the diode and no output is achieved.

In case of negative half of the input cycle, both battery potential and supply voltage reverse biases the diode. Hence we get a complete negative half cycle at the output.

2. The case of negatively biased circuits:



In this case, the negative half of the input cycle and the applied battery potential both causes the diode to be forward biased. Thus current flows through the diode, and no signal is observed across the output.

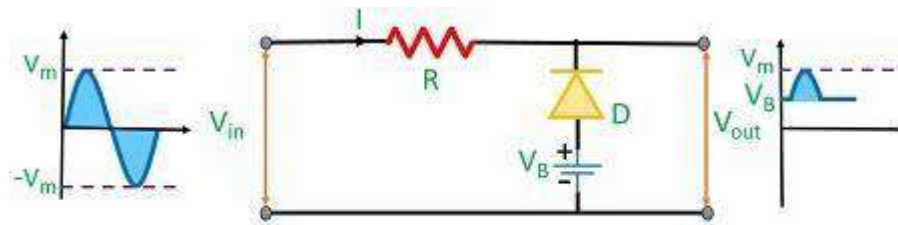
During the negative half of the input cycle, the diode is reverse biased due to supply voltage but is in the forward biased state due to battery potential.

So, here output will not appear at the load until battery potential is higher than the supply voltage. As the supply voltage exceeds the battery potential, the diode stops conduction due to the reverse bias condition. Hence signal appears at the output.

- **Shunt negative clipper circuits with bias**

Now, let's move further and discuss the case of positive and negative biasing separately.

1. The case of positively biased circuits:



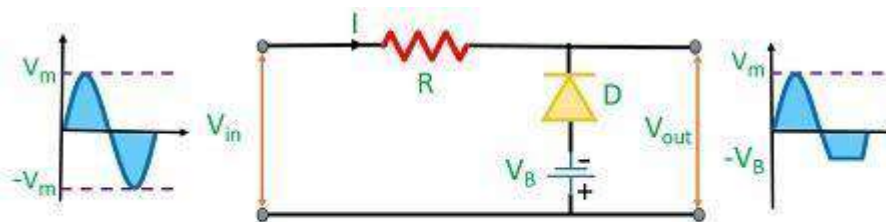
Shunt negative Clipper circuit with positive biasing

Electronics Coach

In case of positive half of the input signal, the diode gets reverse biased. However, the battery potential causes the diode to be in forward biased condition. So, until the battery potential is higher than the supply voltage, the diode conducts due to forward bias. But as the supply voltage exceeds the battery potential the conduction through diode stops. Ultimately signal appears at the output.

During the negative half of the input cycle, the diode gets forward biased because of battery potential and supply voltage. Hence, no output signal appears for the negative half of the input cycle.

2. The case of negatively biased circuits:



Shunt negative Clipper circuit with negative biasing

Electronics Coach

During the positive half of the input cycle, the diode gets reverse biased due to both battery potential and supply voltage. Resultantly, we have complete positive half of input signal at the output.

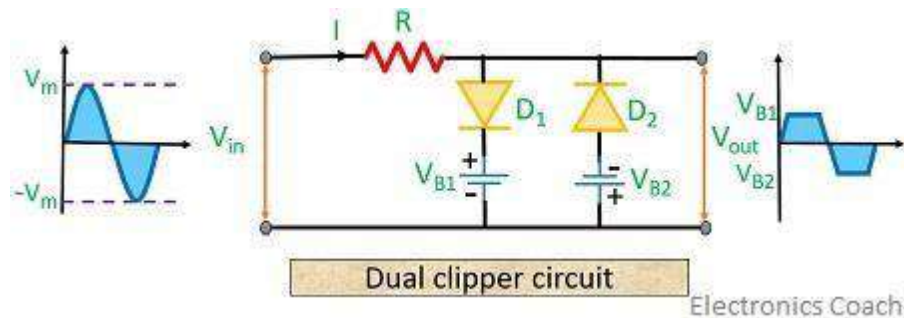
During the negative half of the input cycle, the diode gets forward biased due to supply voltage but is reverse biased by the battery potential. Thus, the output signal is achieved until the battery potential is higher than the supply voltage.

As the supply voltage exceeds the battery potential, the diode completely starts conduction and signal does not appear at the output.

- **Dual clipper circuits**

Whenever there is a need to remove a certain portion of both positive and negative half of the signal, then dual clipper circuits are used.

Let's have a look at the circuit of dual clipper shown below-



During positive half of the input cycle, diode D_1 gets forward biased due to supply voltage but is in reverse bias state due to battery potential V_{B1} . At the same time, diode D_2 is in reverse bias due to both supply voltage and battery potential V_{B2} .

Until battery voltage exceeds the supply voltage D_1 will be in reverse-biased state and D_2 is already in reverse bias condition. Hence signal is achieved at the output. But as the supply voltage exceeds the battery potential, it causes diode D_1 to be forward biased. Hence, no any further signal for the positive half is achieved in this case.

During the negative half of the input cycle, diode D_1 will be reverse biased due to both supply voltage and battery potential.

Conversely, diode D_2 will be forward biased by the supply voltage but is reverse biased by the battery potential V_{B2} .

Until the supply voltage is less than the battery potential, D_2 will be in reverse-biased state and signal appears at the output in this condition. But, as the supply voltage exceeds the battery potential, D_2 will be forward biased. Hence no any signal is obtained at the output.

Applications of Clipper circuits

These are used in numerous applications such as in transmission, in overvoltage protection, in the modification of input waveform, in voltage limiting etc.

P-N DIODE CLAMPER CIRCUIT:

Clamper circuits are the electronic circuits that shift the dc level of the AC signal. Clampers are also known as DC voltage restorers or level shifter. Clampers are basically classified as positive and negative that includes both biased and unbiased conditions

individually. These circuits are used to clamp an input signal to a different dc level. It basically adds dc component to the applied input signal in order to push the signal to either the positive or negative side. Clamper circuit is a combination of a resistor along with a diode and capacitor. It sometimes also employs dc battery so as to have an additional shift in the signal level.

Clamper circuits are constructed in a similar manner as that of clipper circuits. However, clamper includes an extra charging element that is the capacitor in its circuitry. The combination of resistor and capacitor in the clamper circuit is used to maintain different dc level at the output of the clamper.

Operating principle of Clamper circuits

As we have already discussed a clamper consist of the capacitor and a diode in shunt connection with the load.

The working of clamper circuits depends on the variation in the time constant of the capacitor. This variation is the outcome of changing the current path of the diode with the change in input signal polarity.

Here, the magnitude of the time constant is

$$\tau = RC$$

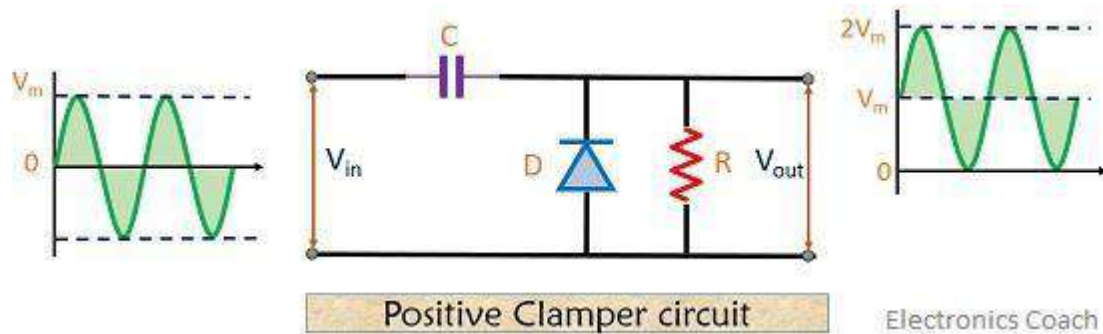
this is chosen large enough in order to assure that voltage across the capacitor does not discharge consequently at the non-conducting interval of the diode. But, such a discharge takes place only when the load resistance is very large. This permits the capacitor to take larger discharge time. Conversely, a smaller value of the capacitor is chosen so that it will charge rapidly at the time of conduction of the diode.

Classification of Clamper Circuits-

Clampers are classified in the following groups:

Positive Clamper circuit

The figure below shows the circuit of a positive clamper-



As we can see here, the diode is in parallel connection with the load. So we can say reverse biasing of the diode will provide the output at the load.

Initially, the positive half of the applied input signal reverse biases the diode but the capacitor is not still charged. So, at this period of time output will not be considered.

For, the negative half of the AC signal, the capacitor now gets fully charged up to the peak of the AC signal but with inverse polarity. This negative half forward biases the diode that results in the flow of the forward current through the diode. The next positive half then reverse biases the diode due to which signal will appear at the output.

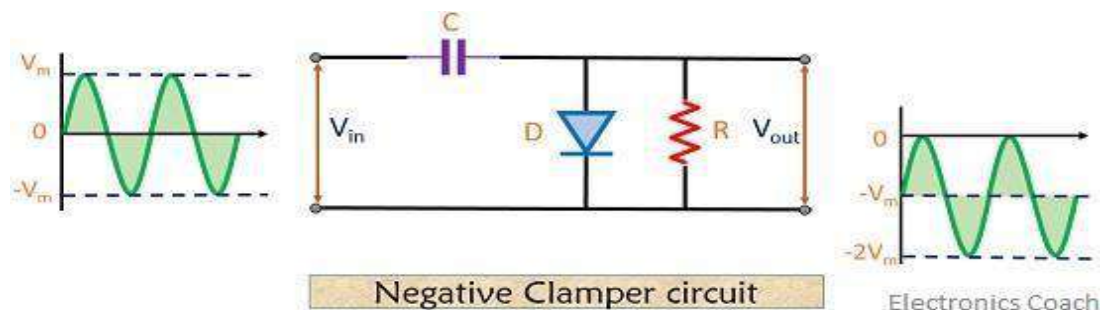
At the beginning of the positive half of the AC signal, the diode is in the non-conducting state that results in discharging of capacitor charge. So, at the output, we will have the summation of the voltage stored across the capacitor and applied the AC input signal. This is given by

$$V_o = V_m + V_m = 2V_m$$

Here as we can in the output waveform shown above, the signal level is shifted upward or positive side. Hence it is named as positive clamper.

Negative Clamper circuit

Let's have a look at the figure shown below of negative clamper in order to understand the detailed operation-



At the time when positive half of the AC input is applied, the diode comes to forward bias condition that results in no-load current at the output. However, a forward current flows through the diode that charges the capacitor to the peak of the ac signal but again with inverse polarity. The capacitor here is charged up to the forward biased condition of the diode.

When negative half of the AC signal is applied, the diode now becomes reverse biased. This allows load current to appear at the output of the circuit. Now, this non-conducting state of the diode discharges the capacitor. So, at the output, a summation of capacitor voltage along with the input voltage is achieved.

Hence at the output, we have,

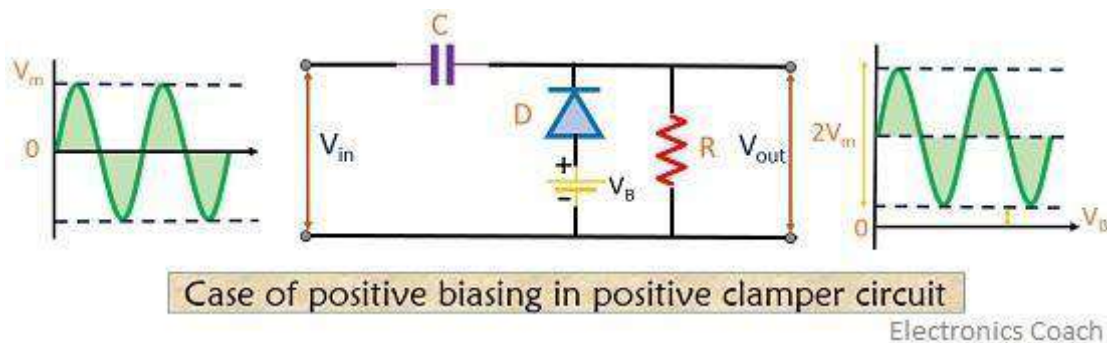
$$V_o = -V_m - V_m = -2V_m$$

This results in the downward shift of the signal. Therefore, it is termed as negative clamper circuit.

Positive clamper circuit with biasing

This is basically done to introduce an additional shift in the level of the signal. Here, the biasing provided to the circuit is of two types. It can be a positively or negatively biased circuit. So, we will discuss both the cases separately.

1. Case of positive biasing



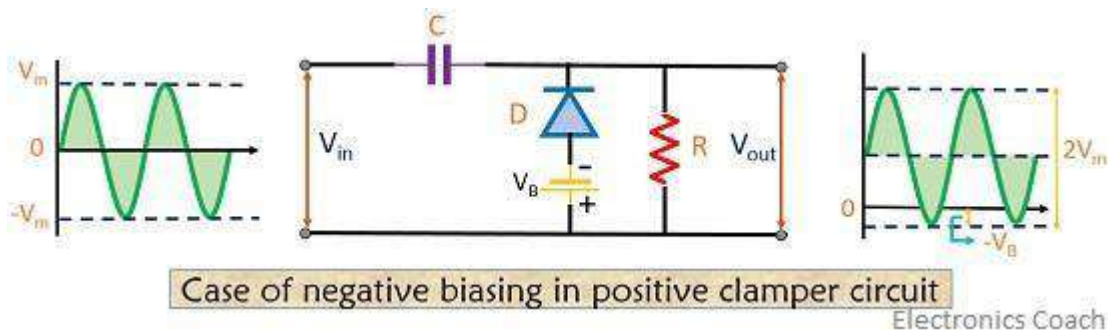
The working is almost similar to the positive unbiased case but here an additional voltage is provided so as to have an additional shift in the level of the signal.

When positive half of the input signal is applied, the diode is reverse biased due to ac input but is forward biased due to battery voltage. So, until the voltage of a battery is greater than the ac input the diode conducts. This forward current through the diode charges the

capacitor but with the battery voltage. As the ac input surpasses battery voltage, the diode now gets reverse biased and hence conduction through the diode stops.

On the application of the negative half of the input signal, the diode is now forward biased due to both ac input and battery voltage and starts conducting. This charges the capacitor with voltage summation of ac input along with battery voltage. Hence such an output voltage level is achieved.

2. Case of negative biasing



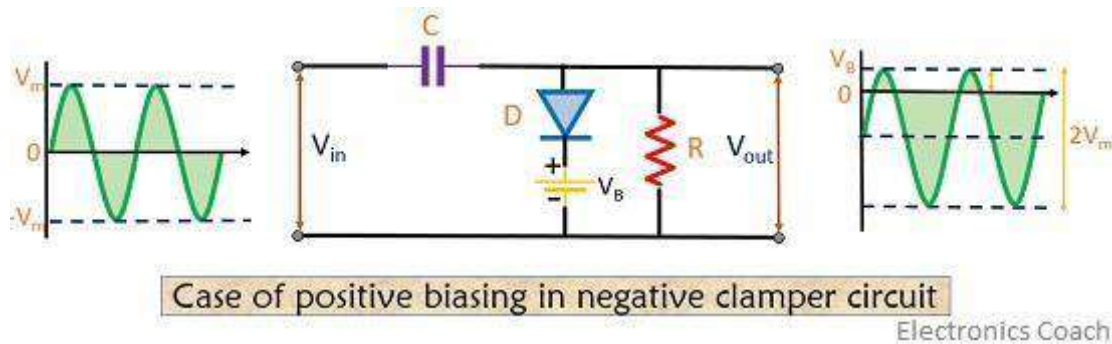
At the time of positive half of the AC signal, the diode gets reverse biased by both ac input and battery voltage. Due to this current flows through the load and combinedly maintain the voltage level.

At the time of the negative half, the diode is in the forward biased condition due to ac input but is in reverse biased condition due to battery voltage. So, the diode conducts only when the AC input dominates the battery voltage. This charges the capacitor hence we get a shifted signal at the output.

Negative Clamper circuit with biasing

In a similar way as previous, positive and negative biasing is provided to the negative clamper circuit. Let's now move further and discuss both the case separately.

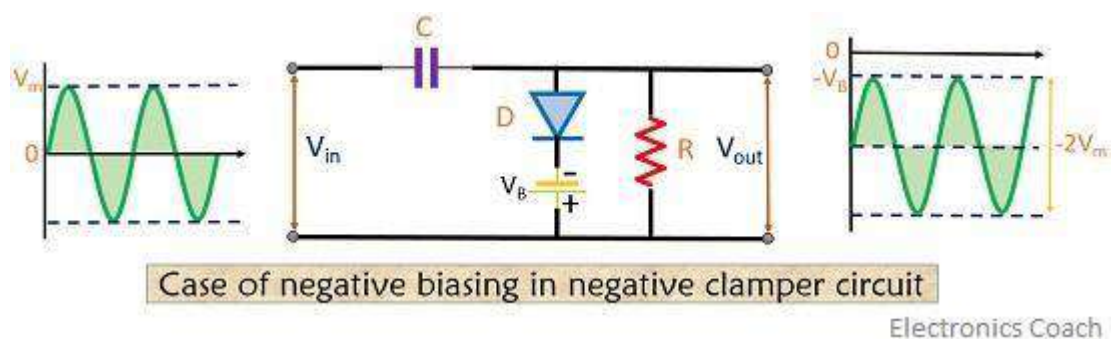
1. Case of positive biasing



As we have already discussed that a negative clamper shifts the signal downward. However, in the case of the positively biased negative clamper, the signal is somewhat raised to a positive level due to the positively applied battery voltage. When positive half of the AC signal is applied, the diode is in the forward biased state due to ac supply but is reverse biased because of battery voltage. So, the diode conducts when ac supply surpasses battery voltage.

Moving further during the negative half, the diode is now in reverse biased state by cause of both the AC supply and battery voltage. This non-conducting state of the diode discharges the capacitor. Thus, the voltage across the capacitor appears at the output.

2. Case of negative biasing



At the time of positive half of ac input, the diode gets forward biased by the cause of ac input and battery voltage. This starts conduction through the diode. Resultantly charges the capacitor.

At the time of the negative half, the diode gets reverse biased but will still conduct due to forward biased condition applied by the battery. The diode current flows until the battery voltage is more than the ac input supply. The time when ac input surpasses battery voltage, the diode gets reverse biased and the capacitor discharges. Thus the voltage across capacitor appears at the load.

Applications of Clamper Circuits

1. Clippers are used to identify the polarity of the circuits.
2. These circuits are used as voltage doublers and help in eliminating distortions.
3. Reverse recovery time can be improved using Clippers.
4. Comparison Chart

PARAMETER	CLIPPER	CLAMPER
Definition	Clipper delimit the amplitude of the output voltage.	Clamper shifts the DC level of the output voltage.
Output Voltage	Less than the input voltage.	Multiples of input voltage.
Energy storage component	Not required	Requires (Capacitor is used as energy storage element)
Shape of Output Waveform	Shape changes (Rectangular, sinusoidal, triangular etc.)	Shape remains same as input waveform.
DC Level	Remains same	DC level get shifted
Applications	In transmitters, receivers, amplitude selector, noise limiter etc.	In voltage multiplying circuits, Sonar, Radar system etc.

