

الالكترونيات التماثلية

المرحلة الثانيه

الكورس الاول

ا.م.د. استبرق طالب عبدالله

Lecture 1

Chapter One

Introduction to Semiconductors

1- Introduction

Electronic devices such as diodes, transistors, and integrated circuits are made of a semiconductive material. To understand how these devices work, you should have a basic knowledge of the structure of atoms and the interaction of atomic particles. An important concept introduced in this chapter is that of the *pn* junction that is formed when two different types of semiconductive material are joined. The *pn* junction is fundamental to the operation of devices such as the solar cell, the diode, and certain types of transistors.

All matter is composed of atoms; all atoms consist of electrons, protons, and neutrons, according to Bohr model, except normal hydrogen, which does not have a neutron. Each element in the periodic table has a unique atomic structure, and all atoms for a given element have the same number of protons. The **nucleus** consists of positively charged particles called **protons** and uncharged particles called **neutrons**. The basic particles of negative charge are called **electrons**, as illustrated in Fig. 1.

Electrons that are in orbits farther from the nucleus have higher energy and are less tightly bound to the atom than those closer to the nucleus. This is because the force of attraction between the positively charged nucleus and the negatively charged electron decreases with increasing distance from the nucleus. Electrons with the highest energy exist in the outermost shell of an atom and are relatively loosely bound to the atom. This outermost shell is known as the **valence** shell, and electrons in this shell are called *valence electrons*. **These valence electrons contribute to chemical reactions and bonding within the structure of a material and determine its electrical properties**. When a valence electron gains sufficient energy from an external source, it can break free from its atom. This is the basis for conduction in materials.

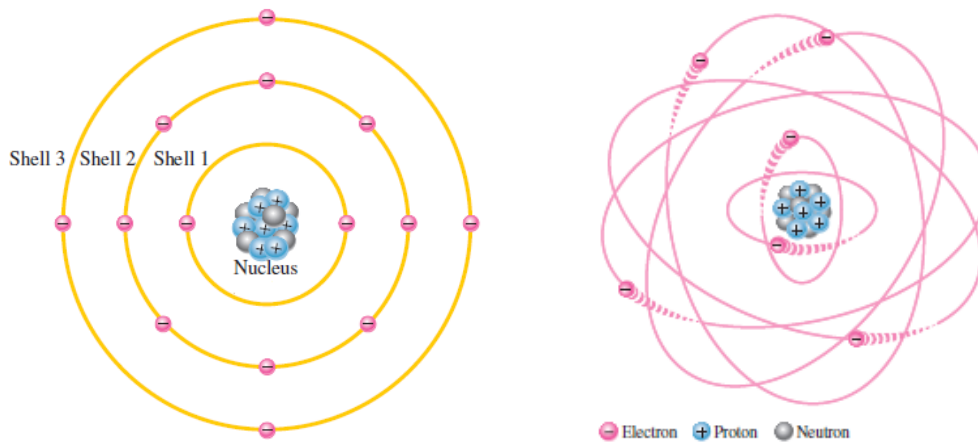


Fig.1: The Bohr model of an atom showing electrons in orbits around the nucleus, which consists of protons and neutrons.

2- Insulators, Conductors, and Semiconductors

In solid materials, interactions between atoms “smear” the valence shell into a band of energy levels called the *valence band*. Valence electrons are confined to that band. When an electron acquires enough additional energy, it can leave the valence shell, become **a free electron**, and exist in what is known as the **conduction band**. The difference in energy between the valence band and the conduction band is called an **energy gap** or **band gap**.

In terms of the electrical properties, materials can be classified into three groups: conductors, semiconductors, and insulators. They are differing in their conductivity depending on the no. of electrons in their outer shell (valance electrons).

Insulators: An **insulator** is a material that does not conduct electrical current under normal conditions. Valence electrons are tightly bound to the atoms; therefore, there are very few free electrons in an insulator (**8 electrons in** outer shell). **It has full valence band, empty conduction band and large energy gap**. Examples of insulators are rubber, plastics, glass, mica, and quartz.

Conductors: A **conductor** is a material that easily conducts electrical current. Most metals are good conductors, such as copper (Cu), silver (Ag), and gold (Au), which are characterized by atoms with only one valence electron very loosely bound to the atom (1 electron in outer shell). There is an overlap between the two bands, no physical distance between the two (no energy gap). Hence there are large numbers of conduction electrons.

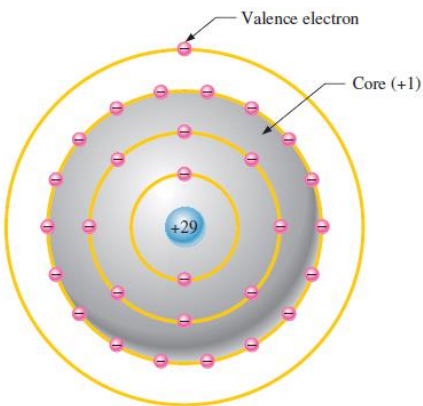


Fig. 2: Copper atom

Semiconductors: A semiconductor material is one whose electrical properties lie in between those of insulators and conductors. Examples are: germanium (Ge) and silicon (Si). In terms of energy bands, semiconductors can be defined as those materials which have almost an empty conduction band and almost filled valence band with a very narrow energy gap (of the order of 1 eV) separating the two. For pure semiconductor at 0°K, there are no electrons in the conduction band, and the valence band is filled, i.e. good insulator. As temperature increases, some electrons are move into the conduction band, leaving behind an empty space (hole) in the valence band. So the electron-hole pair formation.

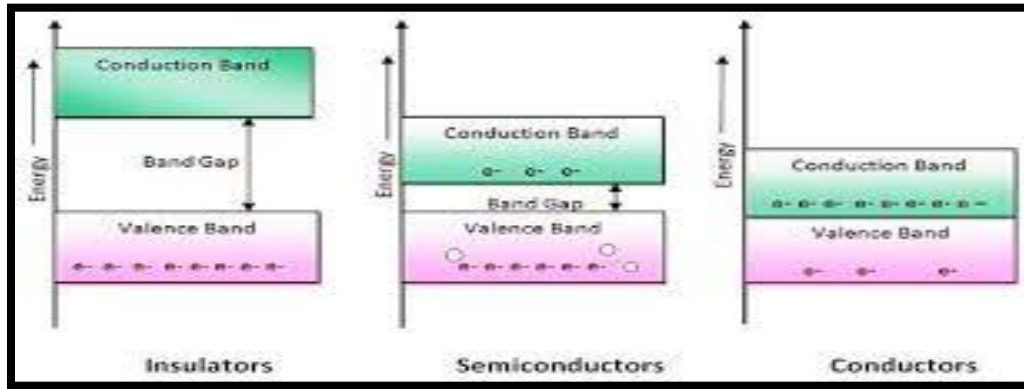
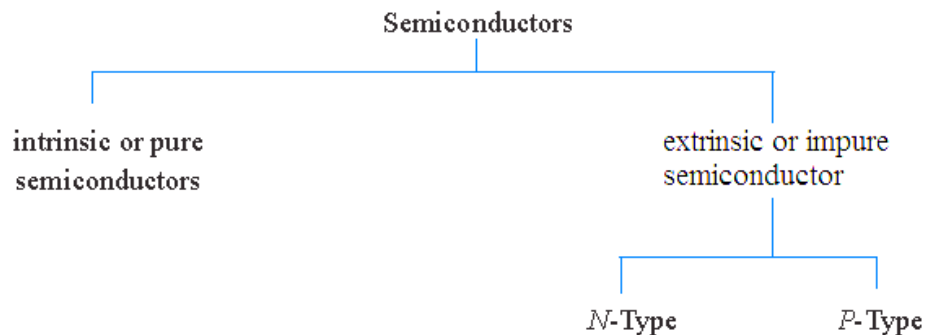


Fig.2 Diagram of energy band.

Types of Semiconductors:

Semiconductor may be classified as under:



a. Intrinsic Semiconductors

An intrinsic semiconductor is one which is made of the semiconductor material in its extremely pure form. Examples of such semiconductors are: pure germanium (Ge) and silicon (Si) which have forbidden energy gaps of 0.3 eV and 0.7 eV respectively. The energy gap is so small that even at ordinary room temperature; there are many electrons which possess sufficient energy to jump across the small energy gap between the valence and the conduction bands.

Alternatively, an intrinsic semiconductor may be defined as one in which the number of conduction electrons is equal to the number of holes. Schematic energy band diagram of an intrinsic semiconductor at room temperature is shown in Fig. below.

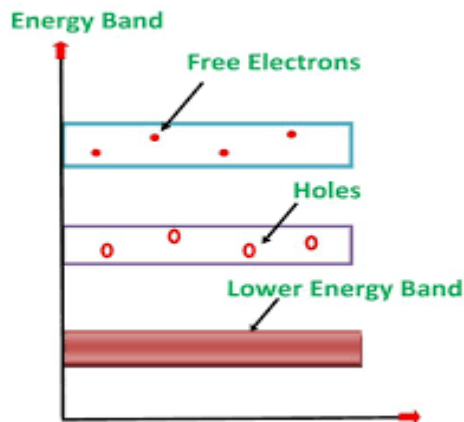


Fig.3 diagram energy band for intrinsic semiconductor

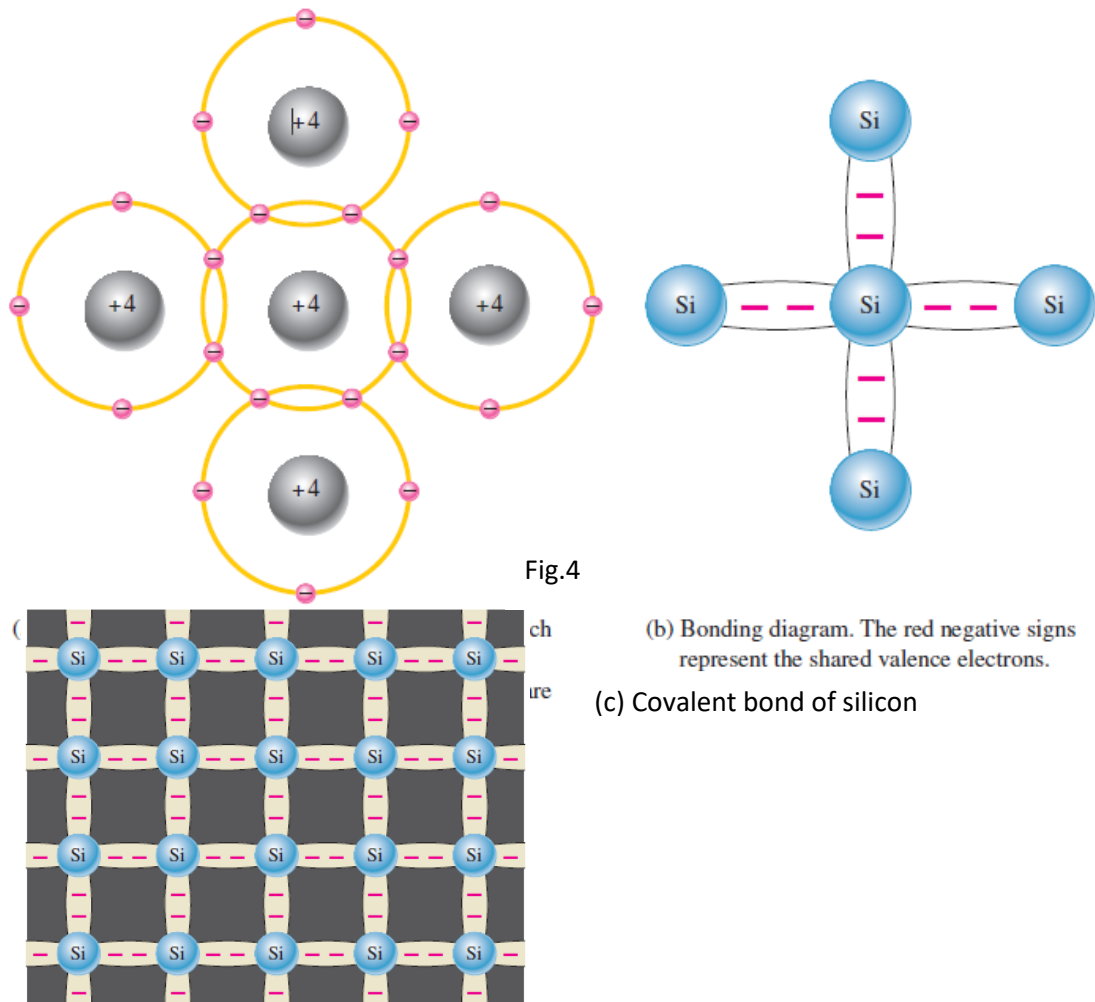
Crystal structure of Si

The outermost orbit of a Si atom is capable to hold up to eight electrons. The atom which has eight electrons in the outermost orbit is said to be completely filled and most stable. The outermost orbit of silicon has only four electrons. It needs four more electrons to fill the outer orbit and so become most stable. So, for Si to become stable, it will share four electrons (by forming four covalent bonds) with four neighbouring atoms. The outermost orbit of silicon is, then, completely filled and electrons are tightly bound to the nucleus (Fig. 4).

In intrinsic semiconductors, at absolute zero temperature, free electrons are not present (since all electrons are engaged in covalent bonds). Therefore, intrinsic semiconductor behaves as perfect insulator.

As temperature increases the electrons in the covalent bond (in the valence band) can gain enough energy to break the bond and become a free electron (in the conduction band). Its place is left without an electron, which is the hole. Hence the generation of an electron-hole pair. These free electrons are free to move randomly

within the lattice, till their energy is consumed and they will fall back into a hole. This is the **recombination process**.



So, in semiconductors there are two types of charge carriers (electrons and holes), as voltage is applied across an intrinsic semiconductor, the holes move towards the negative end of the battery, while the electrons move towards the positive end of the battery, see fig.5.

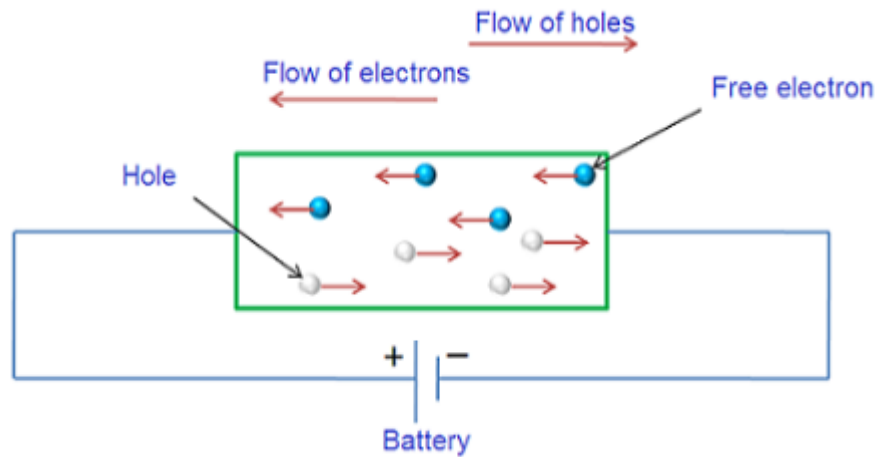


Fig.5: electrons and holes movement under applied voltage.

The ability of an electron to move through a metal or semiconductor, in the presence of applied electric field is called **electron mobility**. Mobility of electron is larger than hole mobility.

Note that

- The electron current is in opposite direction to the hole current.
- Holes do not actually move in the semiconductor. An electron within the valence band may fill the hole, leaving another hole in its place. In this way a hole appears to move.

Summary

- Si and Ge are examples of semiconductor materials.
- Semiconductors are characterized by their small energy band gap.
- At absolute zero, the semiconductor behaves as an insulator.
- As temperature increases, the electrical conductivity increases.
- The increase of temperature produces electron- hole pairs (the number of electrons are equal to the number of holes).

- In semiconductors, there are two type of charge carriers: electrons and holes.
- On applying a potential difference across a semiconductor, two currents are generated the electron current and the hole current, which are in opposite directions.

b. Extrinsic Semiconductors

Intrinsic silicon (or germanium) must be modified by increasing the number of free electrons or holes to increase its conductivity and make it useful in electronic devices. To increase this, very small amounts of dopants (pentavalent or trivalent elements) (impurities) are added to the intrinsic semiconductor. This process is called **doping**.

- ❖ pentavalent elements: have 5 electrons in their outermost orbit.
- ❖ Trivalent elements are those that have three electrons in their outermost orbit.

The semiconductor is said to be **heavily doped** when a large amount of dopants is added. In contrast, it is **lightly doped** when a small amount of dopant is added.

(i) N-type Extrinsic Semiconductor:

If a silicon or germanium atom in its pure form is doped with pentavalent element such as antimony (Sb), arsenic (As) or phosphorus (P). These elements having 5 electrons in their outermost shell react such that they form a covalent bond with the four electrons of silicon, and one electron is left free as a mobile charge carrier, which improves the conduction ability to some extent. The resultant material is known as an n-type semiconductor. Fig. 6.

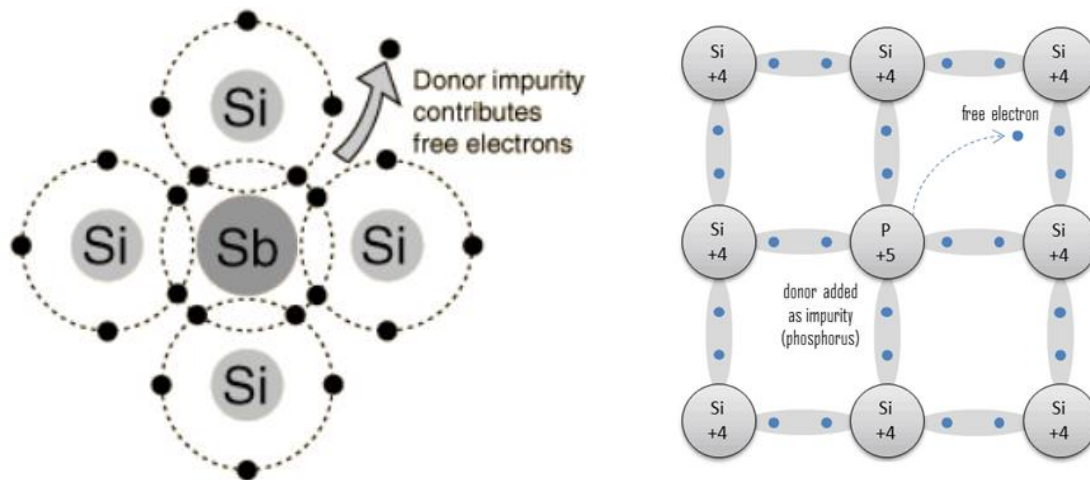


Fig.6 : n-type extrinsic semiconductor

For one dopant atom there will be one free electron. As the number of dopant atoms increases the number of free electrons will increase. The dopant atom is called the **donor atom**, since it donates a free electron to the semiconductor, which is the acceptor. In addition, thermally generated electron-hole pairs are also present in the doped semiconductor. The number of holes (in the valence band) is less than the number of free electrons (in the conduction band).

The electrons are the negative charge carriers; they are called the majority charge carriers. Holes are the positive charge carrier, are called the minority charge carriers.

The majority charge carriers(electrons) are produced by doping, while the minority charge carriers(holes) are due to the increase of temperature.

(ii) P-type Extrinsic Semiconductor:

If a silicon or germanium atom in its pure form is doped with an element of group three in a small amount, such as indium (In), gallium (Ga), Aluminum (Al)

or boron (B), these elements having 3 electrons in their outermost shell react such that they form a covalent bond with the three electrons of silicon, and one hole is left free as a mobile charge carrier, which improves the conduction ability to some extent. The resultant material is known as a p-type semiconductor. Fig.7.

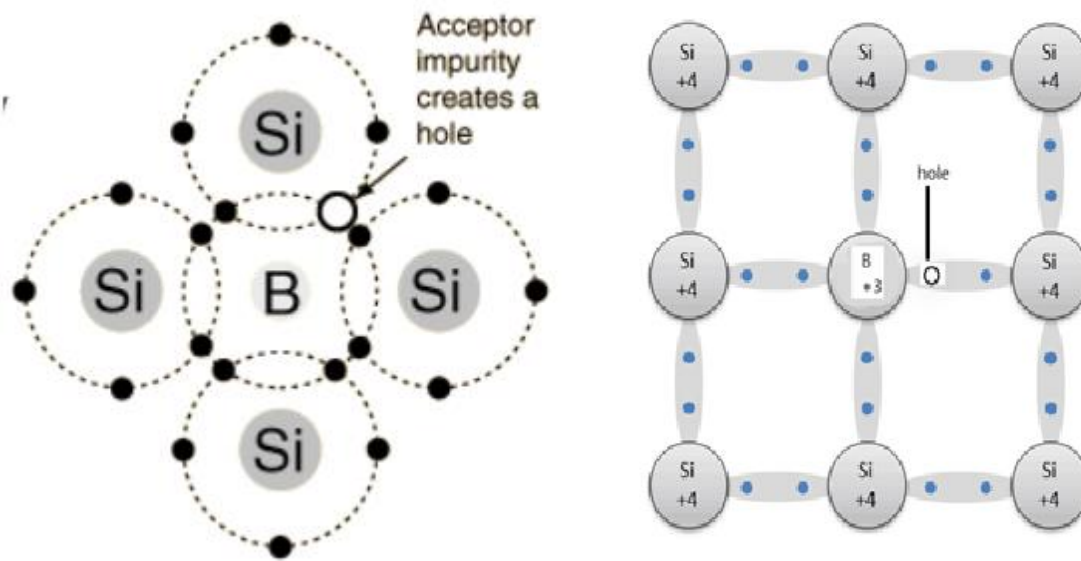


Fig7 : P-type extrinsic semiconductor

For one dopant atom there will be one hole. As the number of dopant atoms increases the number of holes will increase. The dopant atom is called the donor atom, since it donates a hole to the semiconductor, which is the acceptor.

In addition, thermally generated electron-hole pairs are also present in the doped semiconductor. The number of holes is more than the number of free electrons. The holes are the majority charge carriers. Electrons are the minority charge carriers. The majority charge carriers (holes) are produced by doping, while the minority charge carriers (electrons) are due to the increase of temperature.

Lecture 2

Chapter Two P-N Junction (Diode)

1. Introduction

When a piece of semiconductor is doped as P-type from one side and as an N-type from the other side, a PN – junction is formed. The PN- junction can be looked at as if two pieces of semiconductor, one is N-type and the other p-type, are joined together (Fig.1).

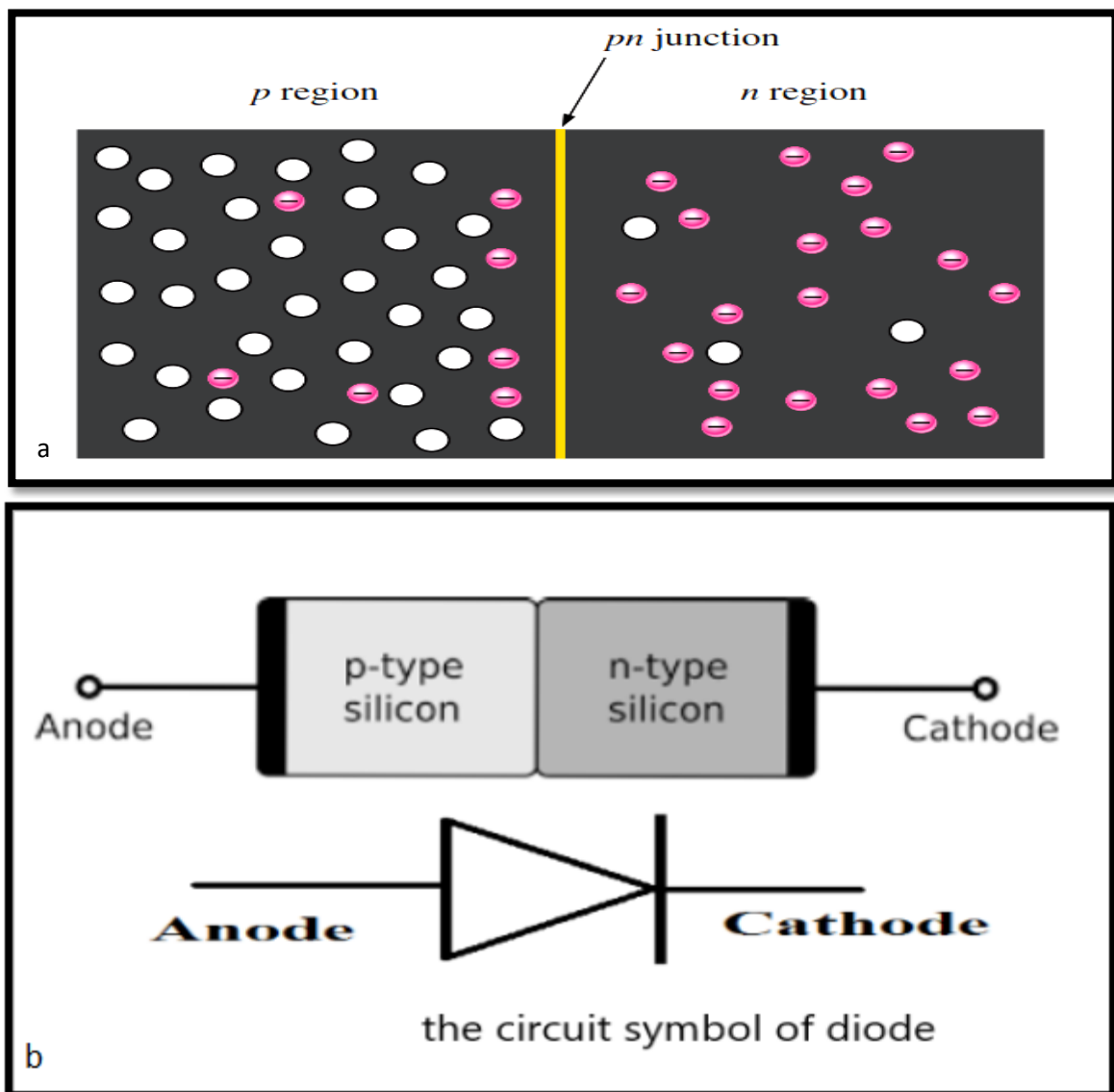


Fig. (1): the p-n junction

2. Formation of the Depletion Region

Since the N-type region has a high free electron concentration and the P-type has a high hole concentration, electrons diffuse from the N-type side to the P-type side. Similarly, holes flow by diffusion from the P-type side to the N-type side. As free electrons move from the N-type side to the P-type side, the number of free electrons in the N-type side will be reduced so positive ions are formed. Likewise, as free electrons go into the P-type side the number of electrons increases creating negative ions. So the movement of electrons and holes across the junction will create positive ions on the N-type side and negative ions on the P-type side, on both side of the junction.

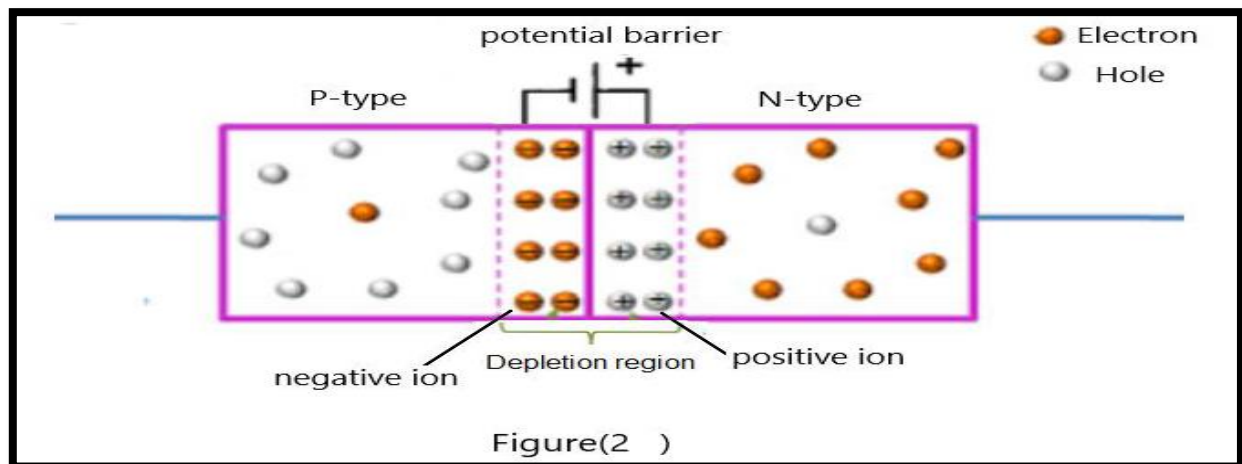


Fig. (2): Formation of the depletion region

This will set up an electric field (which is due to the attraction force between the positive and the negative ions) that opposes the transition of the majority charge carriers across the junction. The direction of this electric field is from the negative to the positive which is opposite to that of the diffusion current for each type of carrier. As the diffusion increases the number of the ions increases so does the strength of the electric field due to the formation of the ions. Equilibrium is reached (the diffusion is stopped) when the electric field is large enough to prevent further diffusion, i.e., As electrons continue to diffuse across the junction, more

and more positive and negative charges are created near the junction as the depletion region is formed. A point is reached where the total negative charge in the depletion region repels any further diffusion of electrons (negatively charged particles) into the p region (like charges repel) and the diffusion stops. In other words, the depletion region acts as a barrier to the further movement of electrons across the junction.

Therefore, at equilibrium, a region of no majority charge carriers appears on both sides of the junction. It is called the **depletion region**, for it is depleted from the majority charge carriers.

3. Barrier Potential

The potential difference of the electric field across the depletion region is the amount of voltage required to move electrons through the electric field. This potential difference is called the **barrier potential** and is expressed in volts. Stated another way, a certain amount of voltage equal (V) to the barrier potential (ϕ) and with the proper polarity must be applied across a pn junction before electrons will begin to flow across the junction. The potential barrier (its value is 0.7V for Si and 0.3V for Ge) which can be represented as a battery across the depletion region, Fig.2. The barrier potential of a pn junction depends on several factors, including the type of semiconductive material, the amount of doping, and the temperature etc.

4. Diode Biasing

Biasing is applying an external voltage source across a PN- junction. This can be done in two ways:

1-Forward biasing

2-Reverse biasing

A-Forward biasing:

This is done when the P-type side of the diode is connected to the positive terminal of the battery (the external voltage source) and the N-type side is connected to the negative terminal of the battery (as shown in Fig.3).

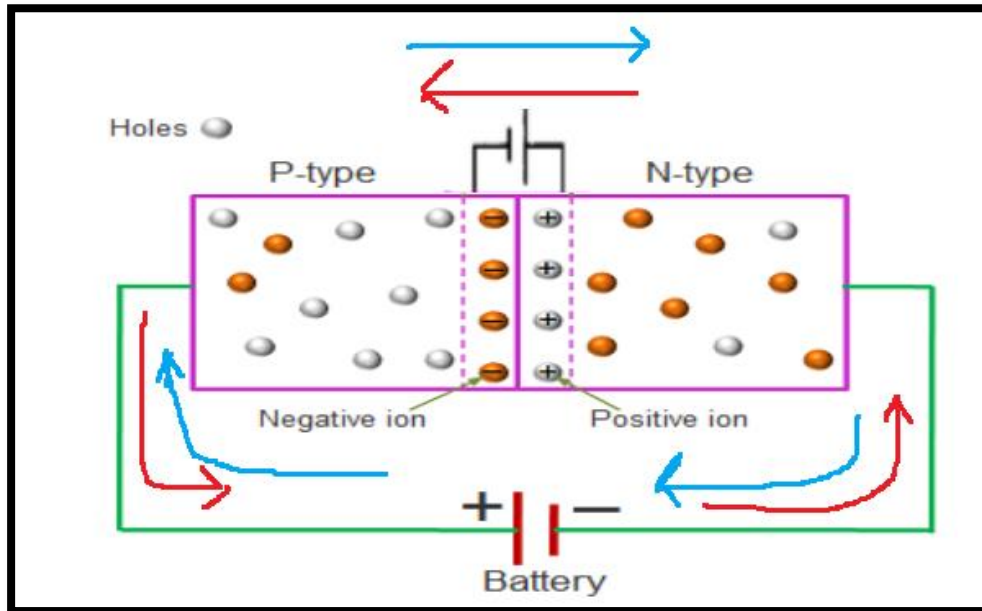


Fig.(3): Forward biasing

From the figure, it can be seen that the polarity of the battery is opposite to the polarity of the potential barrier. The potential barrier is designated by (ϕ) and the battery voltage is designated by (V). So the net voltage applied to the diode is $V - \phi$. If V starts with a small value, smaller than ϕ , the effect of the potential barrier still exists and no current passes through the diode. Current starts to flow when V is greater or equal to ϕ . This means that the applied voltage overcomes the potential barrier, and current starts to flow. The free electrons in the N-region will be attracted towards the positive terminal of the battery, crossing the junction into the P-region to the battery. At the same time, holes in the P-region are attracted towards the negative terminal of the battery. On crossing the junction, the holes and electrons neutralize the negative ions and the positive ions, respectively. This cancels the potential barrier and also reduces the width of the depletion region.

The current versus voltage can be plotted as shown below:

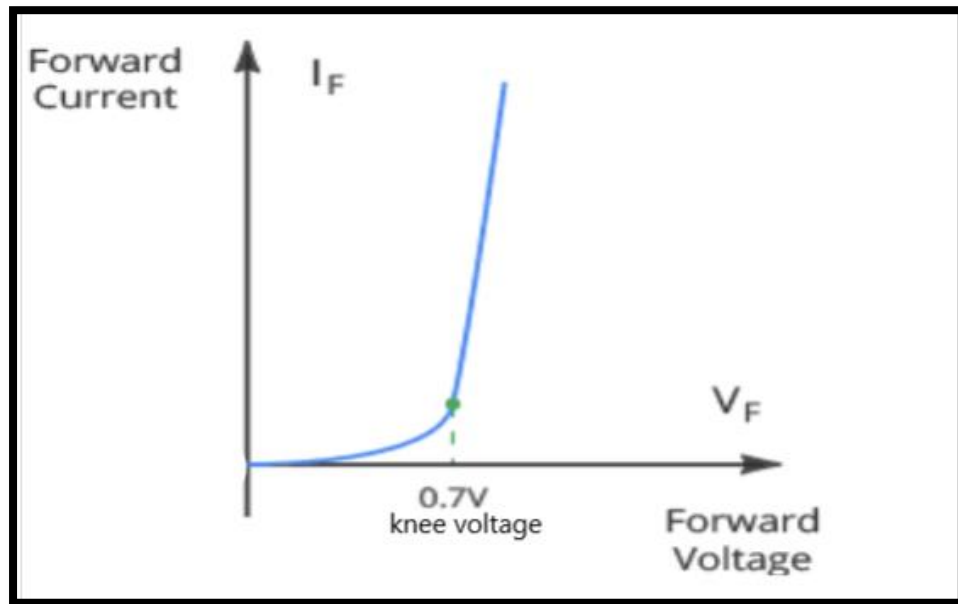


Fig. (4): IV curve of a forward biased diode

b- Reverse Biasing:

In reverse biased p-n junction diode, the positive terminal of the battery is connected to the [n-type semiconductor](#) material and the negative terminal of the battery is connected to the [p-type semiconductor](#) material. (as shown in Fig. 5).

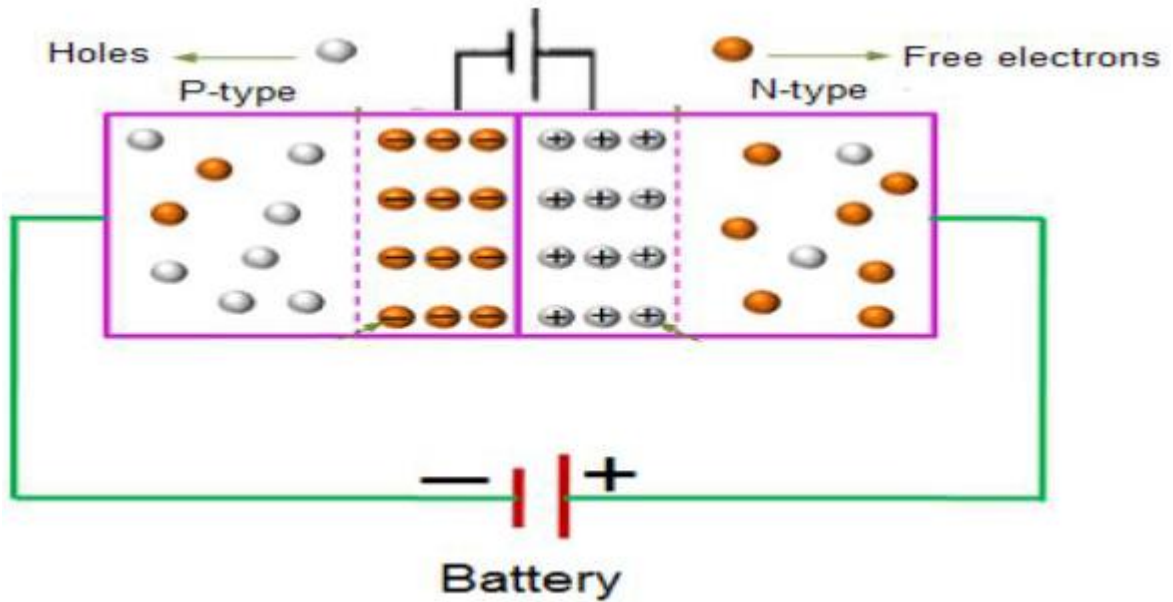


Fig. (5): Reverse Bias of PN junction

As it can be seen in the figure, the polarity of the battery is the same as the polarity of the potential barrier; this means that the voltage of the battery enforces the potential barrier. The positive voltage applied to the N-type side attracts electrons towards the positive electrode and away from the junction, while the holes in the P-type side are also attracted away from the junction towards the negative electrode. If the reverse biased voltage applied on the p-n junction diode is further increased, then even more number of free electrons and holes are pulled away from the p-n junction. **This increases the width of depletion region.** Hence, the width of the depletion region increases with increase in voltage. The wide depletion region of the p-n junction diode completely blocks the majority charge carriers. Hence, majority charge carriers cannot carry the electric current. However, a very small reverse leakage current does flow through the junction which can normally be measured in micro-amperes (μA), see fig. 6.

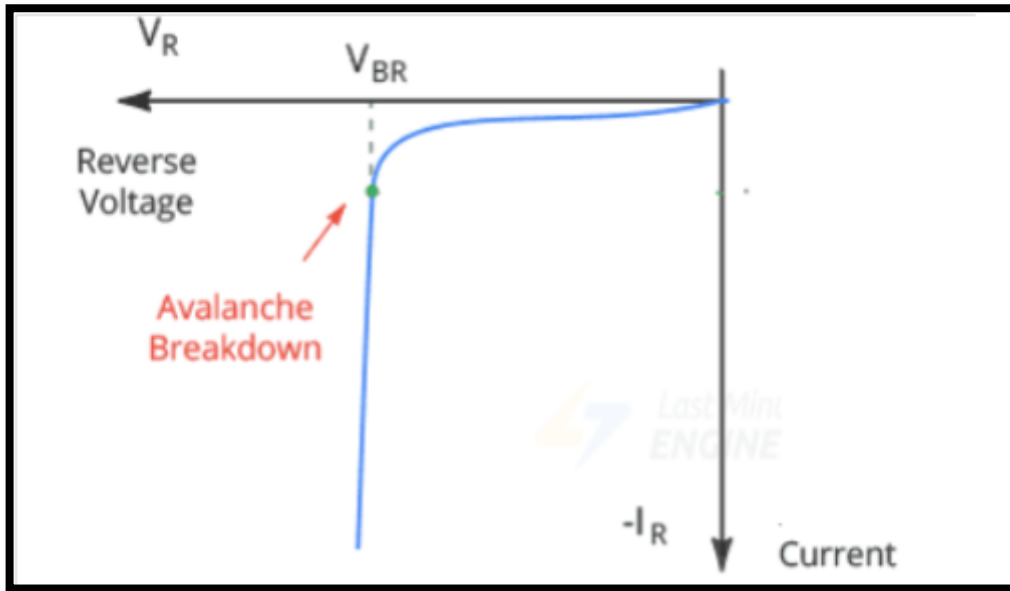


Fig. (6): IV curve of a reverse biased diode

For voltages less than the knee voltage there is no current (or very small) across the diode. The application of a forward biasing voltage on the junction diode results in the depletion layer becoming very thin and narrow which represents a low impedance path through the junction, thereby allowing high currents to flow. This is what is called the **KNEE** point of the curve (Knee voltage is the barrier potential). Continue to increase the forward-bias voltage, the current continues to increase very rapidly.

When $V < \phi$ no current,

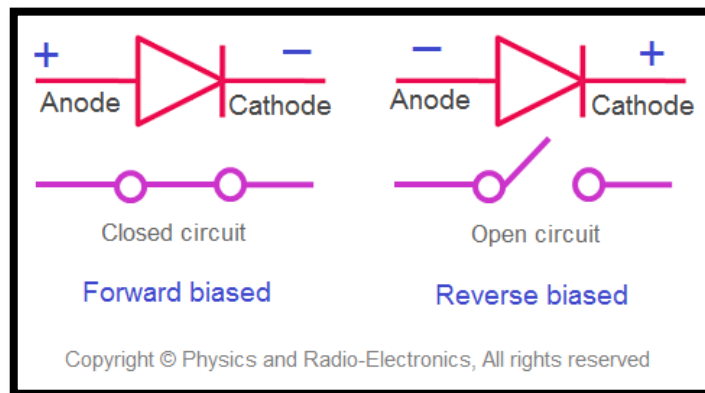
$V \geq \phi$ current increases as the voltage increases, but voltage cannot be increased indefinitely.

Remember that the current passing through the diode will heat the diode and this will break the covalent bonds and generate electron-hole pairs. If this continues, the crystal structure of the semiconductor will breakdown and this will destroy the diode. This means that there is a maximum forward voltage that can be applied across the diode, exceeding this value will destroy the diode.

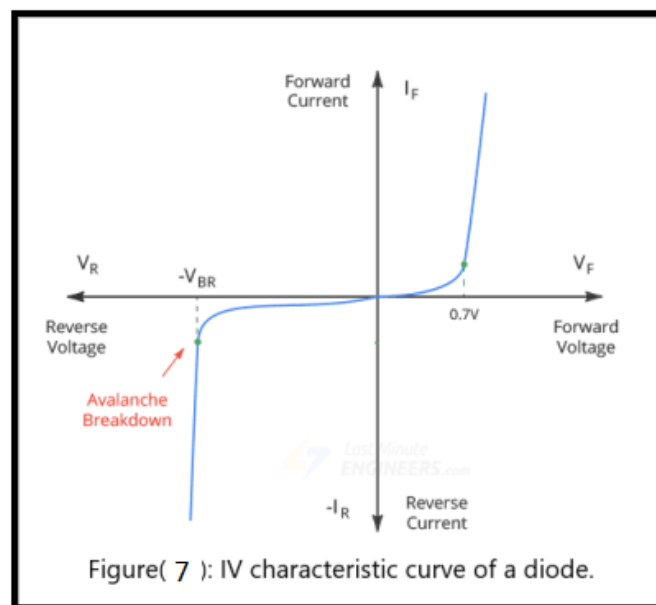
5. I-V (Current-Volt) Characteristics

The diode has:

- ❖ When forward biased: very low resistance, no depletion region, current passes through the diode.
- ❖ When reverse biased: very high resistance, broad depletion region, no current.



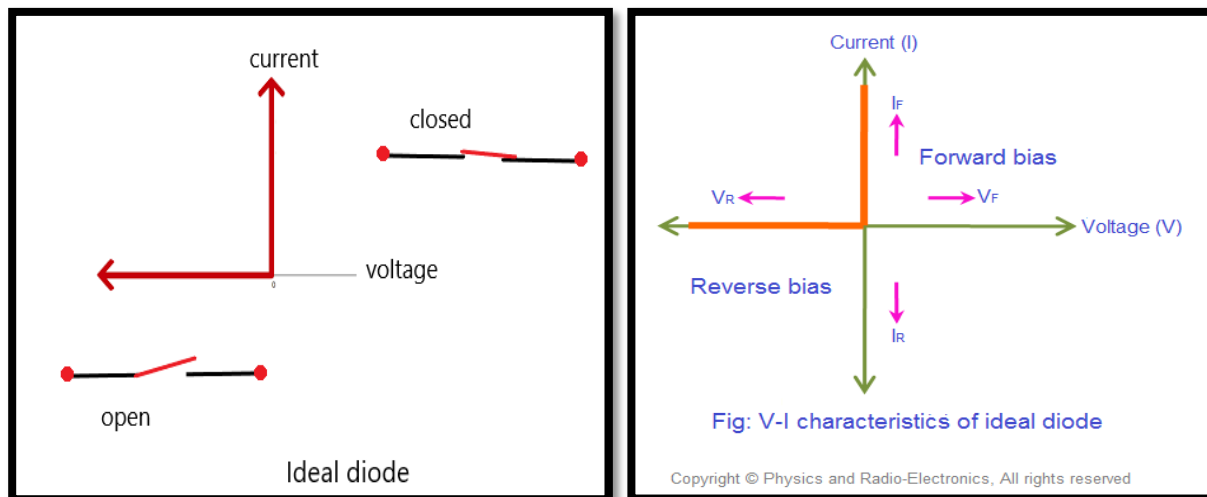
As it can be seen from Figure (7) the relation between current and voltage is not linear which means that the diode does not obey Ohms law (where the relation between voltage and current is linear). So the diode is a non-linear element. The resistance is not constant during the operation of the PN junction.



A diode conducts current in one direction only and that is when it is forward biased (its resistance is very small), but there is no current when it is reversed biased (it has very high resistance).

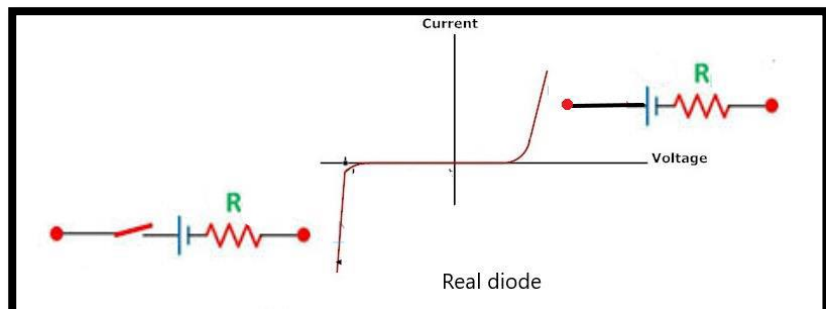
A- Ideal Diode

The ideal diode: we can make this easier by ignoring the voltage drop and the resistance. So the ideal diode has no voltage drop across it, it has zero resistance when it is forward biased and ∞ resistance when it is reversed biased. Now, we can consider the diode as a switch which is closed when it is forward biased and so current can pass through it and it is a closed switch when the diode is reversed biased (its resistance is ∞).

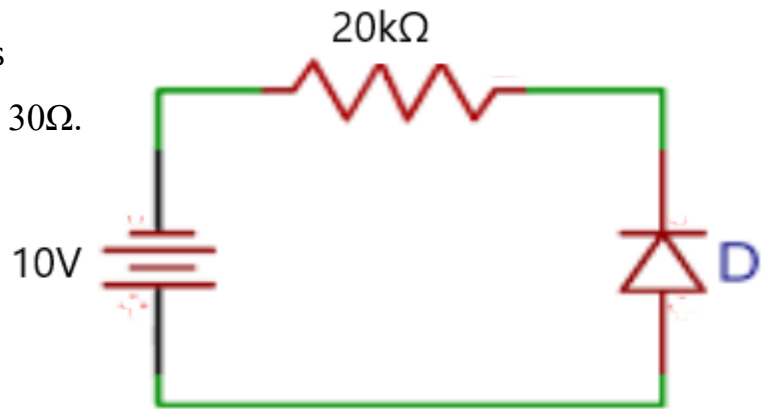


B- Real Diode

It can be represented by a battery (signifying the potential barrier) and a resistance. When it is forward biased the voltage drop across it is either 0.7V for Si or 0.3V for Ge, and the resistance is very small. The resistance is very large when the diode is reversed biased.



Example-1: Find the current in the given circuit if the barrier potential is 0.7V and the resistance of the diode is 30Ω .



Sol.:

First we must check if the diode is forward biased or reverse biased.

Since the anode is connected to the (+) end of the battery, and the cathode to the (-) end of the battery, hence the diode is forward biased.

$$V = IR$$

$$10 - 0.7 = I * (20000 + 30)$$

$I = 0.464 \text{ mA}$ This is the case of a **real diode**

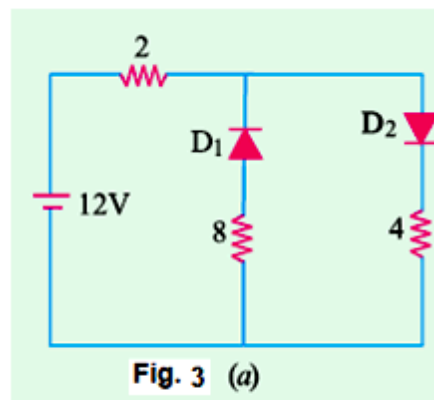
If we assume that the diode is **ideal**. Then $\phi = 0$ and R of the diode = 0 .

$$V = IR$$

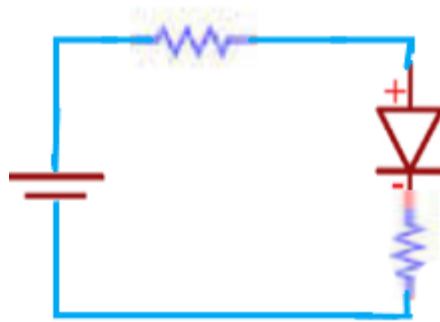
$$10 = I * 20\text{K}\Omega$$

$$I = 0.5\text{mA}$$

Example-2: Show the biasing on each diode in the given circuits.



The cathode of D1 is connected to the (+) end of the battery and its anode is connected to the (-) end of the battery. So the diode D1 is reversed biased (it is an open switch). The cathode of D2 is connected to the (-) end of the battery and its anode is connected to the (+) end of the battery. So the diode D2 is forward biased (it is a closed switch). The circuit becomes as shown.



6- Zener diode

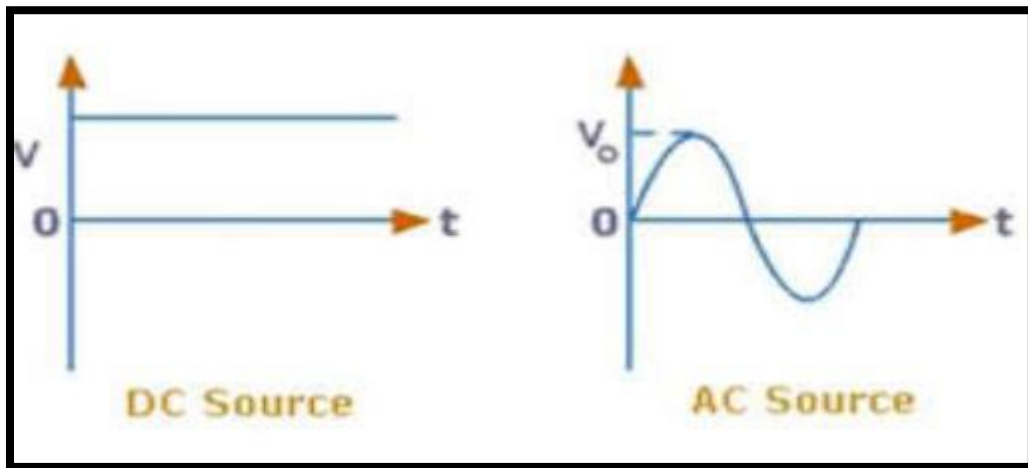
A Zener diode consists of a special, heavily doped PN-junction, that is designed to conduct in the **reverse direction** when a certain specified voltage (Zener breakdown voltage) is reached, without getting damaged.

Zener diode permits current to flow in **either a forward or reverse direction**. It operates just like the normal diode when in the forward-bias mode, and has a turn-on voltage of between 0.3 and 0.7 V. However, when connected in the reverse mode, which is usual in most of its applications, a small leakage current may flow. As the reverse voltage increases to the predetermined breakdown voltage (V_z), a current starts flowing through the diode. The current increases to a maximum, which is determined by the series resistor, after which it stabilizes and remains constant over a wide range of applied voltage. Additionally, the voltage drop across the diode remains constant over a wide range of voltages, a feature that makes Zener diodes suitable for use **in voltage regulation**. This schematic symbol for a Zener diode is shown above

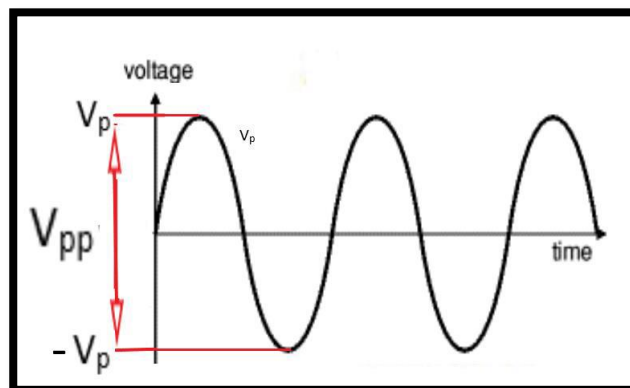


The Zener diode conducts in the reverse direction, if the voltage applied is reversed and larger than the Zener breakdown voltage.

AC and DC voltage



The DC voltage have a constant magnitude and flows in one direction i.e. have the same polarity either negative or positive. While the AC voltage have variable magnitude and direction. As seen in the figure above. The dc voltage is described fully by stating its magnitude and polarity. The AC voltage is described fully by stating it magnitude (V_p , V_{pp} , or V_{rms}) frequency and wave shape (sine, sawtooth, rectangular, square).



Lecture 3

Special Diodes

There are many types of diodes include:

- 1- Light Emitted Diode (LED).
- 2- Schottky Diode.
- 3- The Varactor.
- 4- The tunnel diode.
- 5- The photodiode.
- 6- The thyristor.

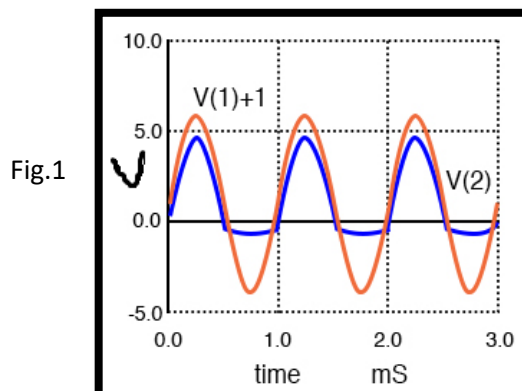
1- Applications of Diodes

The important applications of diodes are

- 1- Clipper Circuit
- 2- Clamping Circuit.
- 3- Voltage Multipliers circuit.
- 4- Logic Gates.
- 5- Rectifier Circuit.

2-1 Clipper Circuit

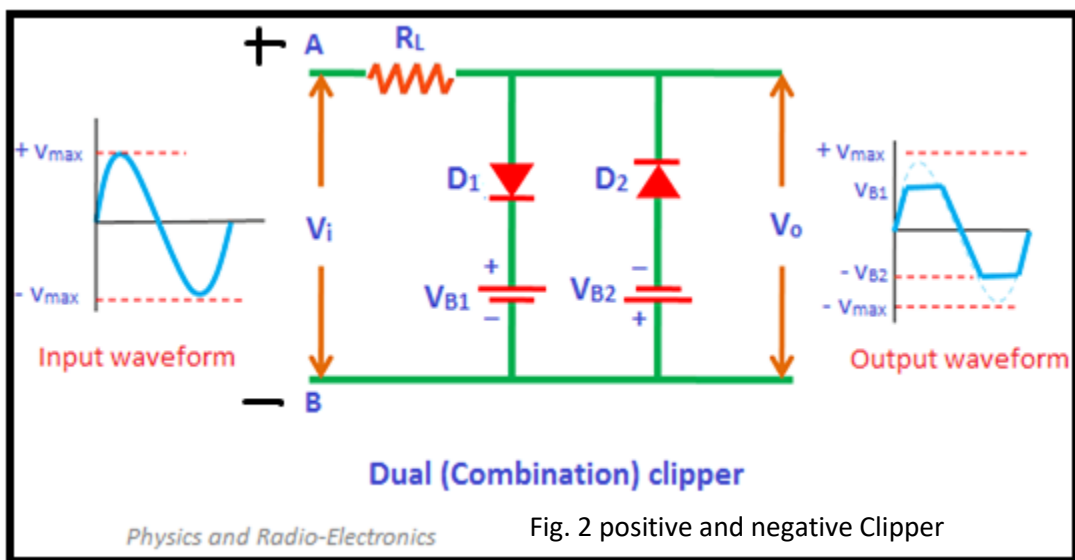
Sometimes, in radar, computers, and other applications it needs to CLIP or to CUT part of a single either positive or negative part. Here the diode act as a switch in order to clip the ac input voltage, Figure (1).



Action of the circuit

During the **positive half cycle** of the input voltage, the diode is forward biased, considering ideal diode, it will act as a closed switch and so the output voltage is zero.

During the **negative half cycle**, the diode is reversed biased and so it will act as an open switch so the whole input negative half cycles will appear at the output across. Illustrated in Fig.2 .

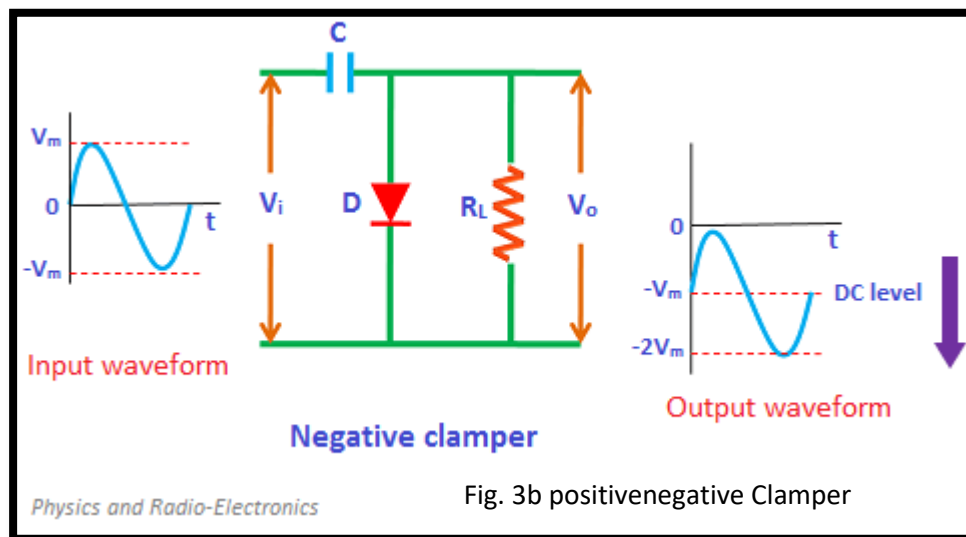
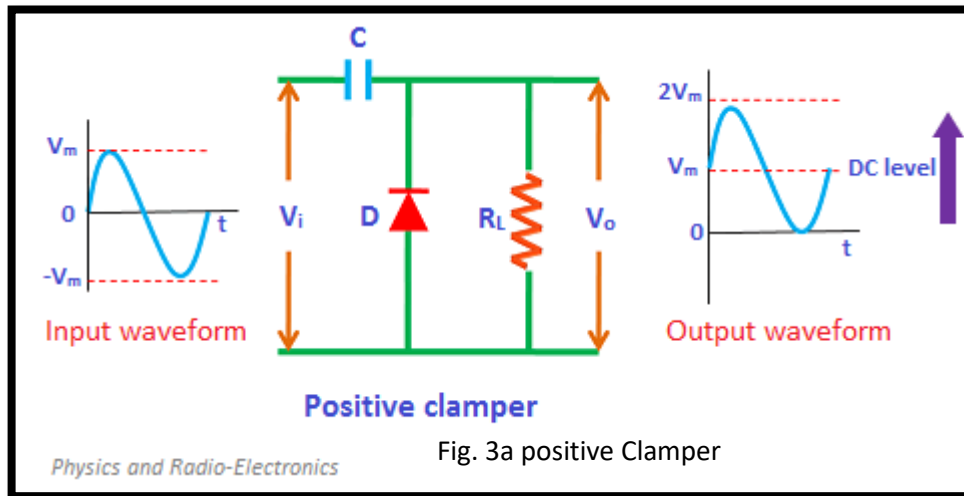


2-2 Clamper Circuit

A clamper is an electronic circuit that changes the DC level of a signal to the desired level without changing the shape of the applied signal. In other words, the clamper circuit moves the whole signal up or down to set either the positive peak or negative peak of the signal at the desired level.

A **positive clamper** circuit (Fig.3a) adds positive dc component to the input signal to push it to the positive side. When the signal is pushed upwards, the negative peak of the signal meets the zero level.

Similarly, a **negative clamper** circuit adds negative dc component to the input signal to push it to the negative side. When the signal is pushed downwards, the positive peak of the signal meets the zero level.



A capacitor is used to provide a dc offset (dc level) from the stored charge.

Action of the circuit: Positive Clamper

The positive clamper is made up of a voltage source V_i , capacitor C , diode D , and load resistor R_L . In the below circuit diagram, the diode is connected in parallel with the output load. So the positive clamper passes the input signal to the output load when the diode is reverse biased and blocks the input signal when the diode is forward biased.

During the **negative half cycle** of the input AC signal, the diode is forward biased (it acts as a closed switch) the diode allows electric current through it, no current passes through R_L and hence no signal appears at the output (see the adjacent figure). This current will flow to the capacitor and charges it to the peak value of input voltage V_m . As input current or voltage decreases after attaining its maximum value $-V_m$, the capacitor holds the charge as long as the diode remains forward biased.

During the **positive half cycle** of the input AC signal, the diode is reverse biased (acts as an open switch) and hence the signal appears at the output. The input current flows towards the output (Fig.3a).

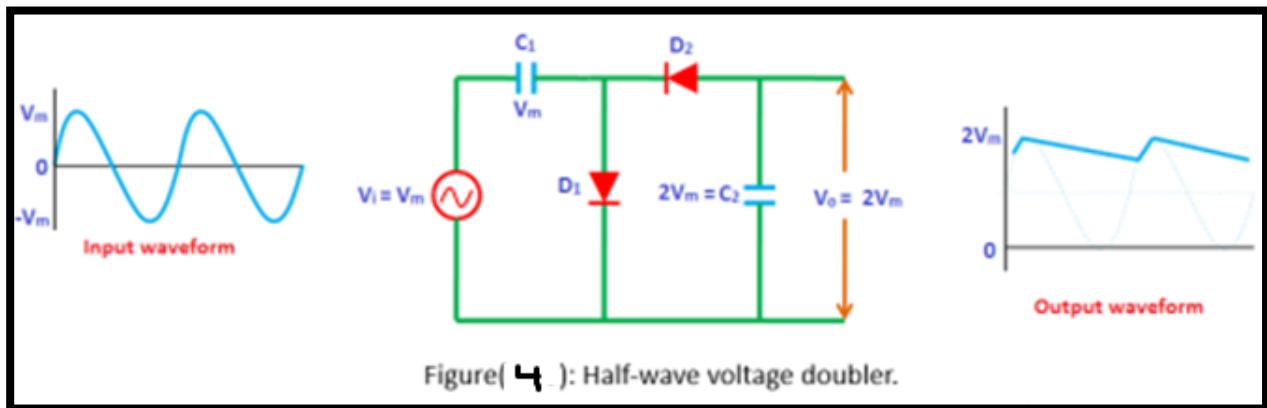
When the positive half cycle begins, the diode is in the non-conducting state and the charge stored in the capacitor is discharged (released). Therefore, the voltage appeared at the output is equal to the sum of the voltage stored in the capacitor (V_m) and the input voltage (V_m) { I.e. $V_o = V_m + V_m = 2V_m$ } which have the same polarity with each other. As a result, the signal shifted upwards.

2- Voltage Multipliers Circuit

It is a special type of diode rectifier circuit which can potentially produce an output voltage many times greater than of the applied input voltage. Example voltage doubler, voltage tripeler, quadrupler,....etc.

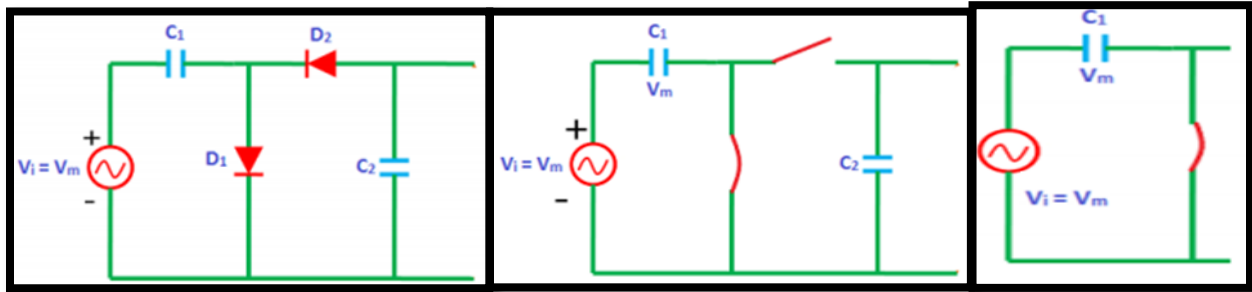
e.g. Although it is usual in electronic circuits to use a voltage transformer to increase a voltage, sometimes a suitable step-up transformer or a specially insulated transformer required for high voltage applications may not always be available. This circuit is a type of AC-to-DC converter, an AC-to-DC boost converter. The circuit intakes an AC voltage signal and gives a larger DC voltage output. By using combinations of diodes and capacitors together we can effectively multiply this input peak voltage to give a DC output equal to some odd or even multiple of the peak voltage value of the AC input voltage.

In this section, the voltage doubler circuit will take as an example. A voltage doubler circuit is a circuit in which the amplitude of the output voltage is double the amplitude of the input voltage. This circuit can be a half-wave or full-wave doubler circuit.

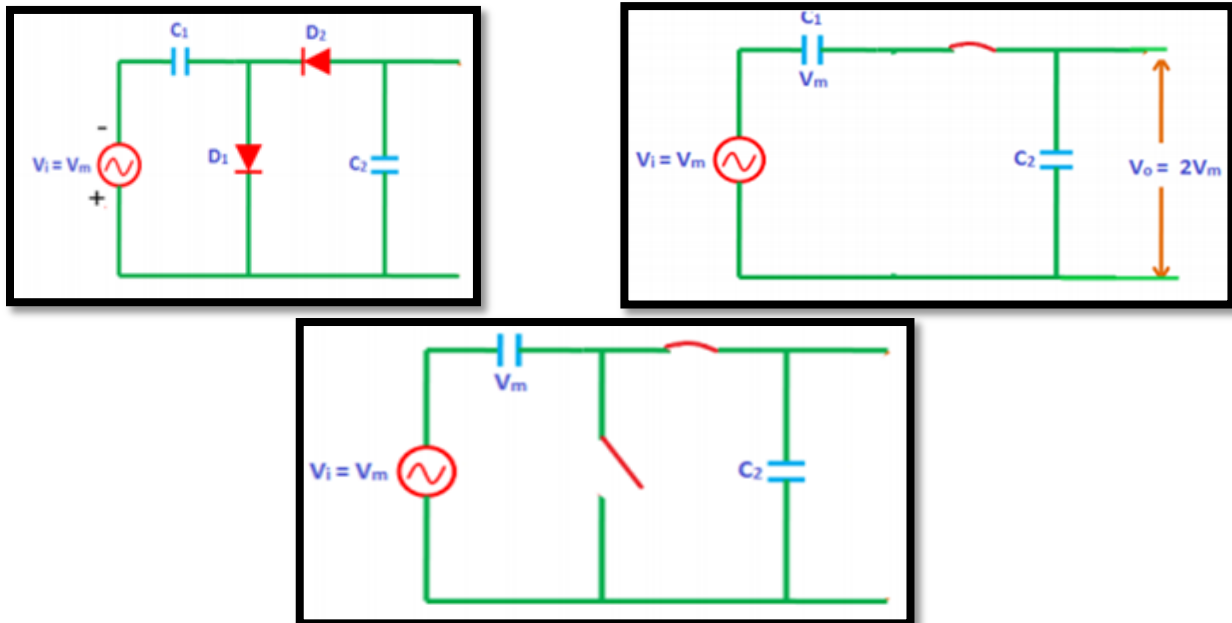


Action of the circuit:

During the positive half cycle of the AC input voltage, diode(D_1) is forward biased so it acts as a closed switch while diode(D_2) is reversed biased and so it is an open switch. The circuit can be redrawn as in the figures below.



So the capacitor C_1 will be charged to the peak value of the input voltage, V_m . During the negative half cycle, (D_1) will be reverse biased (open switch), While D_2 is forward biased (closed switch). See the figures below.



The capacitor C_1 will not be charged. A charge of (V_m) is stored in this capacitor. Now, C_2 is in parallel with C_1 (with a stored charge stored of V_m) and the input supply voltage of peak value of V_m . The capacitor C_2 charges to a value $2V_m$ because the input voltage V_m and capacitor C_1 voltage V_m is added to the capacitor C_2 . Hence, during the negative half cycle, the capacitor C_2 is charged by both input supply voltage V_m and capacitor C_1 voltage V_m . Therefore, the capacitor C_2 is charged to $2V_m$. If a load is connected to the circuit at the output side, the charge ($2V_m$) stored in the capacitor C_2 is discharged and flows to the output.

The shape of the output voltage can be seen in fig. (4). It is DC rectified voltage not AC, since it is of one direction but its value is variable.

3- Logic Circuit

Diodes can perform logic operations such as AND, OR, etc. (these are what are called logic gates). In such gates the diode again acts a switch, open or closed according to the biasing applied to it. We will take the OR logic gate as an example Fig. (5).

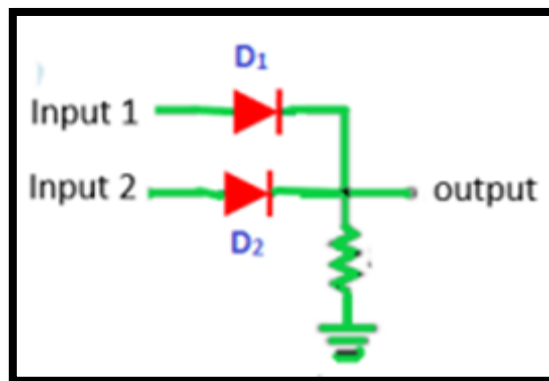
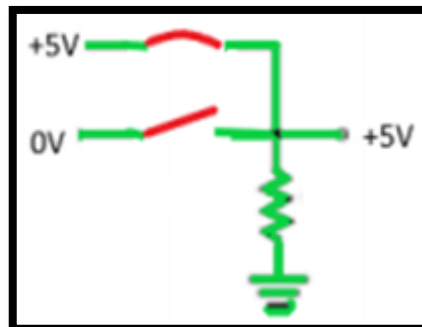
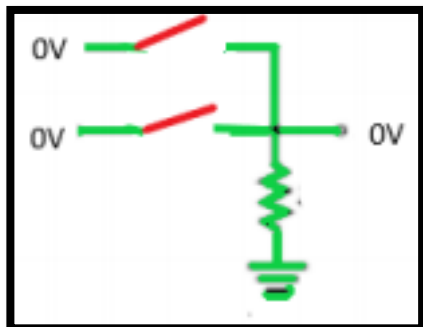


Fig.5 OR gate

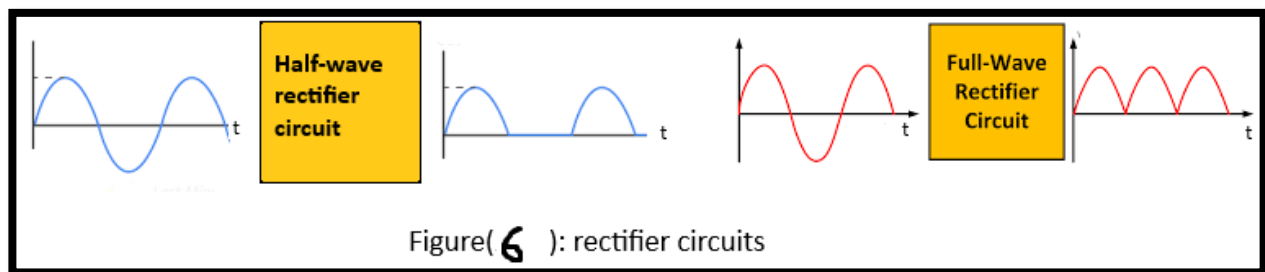
Input 1 and 2 can either be high voltage or low voltage and according to this, the biasing on each diode is determined. The cathode of both diodes is connected to earth via the resistance e.g. it is at 0 volt. if input 1 and input 2 are at low voltages, both diodes are reversed biased and so act as open switches and the output remains at 0 volt. If input 1 is at high voltage and input 2 at low voltage, then D_1 is forward biased and acts as a closed switch, while D_2 is reversed biased and act as an open switch. So the voltage at input 1 will appear at the output. These two cases are represented in the figures below.



4- Rectifier Circuits

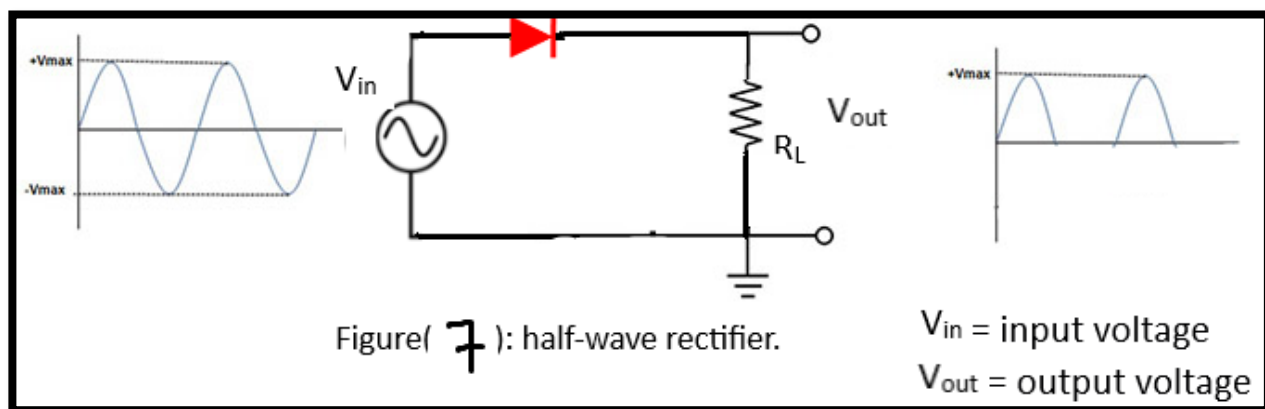
Now we come to the most popular application of the diode: *rectification*. Simply defined, rectification is the conversion of alternating current (AC) into direct current (DC). This involves a device that only allows one-way flow of electric charge. As we have seen, this is exactly what a semiconductor diode does: it allows current to flow in only one direction only. Therefore, if an alternating waveform is applied to a diode, then it will only allow conduction over half the waveform. The remaining half is blocked.

The simplest kind of rectifier circuit is the half-wave rectifier. It only allows one half of an AC waveform to pass through to the load. The other kind is the full-wave rectifier circuit, where both halves of the input voltage will have the same direction at the output (Fig. 6).

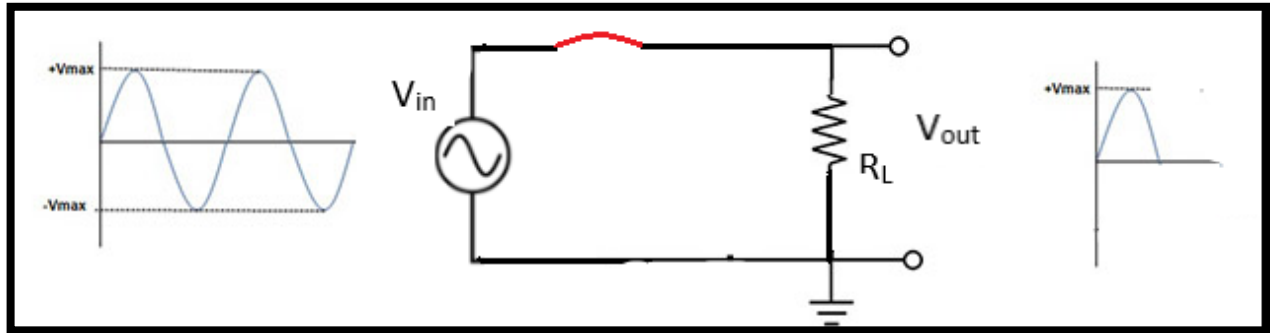


Half-wave rectifier circuit:

This circuit is shown in Fig. (7). the diode is considered as ideal diode (as explained before).

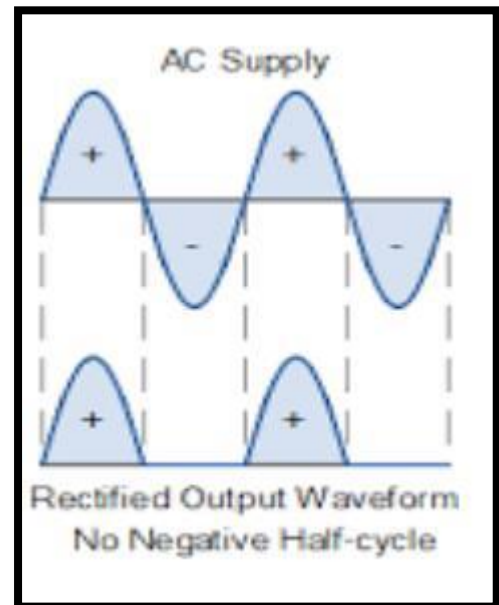


During the positive half cycle of the input voltage, the diode is forward biased, acts as a closed switch ($R_D=0, V_D=0$). So this half of the input voltage appears across the load resistance (R_L), as shown in the figure below.



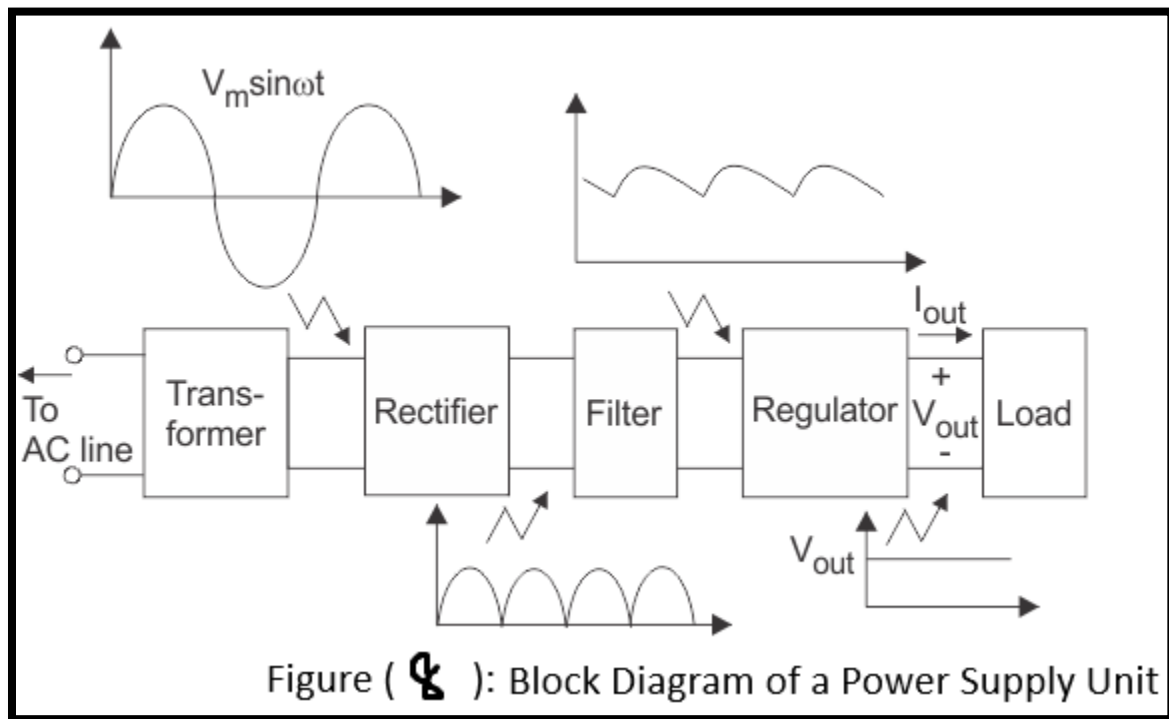
During the negative half cycle, the diode is reversed biased, it is an open switch. So no current passes through R_L and so the output voltage is zero.

Therefore, the output voltage from this circuit has only the positive halves of the input voltage. The full wave rectifier will be explained in the next section.



Power supply

A power supply unit converts the alternating voltage current (AC) into direct current (DC). Alternating voltage periodically changes its value and polarity. While the direct voltage has the same polarity and the same value. So to change AC to DC, its direction must be the same and we must cancel the variation in its value. This is done through many stages as seen in Figure (8).



Parts of the power supply:

1. Transformer. An input step-down transformer.
2. Rectifier. to convert the AC components, present in the signal to DC components.
3. Filter. to smooth the variations, present in the rectified output.
4. Regulator. to control the voltage to a desired output level.
5. Load. The load which uses the pure dc output from the regulated output.

Transformer

A step-down transformer is used to reduce the voltage from 220V of the AC line mains to a value suitable for the diodes used. When a transformer is connected in a circuit, the input supply is given to the primary coil, the secondary coil is connected to the rectifier circuit. Depending upon the number of turns in the secondary winding, a transformer can be classified either as a Step-up or a Step-down transformer. Clearly a step-down transformer is used here.

The output voltage from the secondary coil is the input voltage to the next circuit which will be referred to as V_2 . Usually this voltage is expressed in its rms value (V_{2rms}).

Rectifier circuits

There are two types of rectification, half-wave and full-wave.

The circuits are:

- Half-wave rectifier.
- Full-wave rectifier.

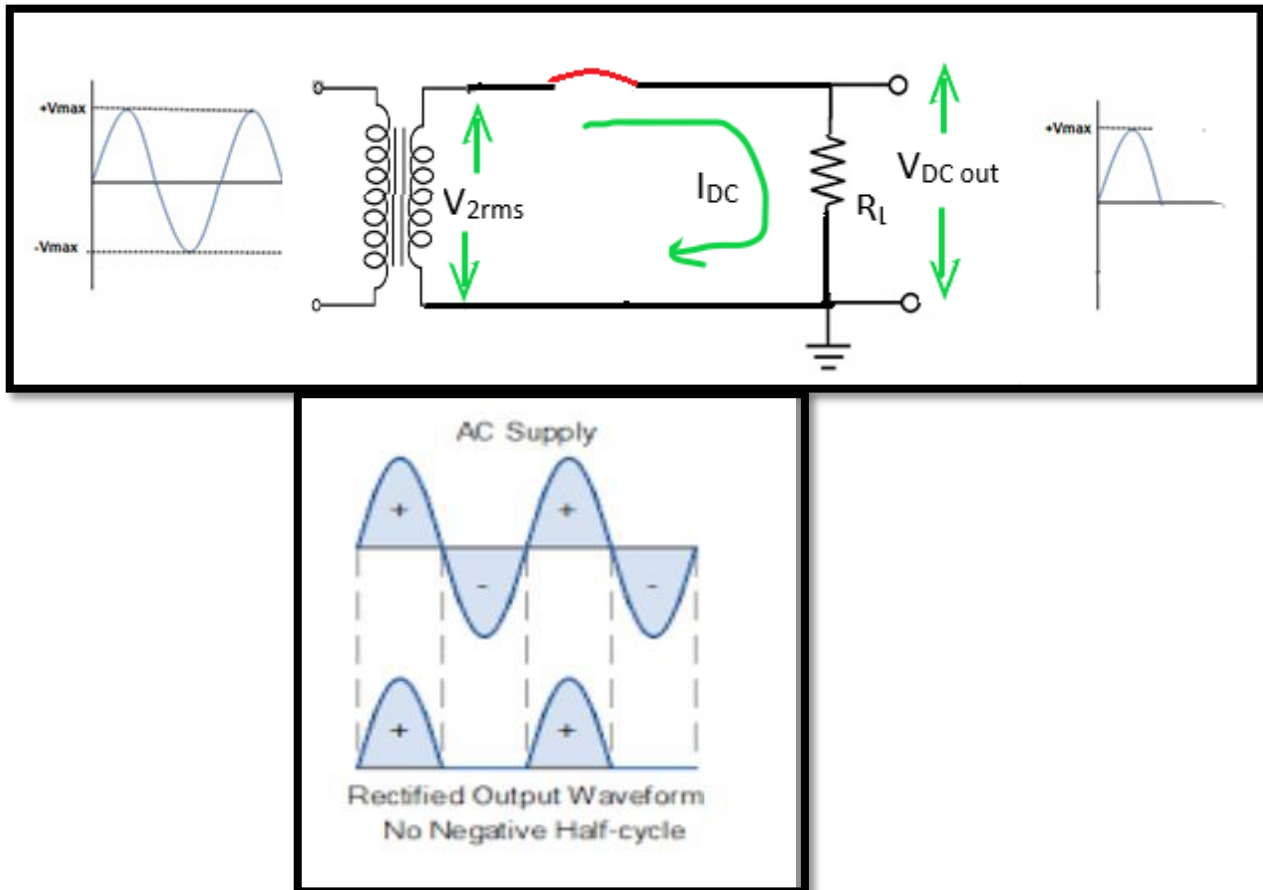
Two circuits are used here

- 1- Full-wave rectifiers using a centre tapped transformer.
- 2- Full-wave rectifiers using bridge circuit.

The output voltage from the rectification circuit is a wave of one direction (same polarity), either positive or negative, but its value is not constant. So it is AC in the sense that its value is variable but it is DC since it has the same polarity. It can be looked at as DC with AC components. For this reason, this voltage is called the pulsating DC voltage or rectified AC voltage.

Half-wave rectifier

We have already explained the action of this circuit and we saw that the output voltage is only half waves of the input voltage, either the negatives or the positives.



The dc output voltage: The AC output voltage ($V_{ac out}$) has the same peak voltage as the input voltage, which is the output voltage of the transformer from the secondary coil(V_{2P}).

$$V_p = \sqrt{2} V_{rms}$$

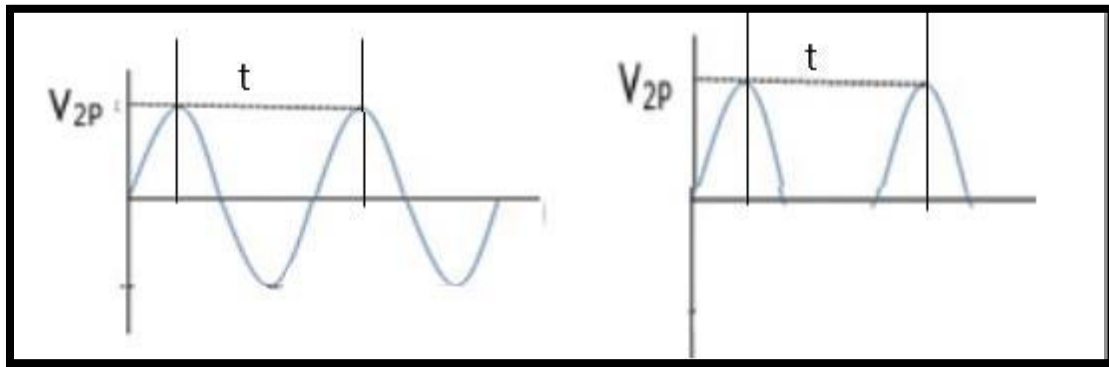
$$V_{ac out} = V_{2P}$$

$$V_{dc out} = V_{ac out} / \pi$$

$$= V_{2P} / \pi$$

$$V_{dc out} = \sqrt{2} V_{2rms} / \pi$$

We should also know the relation between the frequency of the output voltage and that of the input voltage. As seen from the figure below, t , which is the period of the wave, is the same for both voltages so the frequency which is $1/t$ is the same. So $f_{out} = f_{in}$



The DC current that passes through the diode is the same as the current passing through the load resistance (R_L).

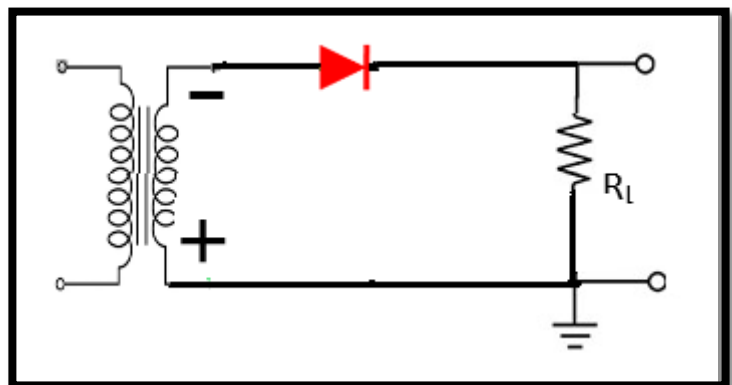
$$I_{dc} = I_{Diode}$$

$$I_{dc} = V_{dc\ out} / R_L$$

Peak inverse voltage (PIV): during the negative half cycle of the input voltage, the diode is reversed biased so it does not conduct. The voltage across it in this case is called the peak inverse voltage.

As seen in the figure below, which illustrates the state of the circuit during the negative half cycle of the AC input voltage. It is clear that the diode is reversed biased, so no current passes through the circuit, (through R_L) hence there is no voltage drop across R_L . therefore, all the output of the secondary coil appears across the diode.

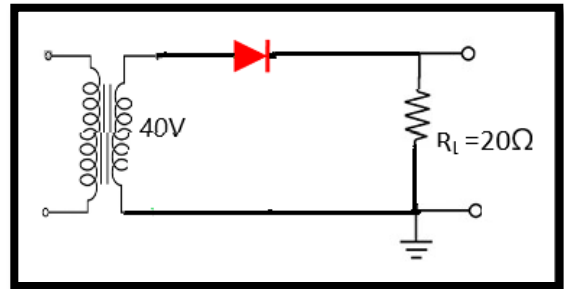
$$PIV = V_{2p}$$



Example: For a half wave rectifier, the ac output voltage from the secondary

coil of the transformer is 40V. Calculate:

1. The dc voltage across the load resistor.
2. The peak inverse voltage across the diode.
3. The dc current through the diode.



1- Usually the voltage stated for the output of the transformer is the rms value, unless otherwise stated.

$$V_{dc \text{ out}} = V_{ac \text{ out}} / \pi$$

$$= V_{2P} / \pi$$

$$V_{dc \text{ out}} = \sqrt{2} V_{2rms} / \pi$$

$$= \sqrt{2} * 40 / \pi$$

$$V_{dc \text{ out}} = 18V$$

$$2- PIV = V_{2P} = \sqrt{2} V_{2rms}$$

$$= \sqrt{2} * 40$$

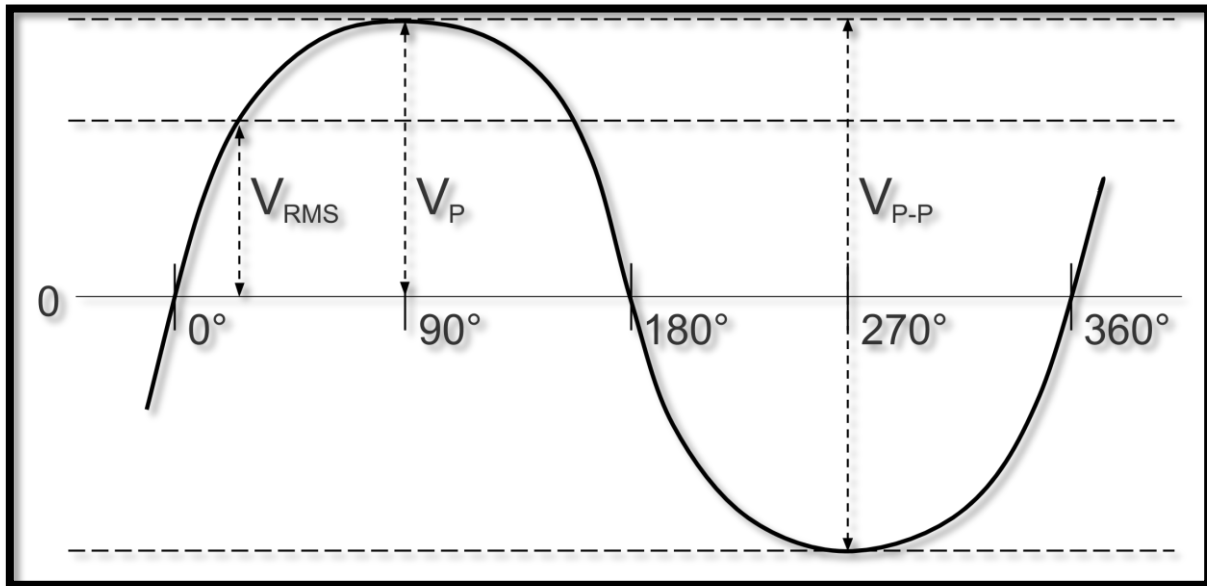
$$PIV = 56.6V$$

$$3- I_{dc} = I_{Diode}$$

$$I_{dc} = V_{dc \text{ out}} / R_L$$

$$= 18 / 20 = 0.9A \text{ which is the DC current passing through the diode.}$$

Lecture 4



Peak to peak value: it is the sum of the positive peak and negative peak values usually written as pp value.

Peak value and maximum value: it is the highest value reached by the current in one cycle.

Root Mean Square (RMS) value: also called the effective value. It is the value of the current of $\theta=45^\circ$ which equals to $0.707 I_m$.

Peak Inverse Voltage (PIV): the maximum reverse voltage across the diode during the cycle.

2- Full-Wave Rectifier using centre-tapped (CT) transformer

What is centre-tapped transformer?

Secondary coil of the transformer is divided, at the “centre tap”, into two halves of the same number of windings, this means that the voltage from each half is equal to the voltage from the second half, but is 180° out of phase with each other. At the centre tap the voltage is 0. This allows the transformer to provide two separate output voltages which are equal in magnitude, but opposite in polarity to each other. In this case both half-cycles of the input are utilized with the help of two diodes working alternatively. Figure below.

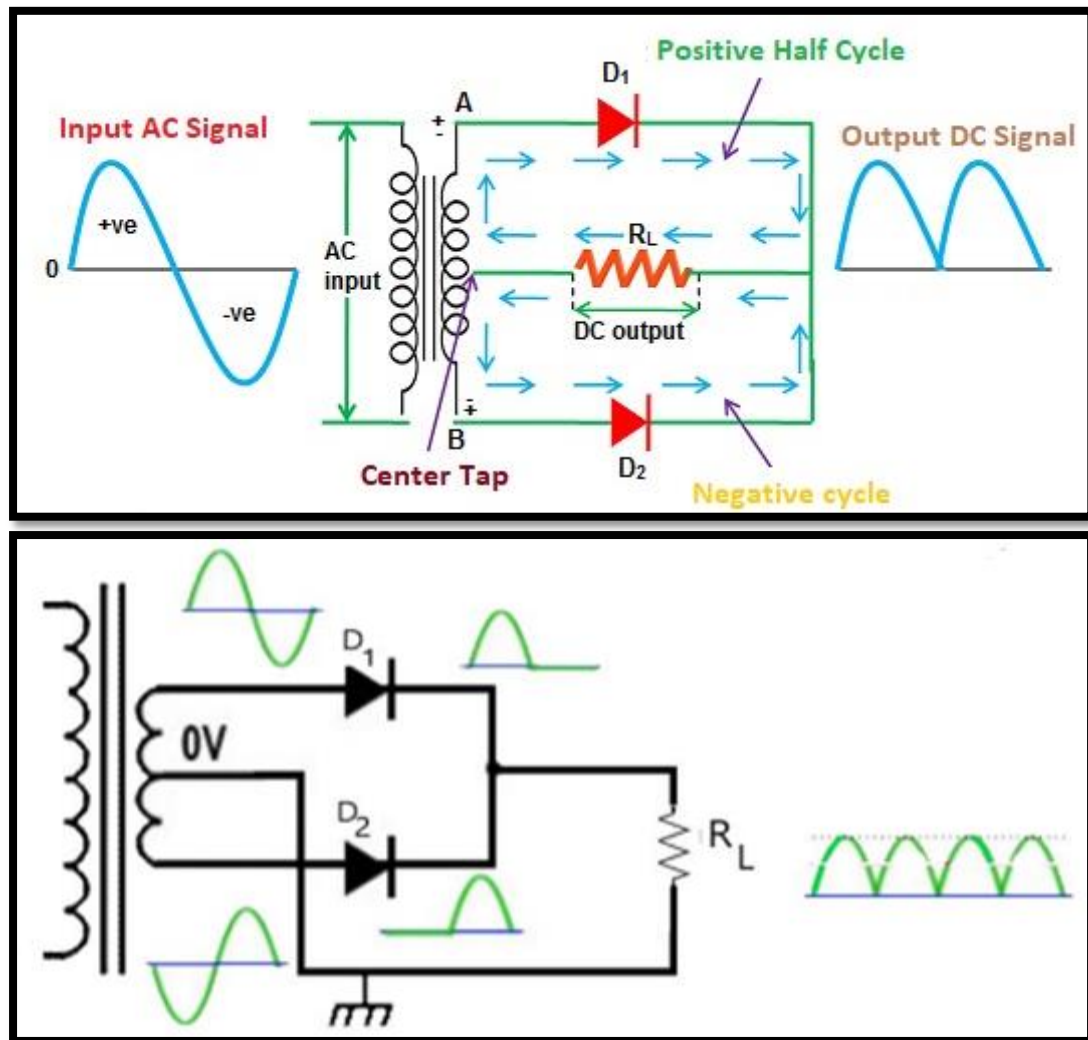
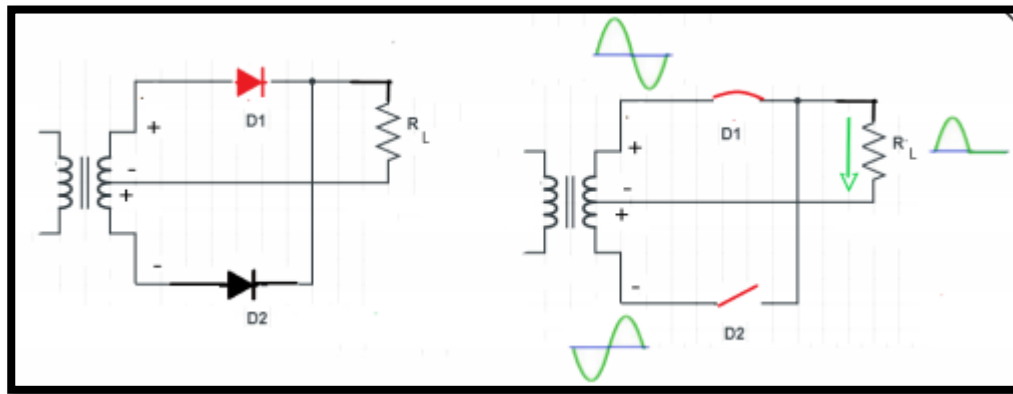


Figure shows the full wave rectifier

Operation

At the positive half cycle of the secondary voltage, D_1 (upper diode) is forward-biased and the lower diode is reversed biased. Therefore, D_1 conducts and current flows through the load resistor.

During the negative half cycle, current flows through the load resistor, D_2 conducts. So we find that current keeps on flowing through R_L in the same direction in both half-cycles of the AC inputs.



Equations

$$V_{2p} = V_{rms}/1.414$$

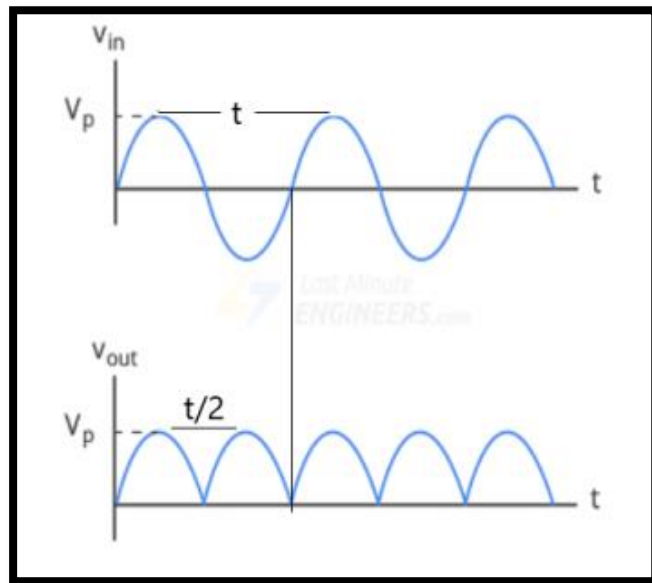
$$V_{dc} = 0.636 V_{out(rms)} = 2/\pi V_{out(rms)}$$

$$V_{out} = V_{2p}/2$$

$$F_{out} = 2 F_{in}$$

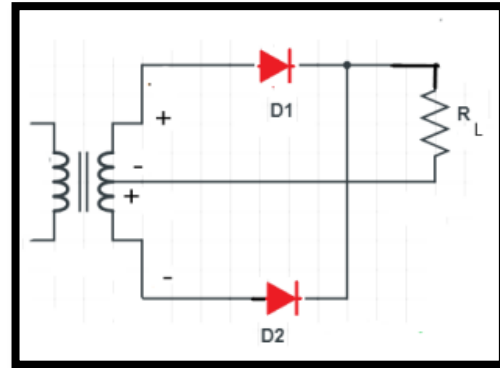
$$PIV = V_{2p}$$

$$I_{diode} = I_{dc}/2$$



Example: For a full wave rectifier using a CT transformer, the ac output voltage from the secondary coil of the transformer is 40V. Calculate:

1. The dc voltage across the load resistor.
2. The peak inverse voltage across the diode.
3. The dc current through the diode.



$$1- V_{dc\ out} = V_{ac\ out} / \pi = V_{2p} / \pi$$

$$V_{dc\ out} = \sqrt{2} V_{2rms} / \pi$$

$$= \sqrt{2} \times 40 / \pi$$

$$V_{dc\ out} = 18V$$

$$2- PIV = V_{2p} = \sqrt{2} V_{2rms}$$

$$= \sqrt{2} \times 40$$

$$PIV = 56.6V$$

$$3- I_{dc} = V_{dc\ out} / R_L$$

$$= 18 / 20 = 0.9A$$

$$I_{dc} = 2 I_{diode}$$

$$I_{diode} = I_{dc} / 2$$

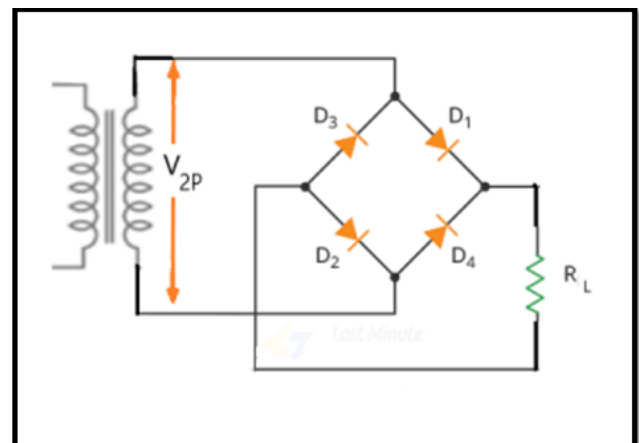
$$I_{diode} = 0.9 / 2$$

$$= 0.45 A$$

3- The Bridge Rectifier

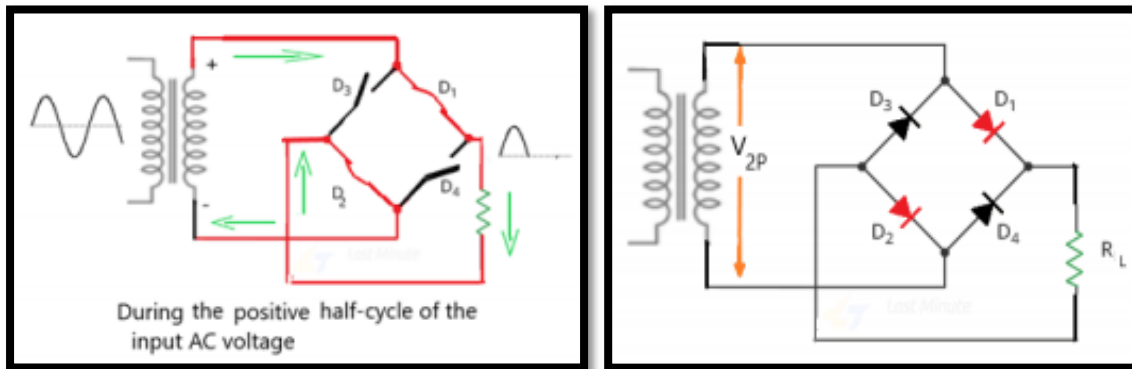
Bridge rectifier is the most widely used of all rectifier circuits.

Four diodes are connected in a bridge configuration. In contrast to the previous rectifier where two diodes were used.

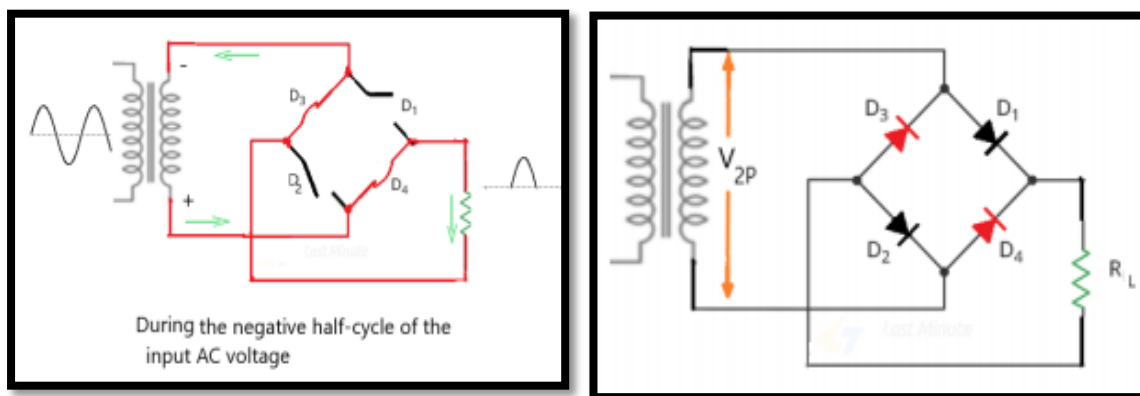


Operation

During the positive half cycle of secondary voltage, diodes D_1 and D_2 are forward biased (ON) as shown in figure below, while diodes D_3 and D_4 are reversed biased (OFF). Because of this electrons flow up through D_2 to the right the load resistance, and up through D_1 (conventional flow in the opposite way).



During the negative half cycle, diodes D_3 and D_4 are conducting, while diodes D_1 and D_2 are (OFF), this time electrons flow down through D_3 to the right through the load resistance, and down through D_4 .



During either half cycle, the electron flow is to the right through the load resistor, because of this the pulse minus polarity of load voltage is the same. This is why the load voltage is the full-wave signal.

Equations

$$V_{dc\ out} = 2V_{ac\ out} / \pi V_{out} = V_{2P}/2$$

$$V_{dc} = 2 V_{2P} / \pi$$

$$V_{dc\ out} = 2\sqrt{2} V_{2rms} / \pi$$

$$F_{out} = 2 F_{in}$$

$$PIV = V_{2P}$$

$$I_{diode} = I_{dc}/2$$

Example: For a full wave rectifier using a bridge circuit, the ac output voltage from the secondary coil of the transformer is 40V of 50Hz frequency. Calculate:

1. The dc voltage across the load resistor.
2. The peak inverse voltage across each diode.
3. The dc current through each diode.
4. The frequency of the output voltage>

1- $V_{2P} = 56.6V$

$$V_{out(peak)} = V_{2P} = 56.6V$$

$$V_{dc} = 2 V_{2P} / \pi = 36V$$

2- $PIV = V_{2P} = 56.6V$

3- $I_{dc} = V_{dc} / R_L = 36V / 20\Omega = 1.8A$

$$I_{diode} = 1.8 / 2 = 0.9A$$

4- $f_{out} = 2f_{in}$

$$f_{out} = 2 \times 50 = 100Hz$$

Comparing the three circuits through the results of the three examples

	$V_{dc \text{ out}}$	I_{dc}	I_{diode}	PIV	f_{out}
Half-wave rectifier	$= V_{2P}=18V$	0.9A	0.9A	56.6V	$= f_{in}$
Full-wave rectifier using CT transformer	$=V_{2P}=18V$	0.45A	0.45A	56.6V	$= 2f_{in}$
Full-wave rectifier using bridge circuit	$=2V_{2P}=36V$	0.9A	0.9A	56.6V	$= 2f_{in}$

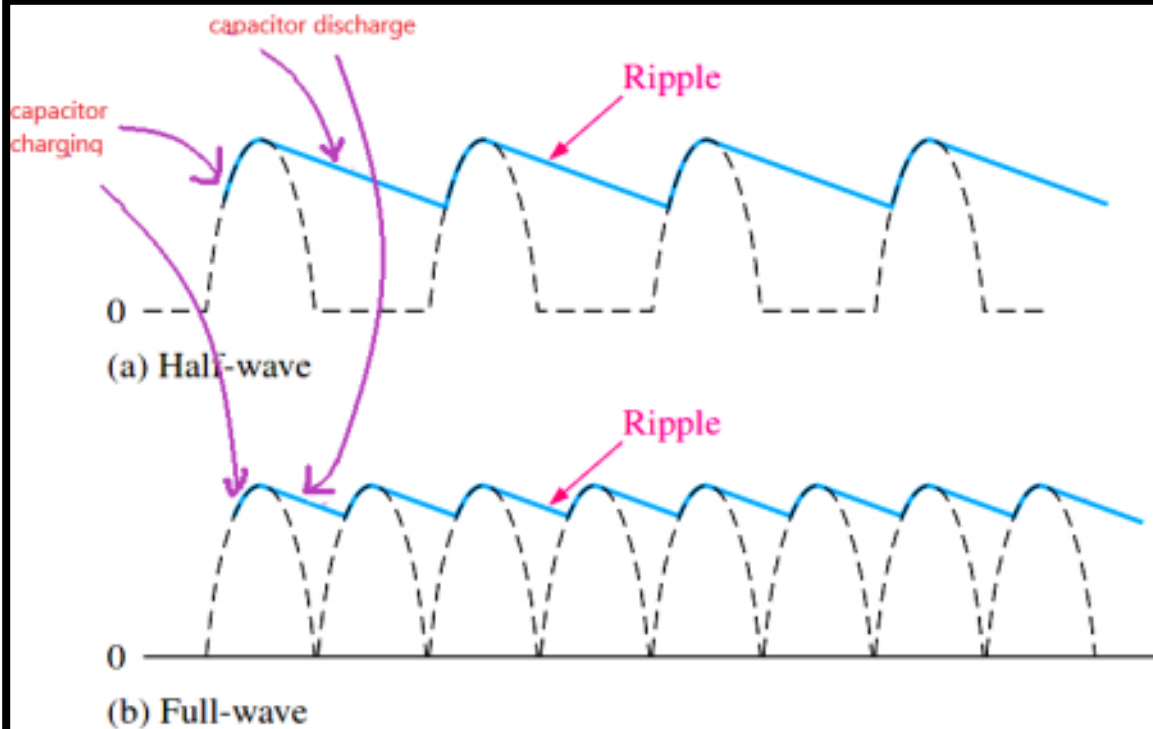
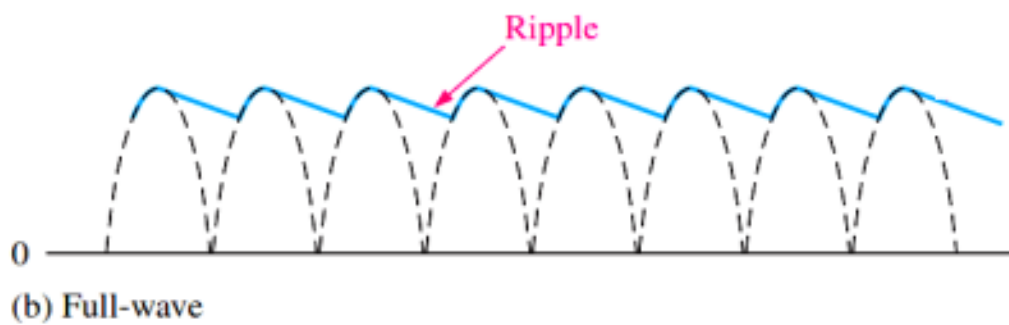
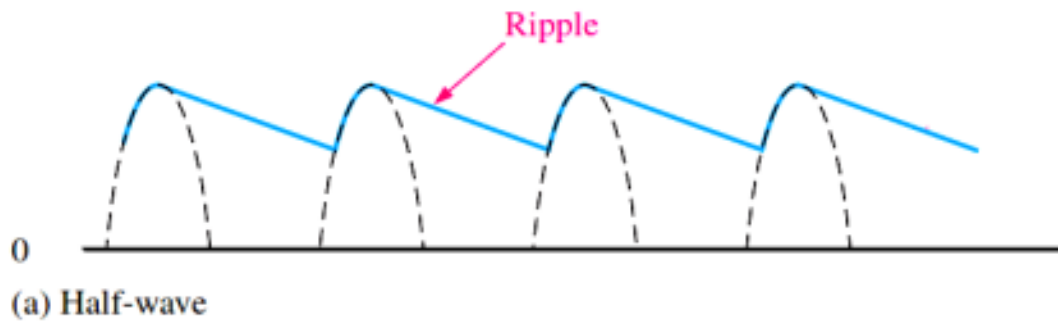
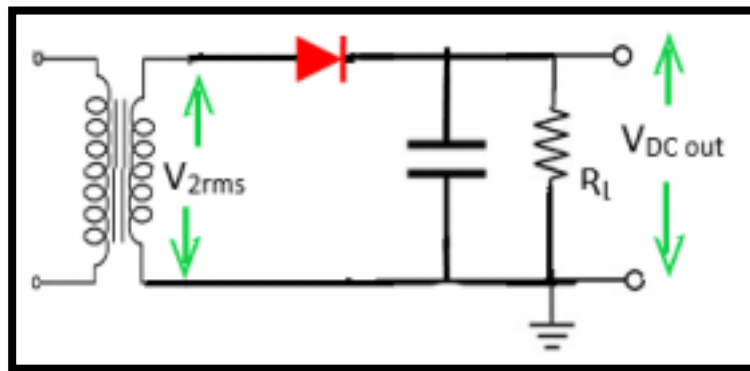
From the above table it can be seen that:

1. Using the same transformer in the three circuits, the peak inverse voltages are the same for the three circuits.
2. The full-wave rectifier using bridge circuit produces the largest dc voltage. The only disadvantage of this circuit is its use of a bridge circuit (the use of four diodes).

It can be concluded that the best rectifier is the full-wave using bridge circuit. In fact, this circuit is the most widely used circuit for rectification.

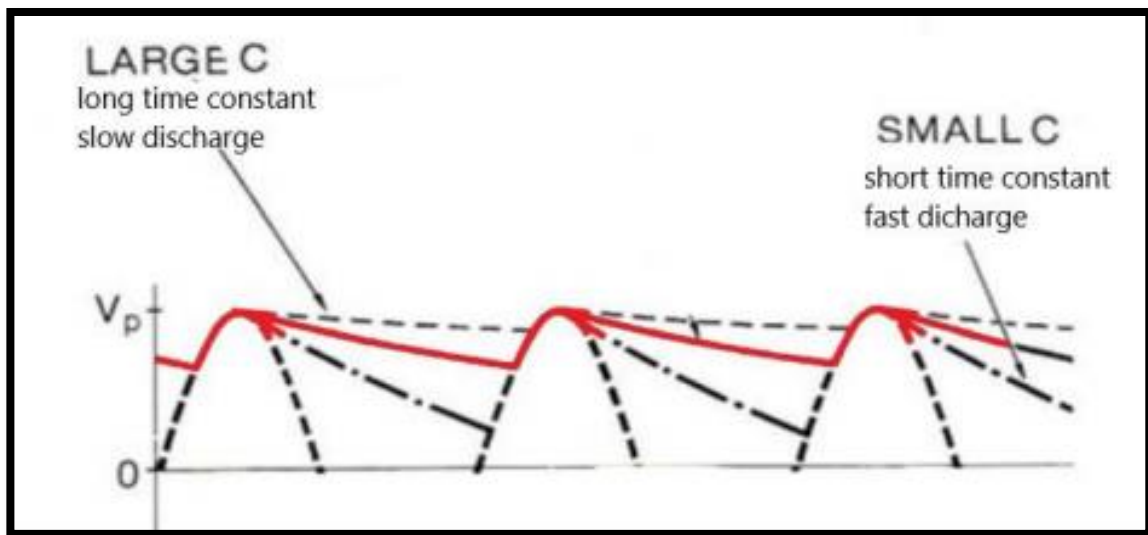
Filter

The main function of a filter circuit is to minimize the ripple content voltage (AC component) in the rectifier output. It is a circuit that converts a pulsing output from a rectifier into a very steady DC level. So it helps to remove the ripple. Many filter types are excited but we will focus on the capacitor filter.

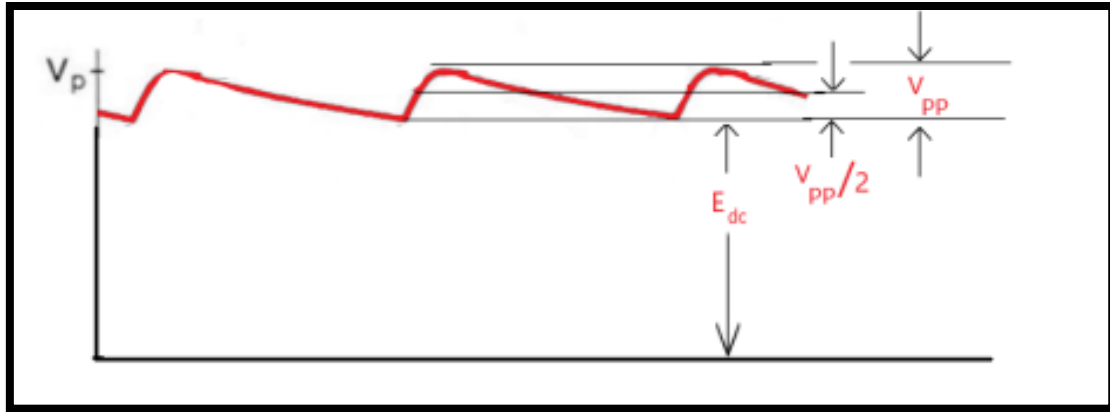


The figures above show the action of smoothing the output voltage for (a) half wave rectified voltage and (b) full wave rectified output voltage.

As the voltage increases, the capacitor charges, as the voltage decreases, the capacitor discharges, it continues discharging until the voltage increases again, the capacitor starts to charge again. From the above figure, it is clear that the ripple voltage (V_{pp}) is less for full wave rectification.



The rate of discharge depends on the time constant RC . Long time constant (large value of C) means slow discharge; V_{pp} is less i.e. the ripple voltage is less. Small C will give low time constant which means fast discharge, V_{pp} is high. This means that a high value of C (high RC time constant) leads to lower ripple of the output voltage. This means that the change in the value of the DC output voltage has decreased (the ripple voltage has decreased). The output voltage is almost a DC voltage equal to the peak value of the input voltage.



$$E_{dcout} = E_{dc} + V_{pp}/2$$

Ripple factor(r): This is a measurement of ripple. It gives an indication of how efficient the filter is in removing the AC components present in the output voltage. Ripple Factor is defined as the ratio of the rms value of ac component present in the rectified output to the average value of the rectified output. Its value should be less than one.

$$\text{Ripple Factor (r)} = \frac{\text{RMS value of AC component present in Rectifier Output}}{\text{Average Value of Rectifier Output}}$$

A high ripple factor means that there is high ripple in the output voltage (undesirable case). The output voltage should be of constant value that is the ac components (ripple voltage) should be very small; this requires a small ripple factor.

Ripple factor of a capacitor filter:

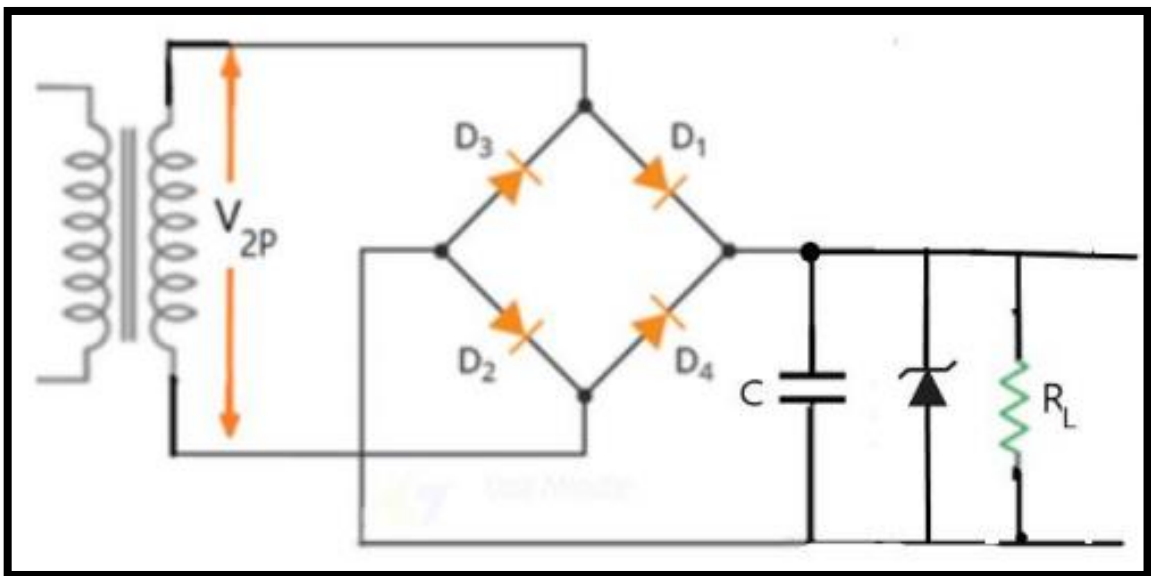
$$r = \frac{1}{2\sqrt{3}} \frac{V_{pp}}{E_{dc}} = \frac{1}{2\sqrt{3}} \frac{V_m}{E_{dc}} \frac{1}{fRC}$$

for small values of r , the time constant must be high and this is done either by a large C or a large R_L .

Voltage Regulator

The change in voltage from no load to full-load condition is called voltage regulation. The aim of the voltage regulation circuit is to reduce these variations to zero, or at least, to the minimum possible value.

One way to stabilize or regulate the DC load voltage is by using Zener diode across the load resistor as shown in figure below.



Lecture 5

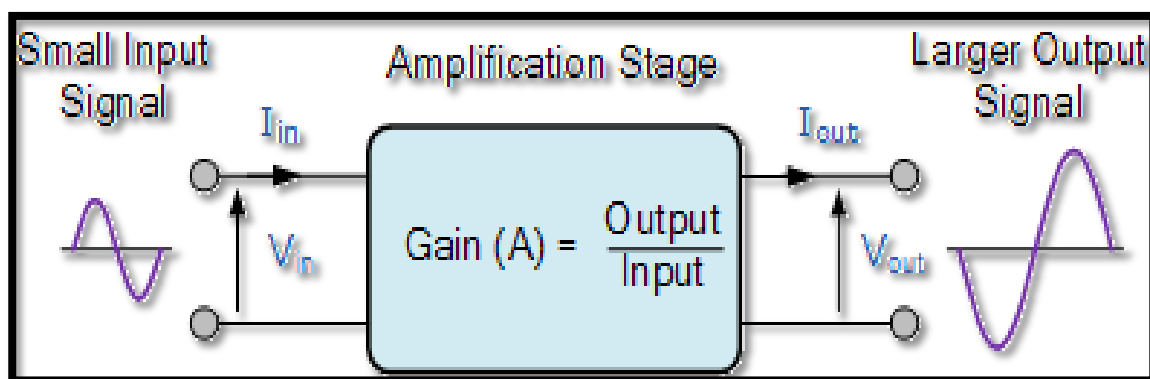
Amplification and Voltage Amplifiers

1- Introduction

An amplifier is an electronic device or circuit which is used to increase the magnitude of a signal (voltage, current or power) applied to its input. In this chapter we will concentrate on the amplification of voltage (voltage amplifier). The components that can be used for amplifications are the transistor and the field effect transistor.

So the output voltage of a voltage amplifier is so many times (amount of amplification) the input voltage. Nothing else of the signal characteristics (frequency and shape) should be changed. The output voltage should be exactly a multiple of the input voltage i.e. if the amplification is 10 times, then the output voltage should be ten times the input voltage. The amount of amplification is called the gain of the amplifier (in our case we mean the voltage gain). It is defined as the ratio of the output voltage to the input voltage.

$$A_V = \frac{V_{output}}{V_{input}}$$



[Fig.1 diagram of ideal amplifier](#)

Characteristic of an **ideal** amplifier

1. Voltage gain $=\infty$
2. Input resistance $=\infty$
3. Output resistance $=0$
4. No noise with the output voltage
5. No distortion of the output voltage
6. No change of frequency
7. Infinite bandwidth.

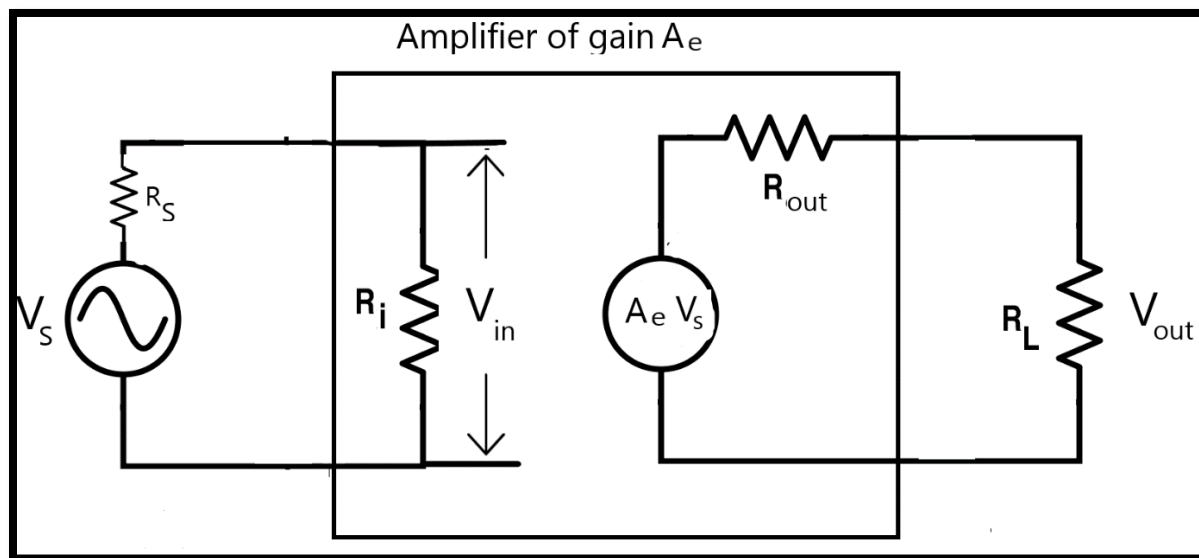


Fig.2 an equivalent circuit of an amplifier

1,2,3- R_i and R_{out} represents the input resistance and the output resistance.

V_S : A source of AC voltage.

R_S : resistance of voltage source, it connected to the input of the amplifier i.e. the output voltage from this source is to be amplified.

The output amplified voltage appears across a load whose resistance is R_L .

As seen from the figure, R_i is connected in parallel with the input voltage source and R_{out} is connected in series with R_L .

For a certain value of R_i , it draws current from the voltage source, this means that there is a voltage drop across R_i which is subtracted from V_S and so

decreasing it. ($V_S - V_{Ri}$) is V_{in} , the input voltage to the amplifier to be amplified. So R_i must be very large ($=\infty$ in the case of an ideal amplifier) so that it will not draw any current from the voltage source and all V_S will be amplified.

R_{out} , which is connected in series with R_L , should be very small ($=0$ in case of ideal amplifier) so that the voltage drop across it is very small ($=0$ in case of ideal amplifier), so a very small voltage is lost from the amplified voltage that will pass through to the load.

- 2- No noise: Noise here means any undesirable voltage. In the case of voltage amplifiers, the output voltage must be the gain of the amplifier times the input voltage, but the output voltage may be larger than this. This is due to any voltages inside the amplifier (called noise) that are not the input voltage which will be amplified and thus are added to the output voltage.
- 3- No distortion: Distortion is the change of the wave shape. For a voltage amplifier, if the input voltage is a sine wave, the output should be a sine wave as well, a replica of the input voltage.
- 4- No change of frequency: the frequency should stay the same as the frequency of the input voltage.
- 5- Infinite bandwidth this means that the gain of the amplifier must be the same for all the input frequencies i.e. the amplifier amplifies all the voltages irrespective of the frequency, the gain is constant for all the input frequencies.

1- Transistors

The transistor is a semiconductor crystal with three doped regions like the NPN transistor fig.1. The bipolar transistor structure consists of three alternating

layers of n- and p- type semiconductor material. These layers are referred to as the emitter (E), base (B), and collector (C) of the transistor.

The emitter is heavily doped; its job is to emit or inject electrons into the base.

The base is lightly doped and very thin, so as to decrease recombination to a minimum, it passes most of the emitter-injected electrons on the collector.

The collector is largest of the three regions because it must dissipate more heat than the emitter or base, the doping level of the collector lies between that of the emitter and the base.

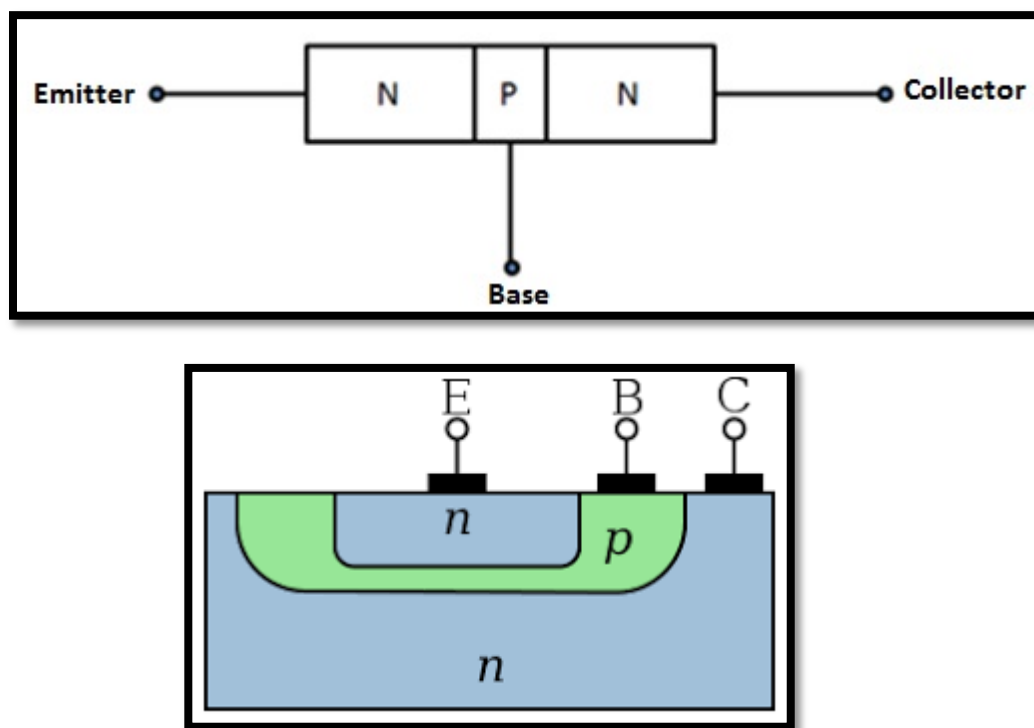


Fig.3 NPN transistor structure

The transistor in the two figures has two junctions, one between the emitter and the base, and another between the base and the collector. Because of this, a transistor is similar to two diodes connected back to back. We call the diode on the left the emitter-base diode, or simply emitter diode. The diode on the right is the collector-base diode or the collector diode. Fig. 4.

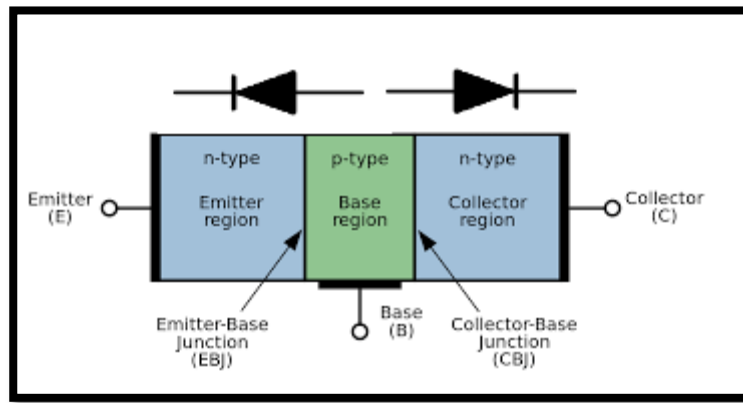


Fig. 4 NPN transistor

2- Schematic Symbol

An arrow head points in the easy direction of conventional flow. The electrons flow against the arrow head. Remember that the arrow head points to where the electrons are coming from. Fig.5

If the arrow is going out of the base, the transistor is an NPN.

If the arrow is pointing to the base, the transistor is a PNP transistor.

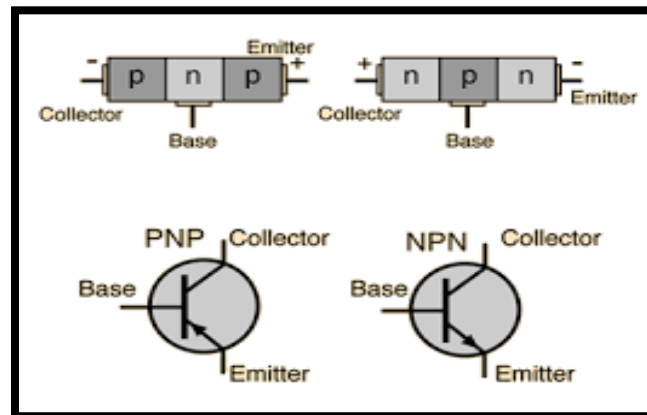


Fig. 5 Transistor Symbol

3- Biasing of Transistor

In order to let the current pass through the transistor, I will need two biasing. The forward bias is applied to the emitter diode; free electrons in the emitter have not yet entered to the base region. If the applied voltage exceeds approximately

0.7V, many free electrons will enter the base. The free electrons in the base can now flow in either of two directions, down the thin base into the external base lead or across the collector junction into the collector region.

The thin and lightly doped base gives almost all of these free electrons enough time to diffuse in to the collector before they can recombine with holes in the base. The free electrons then flow out of the collector in to the positive terminal of the collector supply.

Remember

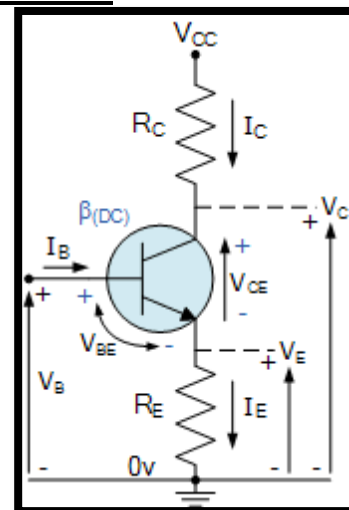
- 1- For normal operation, forward-bias the emitter diode. $V_{EE} > V_{BE}$. So EB has low resistance.
- 2- Reversed-bias the collector diode, in order to pass I_E . That will let I_C to flow through the circuit. The potential barrier of this junction is V_{CB} . So CB has high resistance.
- 3- Collector current is almost equal to emitter current $I_C = I_E$.
- 4- Base current is very small.

$$I_B = I_E - I_C$$

4- Transistors Circuits

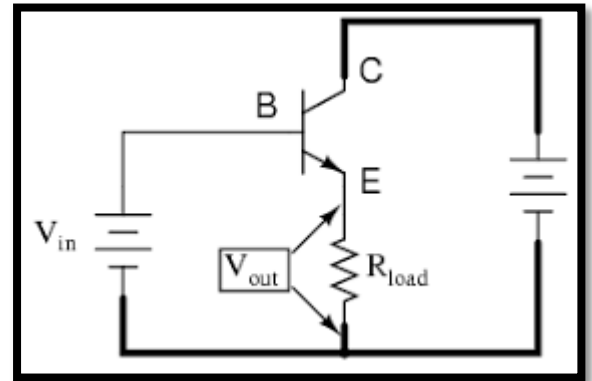
A-Common Emitter (CE) Circuit Characteristics

- 1- High voltage gain.
- 2- High current gain.
- 3- Phase reversal of the input voltage.
- 4- Low input resistance.
- 5- High output resistance.



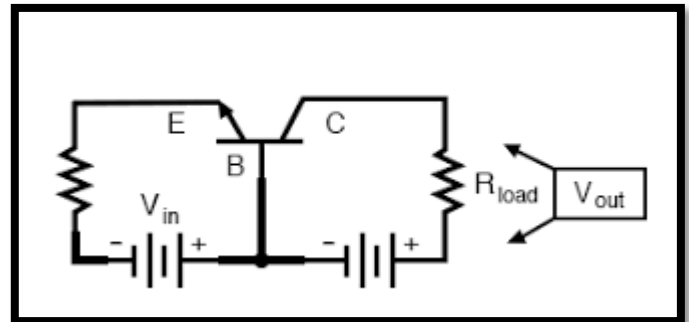
B- Common Collector (CC) Circuit Characteristics

- 1- Voltage gain =1.
- 2- High current gain.
- 3- High input resistance.
- 4- Low output resistance.
- 5- No reversal of the input voltage.

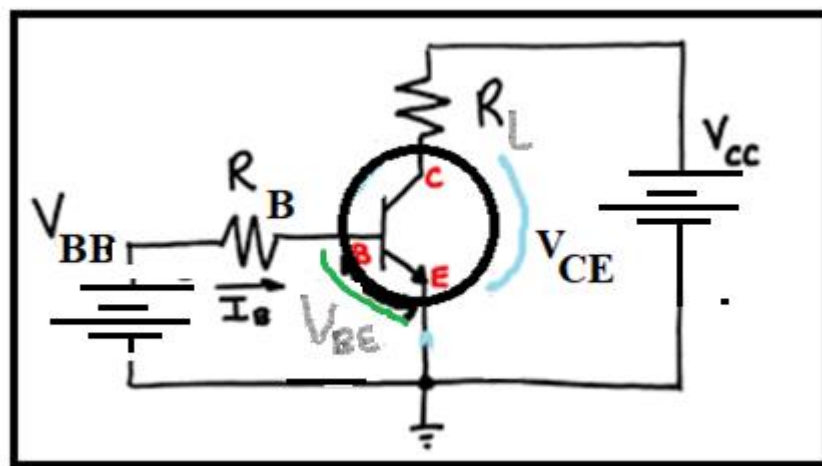


C- Common Base (CB) Circuit Characteristics

- 1- High voltage gain.
- 2- Current gain =1.
- 3- Very low input resistance.
- 4- Very high output resistance.



5- Common Emitter (CE) Circuit Characteristics



V_{CE} should be at least 1V or more to reverse-bias the collector diode of most Si transistors, since there may be an additional few tenths of a volt dropped across the bulk resistance of the collector-base region.

$$V_{CE} = V_{CB} + V_{BE} \text{ (0.7 or 0.3V)}$$

V_{BB} = is the base supply voltage.

R_B = is the current limiting resistance in the base circuit.

V_{BE} = denotes the voltage between the base and emitter.

$V_{BB} > V_{BE}$ by changing V_{BB} and R_B we can control the base current, which I_B controls I_C .

V_{CC} = is the voltage between collector and emitter.

R_C = is the current limiting resistance in the collector circuit.

6- DC Beta

The collector current is large and the base current is small. The dc beta of a transistor, also called the dc current gain for the CE connection.

Relation between α_{dc} and β_{dc}

The emitter current equals the sum of the collector current and the base current:

$$I_E = I_C + I_B \text{(1)}$$

Dividing both sides by I_C given

$$I_E / I_C = 1 + I_B / I_C \text{(2)}$$

But

$$\alpha_{dc} = I_C / I_E \text{ and } \beta_{dc} = I_B / I_C \text{(3)}$$

then by substituting 3 in 2

$$1/\alpha_{dc} = 1 + 1/\beta_{dc} = \frac{\beta_{dc} + 1}{\beta_{dc}} \text{(4) Or}$$

$$\alpha_{dc} = \frac{\beta_{dc}}{\beta_{dc} + 1} \text{(5)}$$

Lecture 6

7- Hybrid Parameters

(also known as **h parameters**) are known as 'hybrid' parameters is used 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.

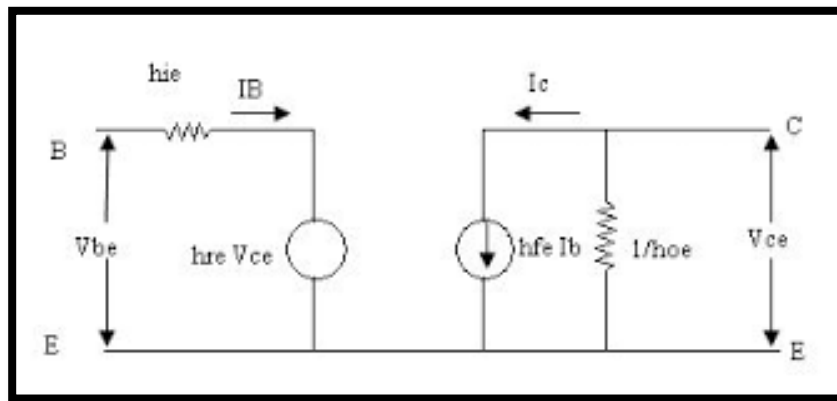
The four h **parameters of the CE connections are:**

$$h_{ie} = \left(\frac{\partial V_{BE}}{\partial I_B} \right)_{V_{CE}} = \text{input impedance (output short circuit)} (\Omega).$$

$$h_{fe} = \left(\frac{\partial i_c}{\partial i_B} \right)_{V_{CE}} = \text{forward current amplification factor}.$$

$$h_{re} = \left(\frac{\partial V_{BE}}{\partial V_{CE}} \right)_{I_B} = \text{reverse voltage amplification factor}.$$

$$h_{oe} = 1/\text{output resistance}$$



$$V_{BE} = f_1(i_B, V_{CE})$$

معادلة دائرة الدخول

$$dV_{BE} = \left(\frac{\partial V_{BE}}{\partial i_B} \right)_{V_{CE}} di_B + \left(\frac{\partial V_{BE}}{\partial V_{CE}} \right)_{i_B} dV_{CE}$$

$$v_{be} = h_{ie} i_B + h_{re} v_{ce}$$

$$i_c = f_2(i_B, V_{CE})$$

معادلة دائرة الخروج

$$di_c = \left(\frac{\partial i_c}{\partial i_B} \right)_{V_{CE}} di_B + \left(\frac{\partial i_c}{\partial V_{CE}} \right)_{i_B} dV_{CE}$$

$$i_c = h_{fe} i_B + h_{oe} V_{ce}$$

$$A_e = -\frac{V_o}{V_i} = \text{voltage gain}$$

$$A_e = -\frac{h_{fe} R_L}{h_{ie}}$$

$$A_i = h_{fe}$$

$$R_{in} = h_{ie}$$

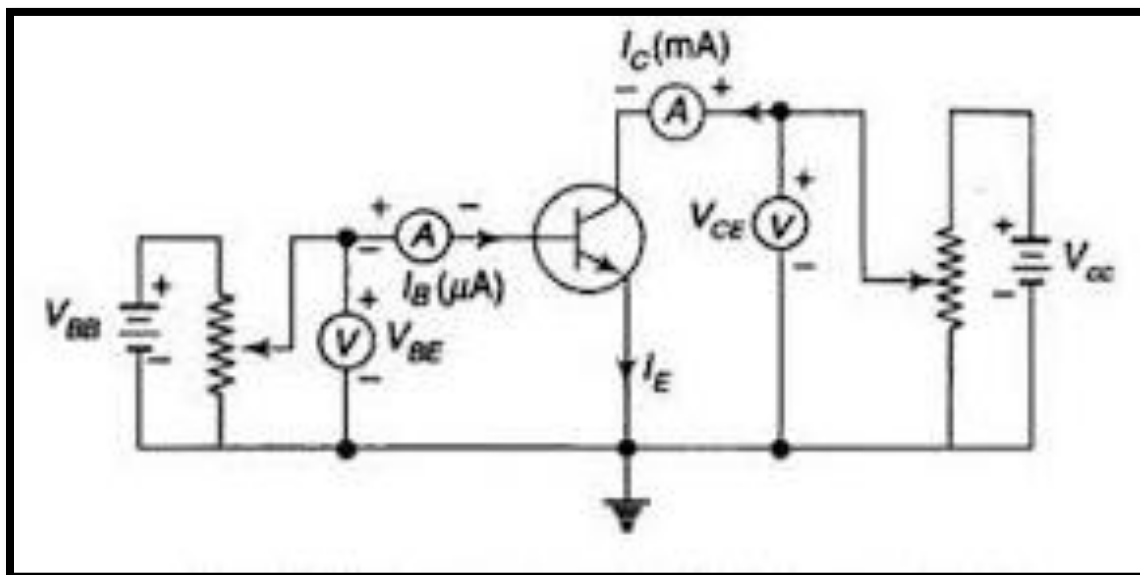
$$R_o = 1/h_{oe}$$

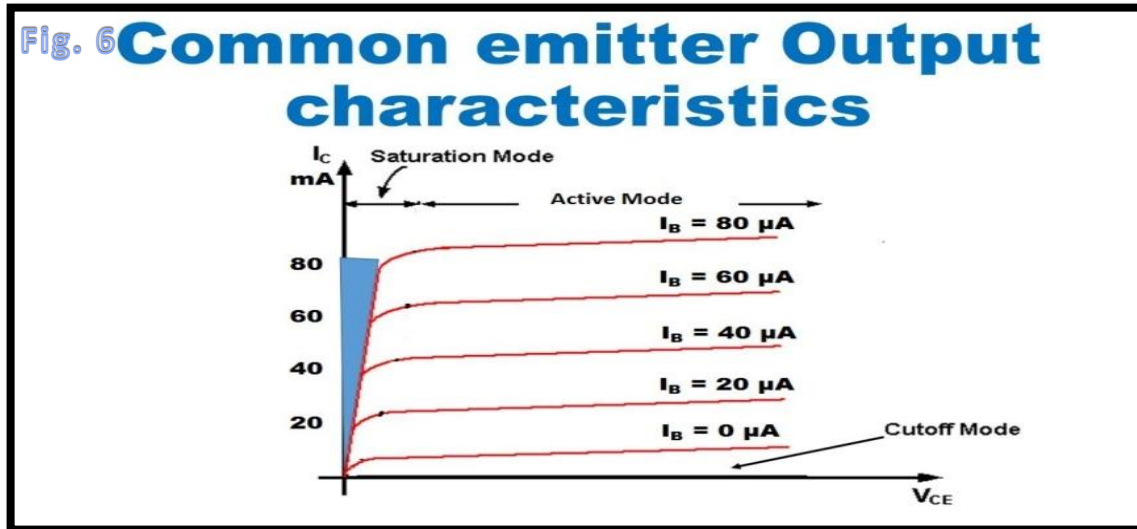
8- Output Characteristics

Fig. below shows the circuit needed to study the characteristic curves of a transistor connected in the common emitter. These characteristics depicts the relation between collector current I_C and collector –to –emitter voltage (V_{CE}) for different values of base current I_B .

V_{BB} supplies the forward bias voltage for the base-emitter junction.

V_{CC} supplies the reverse bias voltage for the collector–emitter junction.





In general the output characteristic curves show the relation between the two variables (V_{CE}, I_C) of the output circuit (CE junction), in relation to the changes of I_B in the input circuit (BE junction). So, we need to draw I_C versus V_{CE} for different values of I_B .

How we began ?

- When V_{CE} is zero, I_C is zero because the collector diode is not reverse-biased. When V_{CE} initially increases, I_C increase rapidly. In this region both junctions are forward biased, because V_{CE} is less than 0.7V. (Curve 1 of fig.6).
- The first step is to increase V_{BB} so that V_{BE} is more than the knee voltage and notice the value of I_B . For this value, increase the value of V_{CC} and for each value record the values of V_{CE} and I_C . Fig.6
- As I_B is increased, the curves go up i.e. I_C increases.

By this way, a set of curves (Figure (3.9)) is obtained of the relation between V_{CE} and I_C (in the output circuit) for different values of I_B (in the output circuit). This shows the effect of the input circuit on the output circuit.

1- The saturation region

Also called the ohmic or the linear region, the relation between V_{CE} and I_C is linear and obeys Ohm's law. It is for small values of V_{CE} . The almost vertical part of the curves near the origin is called the saturation region.

2- The active region

I_C is constant irrespective of the increase of V_{CE} values. It's the region which the curves are almost horizontal. Only if V_{CE} beyond 0.7V the collector diode becomes reversed-biased. When this happens, the collector current level off and become almost constant. The emitter junction is forward biased. When a transistor is used as an amplifier, it must operate in this region only. The base current is:

$$I_B = I_E - I_C$$

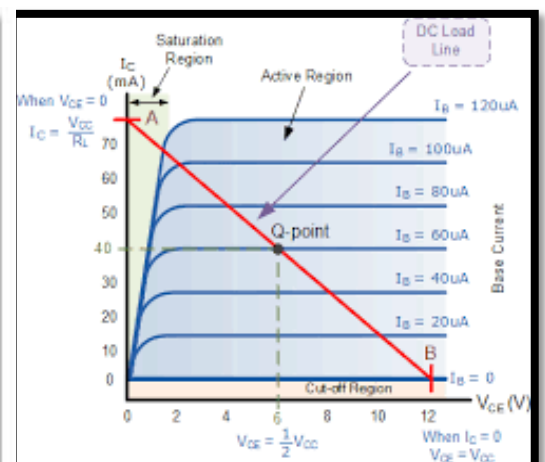
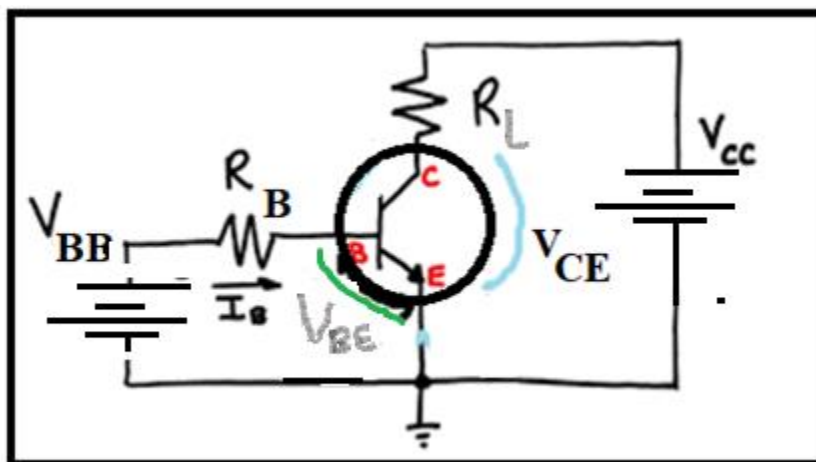
$$I_E = I_C + I_B$$

In this region, changes in V_{CE} have little effect on I_C . To change I_C , you have to change I_B .

3- Breakdown region

When the collector voltage is too large, the collector diode breakdown, indicated by a rapid increase of collector current, here the excessive power dissipation may destroy the transistor.

9- DC load Line



Remember

1- The input circuit (between the base and the emitter) whose equation is

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

2- The output circuit (between the collector and the emitter) whose equation is

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

The load line will be drawn on the output characteristics curves of the transistor, so the output circuit equation is relevant here. The voltage equation of the collector-emitter circuit is

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

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

$$I_C = \frac{V_{CC}}{R_L} - \frac{V_{CE}}{R_L}$$

$$I_C = -\frac{1}{R_L} V_{CE} + \frac{V_{CC}}{R_L}$$

This is a straight equation

$$y = -mx + b$$

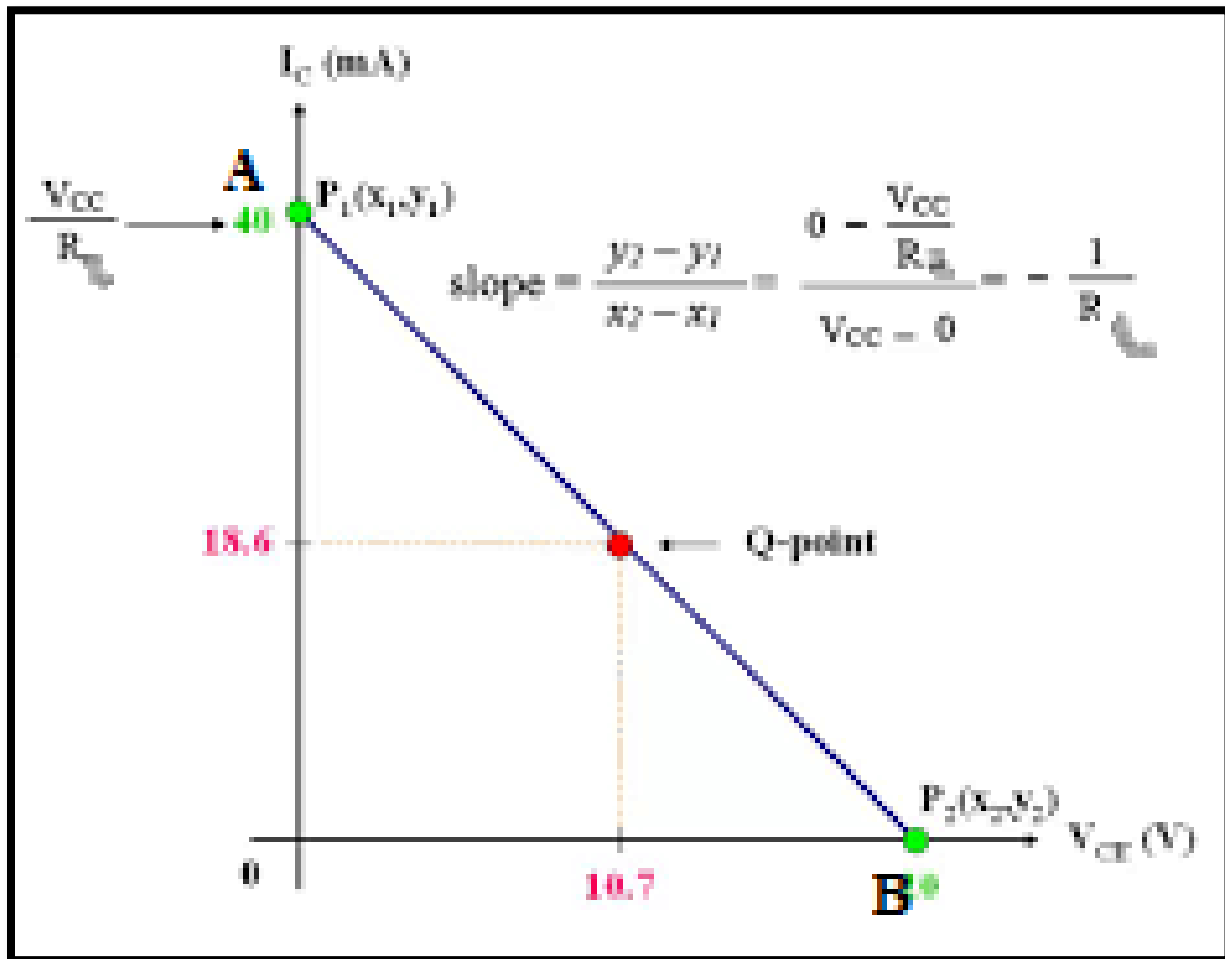
by comparing the two equations

I_C is on the Y-axis

V_{CE} is on the X-axis

$$I_C = \underbrace{-\frac{1}{R_L}}_m V_{CE} + \underbrace{\frac{V_{CC}}{R_L}}_b$$

You can notice the slope is negative (m). The b is the intersection point of the load line on the y-axis.



- A is the point of intersection at the upper end of the dc load line, it called saturation point. At this point, the collector current is $= \frac{V_{CC}}{R_C}$. where $V_{CE}=0$ so all V_{CC} is converted to current and it is the maximum current (saturation current) that can pass through this circuit.
- B is the cutoff point, where $V_{CE}=V_{CC}$, where $I_C =0$ there is no current through the transistor, it is not conducting, the transistor is cut-off.
- the quiescent point (Q-point), which is the intersection of the load line with any of the characteristics curves.

Consider the following particular cases:

1- When $I_C=0$, $V_{CC}=V_{CE}$

CUTOFF POINT

2- When $V_{CE}=0$, $I_C = \frac{V_{CC}}{R_C}$

SATURATION POINT

Q-point (Quiescent)

It is a point on the dc load line which represents values of I_C and V_{CE} that exist in a transistor circuit when no input signal is applied. It is also known as the dc operating point or working point.

You can notice there are many Q- points that intersection the load line with curve $I_B = 0$. So how can we choose?

The best position of the Q-point should be at the middle of the curves at the active region and in the middle of the active region (a position corresponding to the middle of the load line). The Q-point is also known as the DC operating point.

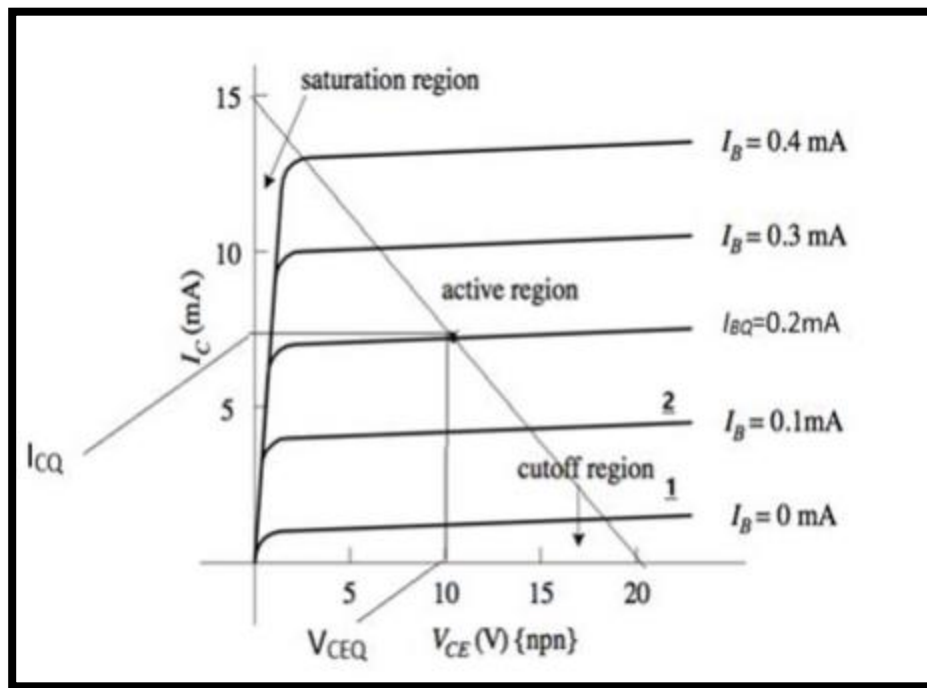


Fig.7

According to this the Q-point is chosen as seen in fig.7 which is the interception of the load line with the curve of $I_B=0.2\text{mA}$, So the position of the Q-point in this example is at $I_B=0.2\text{mA}$, $I_C=7\text{mA}$, $V_{CE}=10\text{V}$

The position of the Q-point is defined by three quantities: I_{BQ} , I_{CQ} , V_{CEQ} .

Example-1:

For the given circuit, the transistor used is NPN-Si with $\beta = 100$. Find the position of the Q-point. Knowing that $I_B = 20\mu A$, $R_L = 4K\Omega$, $V_{CC} = 12V$.

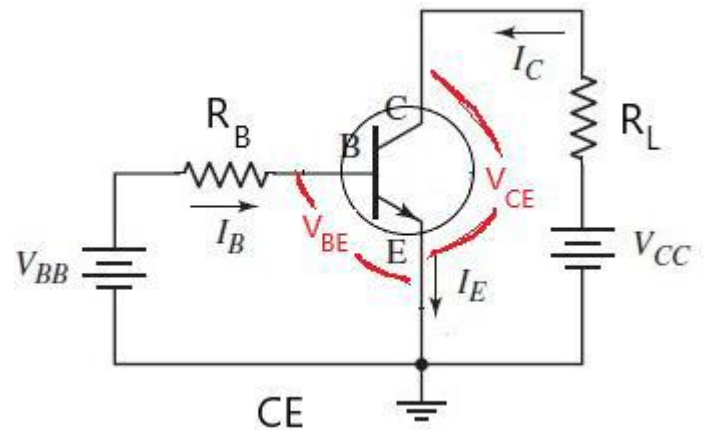
$$\beta = \frac{I_C}{I_B}$$

$$100 = I_C / 20\mu$$

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

$$12 = 2m \times 4k\Omega + V_{CE}$$

So the position of the Q-point is at: $V_{CE} = 4V$ and $I_C = 2mA$



$$I_C = 2mA$$

$$V_{CE} = 4V$$

Example 2:

For the given circuit, the transistor is NPN-Ge transistor with $\beta = 150$. The position of the Q-point is at $V_{CE} = 10V$. Find the value of R_B . Knowing that $V_{BB} = 10V$, $V_{CC} = 25V$, $R_L = 3k\Omega$

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

$$25 = I_C \times 3k\Omega + 10$$

$$150 = 5m / I_B$$

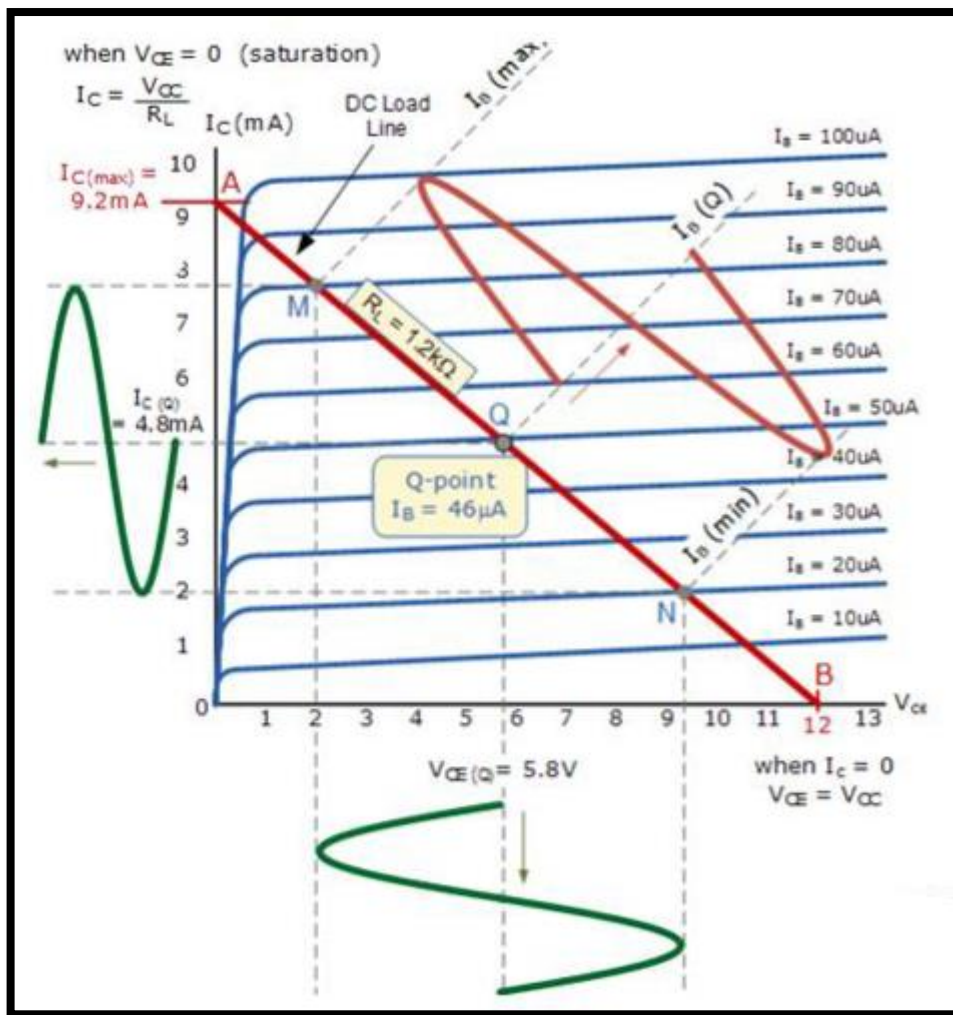
$$V_{BB} = I_B R_B + V_{BE} \quad 10 = 33\mu \times R_B + 0.3$$

$$I_C = 5mA$$

$$I_B = 33\mu A$$

$$R_B = 293K\Omega$$

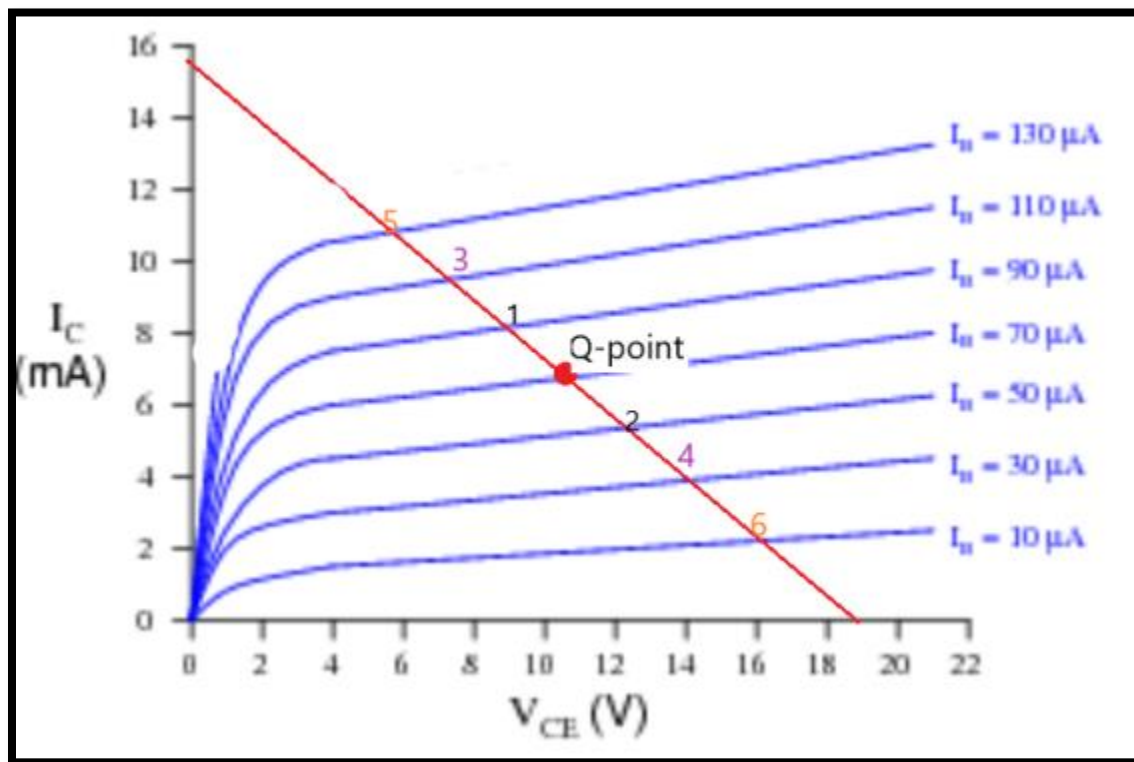
After choosing the Q-point and knowing the dc values of voltage and currents at the Q-point which are I_{BQ} , I_{CQ} and V_{CEQ} , an ac voltage to be amplified is introduced to the circuit between the emitter and base (the input circuit), this will cause the I_B to change according to the change of the input ac voltage, accordingly V_{CE} and I_C will change(in the output circuit) i.e. the Q-point, which is now called the **operating point**, will change its position along the load line.



According to Fig. (7), the position of the Q-point is at $I_{BQ} = 50 \mu\text{A}$, $I_{CQ} = 4.8 \text{ mA}$, $V_{CEQ} = 5.8 \text{ V}$.

An ac input signal is introduced such that it causes I_B to rise up to $80\mu\text{A}$ (peak value in the positive direction) and go down to $20\mu\text{A}$ (peak value in the negative direction) i.e. $I_B=60\mu\text{A}$, the Q-point has moved along the load line. As the Q-point goes up to $I_B=80\mu\text{A}$, I_C changes from its dc quiescent value at 4.8mA to the peak value at the positive direction = 8mA and down to the peak value in the negative direction = $2\mu\text{A}$. But, as the I_B values increases from $50\mu\text{A}$ to $80\mu\text{A}$, V_{CE} will decrease will change between 2V and 9.2V , (from 5.8V to 2V (the negative half) and up to 9.2V in the positive half as I_B decreases to $20\mu\text{A}$ in the negative half). This shows that when there a positive half cycle of the voltage at the input, there will be a negative half of the voltage at the output, and vice versa. This explains the minus sign of the gain, which is the 180° phase difference between the input and the output voltage.

The 180° phase difference between the input and the output voltage.



The input voltage should be of value such that the change of I_B is between points 5 and 6, which is the capacity of the amplification of the amplifier; otherwise distortion will result (as will be explained later).

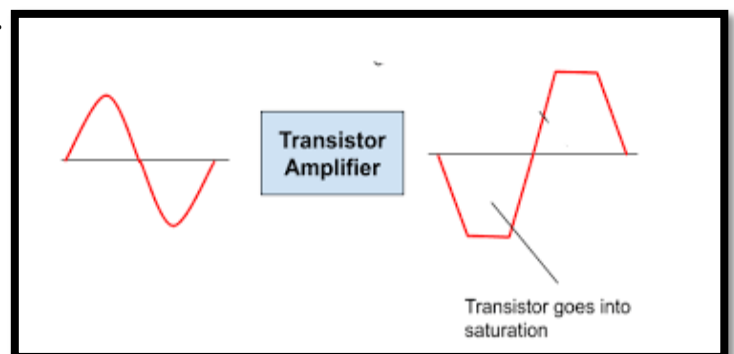
So far, the amplification process has been explained but still the circuit is not suitable to be used as an amplifier for it lacks the correct biasing, which is important for a fixed position of the Q-point and so to eliminate the effects of the increase of temperature.

The position of the Q-point must not change (fixed values of I_B , I_C , V_{CE}) during the process of amplification. This is done by correct and suitable biasing of the amplifier circuit (this will be discussed in the biasing section). Good biasing and good position of the Q-point will ensure that the output voltage is the exact replica of the input voltage (the output voltage waveform is the same as the input voltage, only bigger, amplified).

Reasons for distortion of the output voltage

1- Incorrect position of the Q-point: this will cause clipping either on the negative cycle (when the Q-point is in the lower half of the load line) or on the positive cycle (when the Q-point is in the upper half of the load line).

2- Large input signal, higher than what the amplifier can amplify. In this case saturation will occur. Clipping in both the positive and negative halves of the input voltage occurs and the output wave is not a sine wave as the input. This can happen even with correct biasing of the circuit.



Lecture 7

Transistors

Stabilization

It is the process of making Q-point independent of temperature changes or variation in transistor parameters. i.e. the aim of stabilization is to make the position of the Q-point is independent of:

- Temperature changes.
- Individual variations (hybrid parameters).
- Thermal runaway

Temperature dependence of I_C & Thermal runaway

As mentioned before, the position of Q-point should be stable, hence no change in its position, so the values of I_B , I_C , V_{CE} should not change and stay fixed at the value chosen, which should be change only when the AC voltage to be amplified is introduced.

But in reality, this is not what happened, because the changes of temperature will arise the minority carriers. The collector current is organized from the emitter depending on the relation:

$$I_C = \beta I_B$$

To this, is added a small current that is due to the minority carriers (leakage current) which are produced because of the rise of temperature of the transistor. So emitter current becomes:

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

where I_{CO} is the leakage current.

This leakage current causes many problems one of which is the instability of the position of the Q-point. Therefore, a solution must be found to overcome this.

Thermal Runaway

- As current passes through the transistor, its temperature increases, so leakage current (I_{CO}) will be produced
- I_{CO} is strong function of temperature. The flow of I_C produces heat within the transistor and raises the transistor temperature further and therefore, further increases in I_{CO} .
- It doubles for every 10°C rise in temperature in case of Si. An increase of I_{CO} is magnified by $(1+\beta)$. So even a very small I_{CO} will cause a large increase of I_C .
- This effect is cumulative and in few seconds, the I_C may become larger enough to cause the destruction of the crystal structure of the transistor burn out the transistor.
- The self-destruction of an unstablized transistor is known as thermal runaway.
- Obviously, this process has to be stopped and this is done by the thermal stabilization of the transistor circuit.

Stability Factor

The rate of change collector current I_C with respect to the collector leakage current I_{CEO} is called **stability factor**, denoted by S .

$$S = 1 + \frac{R_B}{R_E}$$

Biasing

1. To eliminate one of the batteries, V_{BB} is eliminated.
2. Forward biasing on the BE junction must be maintain and so for the reverse biasing of the BC junction.

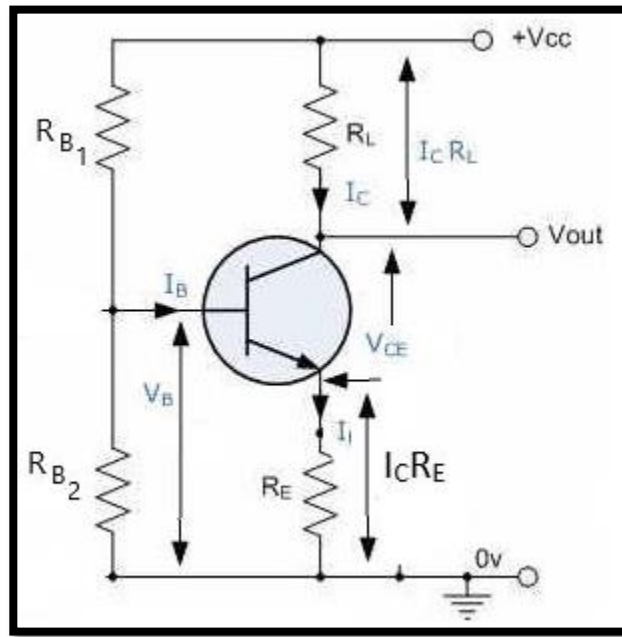
Biasing a transistor means applying external voltages to produce a desired collector current, because of the collector resistance R_L , the value of V_{CE} is less than the supply voltage V_{CC} .

There are many circuits to satisfy the aims of biasing and the thermal stabilization (mentioned above) but the best is “the voltage divider self-biased circuit”

Voltage-Divider Self-Biased Common Emitter Circuit

The single stage common emitter amplifier circuit shown below uses what is commonly called “Voltage Divider Biasing”. This type of biasing arrangement uses two resistors as a potential divider network across the supply with their center point supplying the required Base bias voltage to the transistor. Voltage divider biasing is commonly used in the design of bipolar transistor amplifier circuits.

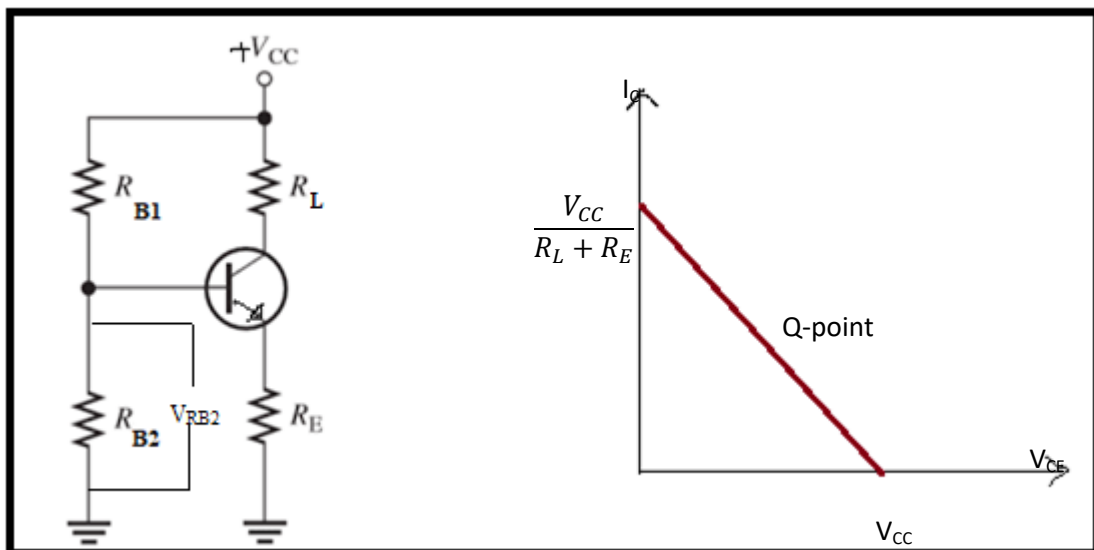
This method of biasing the transistor greatly reduces the effects of varying Beta, (β) by holding the Base bias at a constant steady voltage level allowing for best stability. The quiescent Base voltage (V_B) is determined by the potential divider network formed by the two resistors, R_1 , R_2 and the power supply voltage V_{CC} as shown with the current flowing through both resistors.



From this figure:

- V_{BB} was eliminated.
- R_{B1} & R_{B2} were added. They are connected to V_{CC} what is known as voltage divider circuit.
- The voltage level generated at the junction of resistors R_{B1} and R_{B2} hold the Base voltage (V_B) constant at a value below the supply voltage.

When the circuit is properly designed, the voltage across R_{B2} forward-biases the emitter diodes and produces a collector current that is almost independent of β_{dc} . This is the main reason for the great popularity of voltage-divider bias.



Equations

This bias reference voltage can be easily calculated using the simple voltage divider formula below:

The voltage across R_{B2} is approximately, V_{CC} = voltage divider circuit

$$V_{R_{B2}} = \frac{R_{B2}}{R_{B2} + R_{B1}} V_{CC} \dots\dots\dots(1)$$

In this equation, the V_{CC} voltage connected to the circuit is supplies the voltage to the input circuit hence $V_{R_{B2}}$ it is a part of V_{CC} with a certain partial which is equal to the amount of the resistance R_{B2} to the summation of the self-bias resistance ($=R_{B2} + R_{B1}$) because the resistance connected parallel.

The voltage across the emitter resistor equals:

$$V_E = V_{R_{B2}} - V_{BE}$$

$$V_{R_{B2}} = I_C R_E + V_{BE} \dots\dots\dots(2) \quad \text{Input circuit (I}_C \sim I_E)$$

Therefore, the dc emitter current is

$$I_E = \frac{V_E}{R_E} \quad \text{or} \quad I_E = \frac{V_{R_{B2}} - V_{BE}}{R_E}$$

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

$$V_{CC} = I_C (R_L + R_E) + V_{CE} \quad \text{output circuit}$$

$$R_B = \frac{R_{B1} R_{B2}}{R_{B1} + R_{B2}}$$

The benefit of R_E is to control the thermal stabilization.

.

EX1: what is the dc voltage from the collector to the ground? $R_{B1}=10K\Omega$, $R_{B2}=20K\Omega$, $R_C=4K\Omega$, $R_E=5K\Omega$, $V_{CC}=30V$

$$V_{12}=V_1-V_2=IR$$

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

$$V_{CC}-V_C=I_C R_L$$

$$V_{R_{B2}} = I_C R_E + V_{BE}$$

$$\frac{R_{B2}}{R_{B2}+R_{B1}} V_{CC} = I_C R_E + V_{BE}$$

$$\frac{10K\Omega}{(10+20)K\Omega} 30V = I_C (5K\Omega) + 0.7V$$

$$10V - 0.7V = I_C (5K\Omega)$$

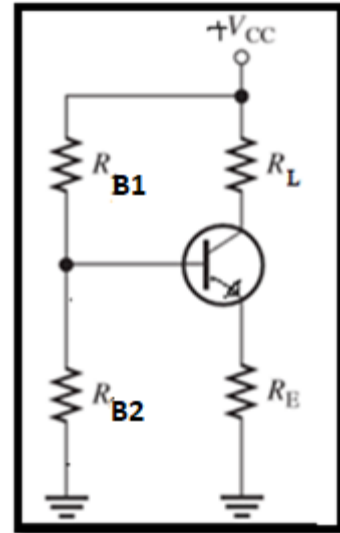
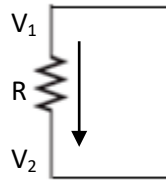
$$I_C = 1.86mA$$

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

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

$$30V - V_C = (1.86mA)(4K\Omega)$$

$$V_C = 22.6V$$



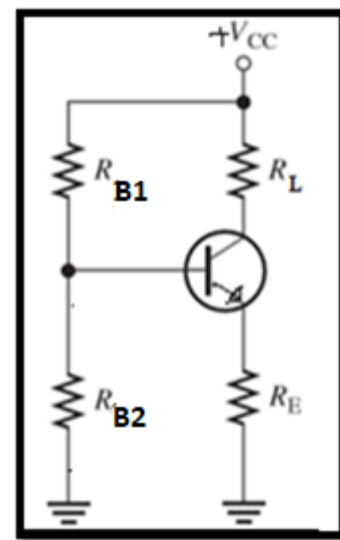
EX: Design a voltage divider self biased common emitter circuit using Si NPN transistor. The circuit is supplied with $V_{CC}=15V$, and it biased at $I_{CQ}=10mA$ and $V_{CEQ}=7V$. Knowing that $R_E=330\Omega$, and $R_B=8.3K\Omega$.

$$V_{CC} = I_{CQ}(R_L + R_E) + V_{CEQ}$$

$$15V = 10mA(R_L + R_E) + 7V$$

$$R_L + R_E = 800\Omega$$

$$R_L = 800\Omega - R_E = 800\Omega - 330\Omega = 470\Omega$$



$$V_{R_{B2}} = I_C R_E + V_{BE}$$

$$V_{R_{B2}} = (10\text{mA} * 470\Omega) + 0.7\text{V}$$

$$V_{R_{B2}} = 5.4\text{V}$$

$$V_{R_{B2}} = \frac{R_{B2}}{R_{B2} + R_{B1}} V_{CC}$$

$$5.4\text{V} = \frac{R_{B2}}{R_{B2} + R_{B1}} * 15\text{V}$$

$$\frac{R_{B2}}{R_{B2} + R_{B1}} = 0.36$$

$$R_B = \frac{R_{B2}}{R_{B1} + R_{B2}} R_{B1}$$

$$8.3\text{K}\Omega = 0.36 R_{B1}$$

$$R_{B1} = 23.1\text{ K}\Omega$$

$$R_{B2} = 12.95\text{ K}\Omega$$

Important equations

$$V_{BB} = I_B R_B + V_{BE} \text{ (input)}$$

$$\text{Active region, } I_C = \beta_{dc} I_B$$

This collector current produces a voltage drop of $I_C R_C$ across the collector resistance. Therefore, the collector-emitter voltage is

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

$$V_{CC} = I_C R_L + V_{CE} \text{ (output)}$$

Lecture 8

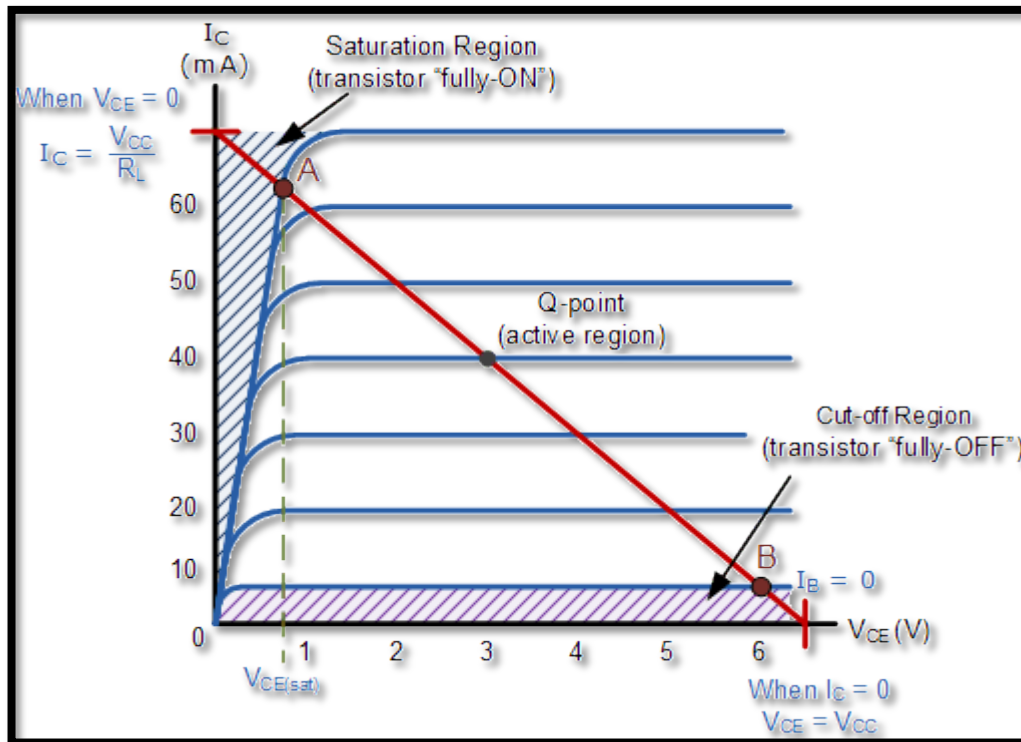
Transistor as a Switch

Transistor switches can be used to switch a low voltage DC device (e.g. LED's) ON or OFF by using a transistor in its saturated or cut-off state. When used as an AC signal amplifier, the transistor is always operates within its “active” region, that is the linear part of the output characteristics curves .Any point along the load line can be used as the Q-point..

However, both the NPN & PNP type bipolar transistors can be made to operate as “ON/OFF” type solid state switch by suitable biasing, so that only two points are employed, the cut-off point (open switch) or the saturation point (closed switch) the transistors Base terminal differently to that for a signal amplifier.

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

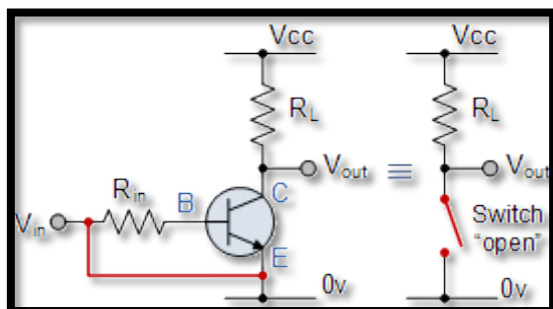
The areas of operation for a transistor switch are known as the **Saturation Region** and the **Cut-off Region**. This means then that we can ignore the operating Q-point biasing and voltage divider circuitry required for amplification, and use the transistor as a switch by driving it back and forth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



1. Cut-off Region

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

Cut-off Characteristics



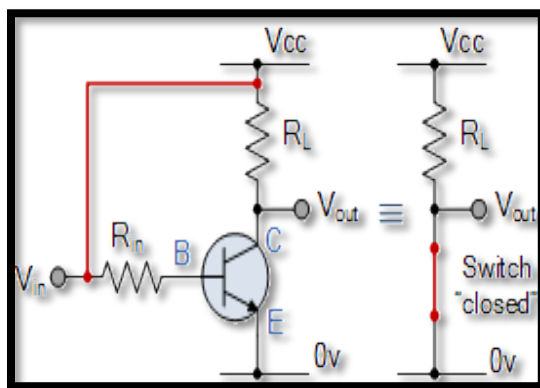
- The input and Base are grounded (0v)
- Base-Emitter voltage $V_{BE} < 0.7v$
- Base-Emitter junction is reverse biased
- Base-Collector junction is reverse biased
- Transistor is "fully-OFF" (Cut-off region)
- No Collector current flows ($I_C = 0$)
- $V_{OUT} = V_{CE} = V_{CC} = "1"$
- Transistor operates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

Then we can define the “cut-off region” or “OFF mode” when using a bipolar transistor as a switch as being, both junctions reverse biased, $V_{BE} < 0.7\text{v}$ and $I_C = 0$.

2. Saturation Region

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

Saturation Characteristics



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

Then we can define the “saturation region” or “ON mode” when using a bipolar transistor as a switch as being, both junctions forward biased, $V_{BE} > 0.7\text{v}$ and $I_C = \text{Maximum}$. With a zero signal applied to the Base of the transistor it turns “OFF” acting like an open switch and zero collector current flows. With a positive signal applied to the Base of the transistor it turns “ON” acting like a closed switch and maximum circuit current flows through the device.

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

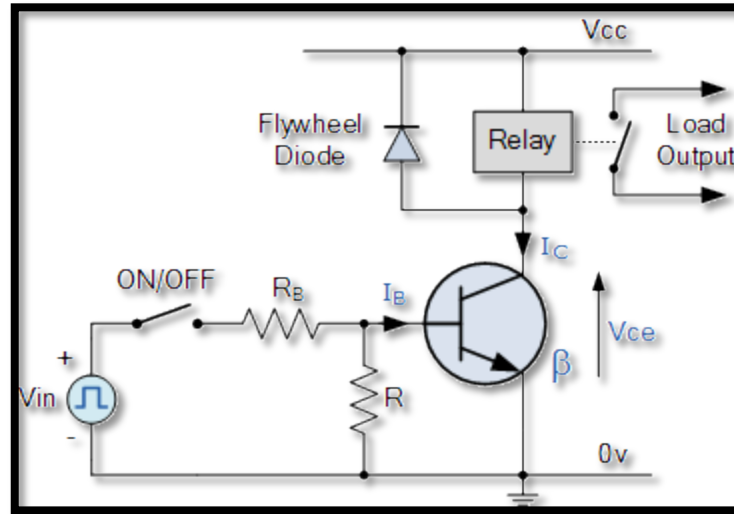
Basic NPN Transistor Switching Circuit

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

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

In order for the Base current to flow, the Base input terminal must be made more positive than the Emitter by increasing it above the 0.7 volts needed for a silicon device. By varying this Base-Emitter voltage V_{BE} , the Base current is also altered and which in turn controls the amount of Collector current flowing through the transistor as previously discussed.

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



Transistor as a Switch Example No1

Using the transistor values of: $\beta = 200$, $I_C = 4\text{mA}$ and $I_B = 20\mu\text{A}$, find the value of the Base resistor (R_B) required to switch the load fully “ON” when the input terminal voltage exceeds 2.5v.

$$R_B = \frac{V_{in} - V_{BE}}{I_B} = \frac{2.5\text{V} - 0.7\text{V}}{20 \times 10^{-6}} = 90\text{k}\Omega$$

Transistor as a Switch Example No2

Again using the same values, find the minimum Base current required to turn the transistor “fully-ON” (saturated) for a load that requires 200mA of current when the input voltage is increased to 5.0V. Also calculate the new value of R_B .

Transistor Base current:

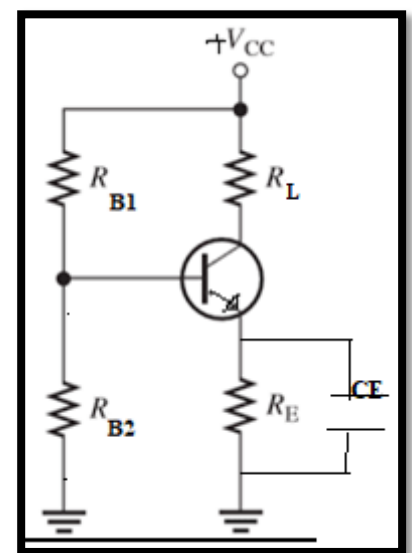
$$I_B = \frac{I_C}{\beta} = \frac{200\text{mA}}{200} = 1\text{mA}$$

Transistor Base resistance:

$$R_B = \frac{V_{in} - V_{BE}}{I_B} = \frac{5.0\text{V} - 0.7\text{V}}{1 \times 10^{-3}} = 4.3\text{k}\Omega$$

Transistor switches are used for a wide variety of applications such as interfacing large current or high voltage devices like motors, relays or lamps to low voltage digital IC's or logic gates like AND gates or OR gates. Here, the output from a digital logic gate is only +5v but the device to be controlled may require a 12 or even 24 volts supply. Or the load such as a DC Motor may need to have its speed controlled using a series of pulses (Pulse Width Modulation). Transistor switches will allow us to do this faster and more easily than with conventional mechanical switches.

C_E is called bypass capacitance, it connected in parallel with R_E to prevent the A.C. current to enter this resistance. it has large capacity and very small impedance.



Lecture 9

Field Effect Transistor

Field Effect Transistor (FET) is a unipolar device and its operation depends on only one type of charge carriers either holes or free electrons, thus FET has majority carriers not minority carriers. While BJT has two types of charge carriers, i.e. holes and free electrons and thus we call it a bipolar device.

The Field Effect Transistor, FET, is a three terminal active device that uses an electric field to control the current flow and it has high input impedance which is useful in many circuits. It is a key electronic component using within many areas of the electronics industry.

A FET is a voltage-controlled semiconductor device. Its advantages over bipolar junction transistors are:

- 1- FET is a unipolar device. Transistor is bipolar.
- 2- FET is a voltage controlled device. Transistor is a current controlled device.
- 3- FET has very high input resistance. Transistor has low input resistance.
- 4- FET has small voltage gain. Transistor has higher voltage gain than the FET.
- 5- The gate in the FET is reverse biased. The base in the transistor is forward biased.

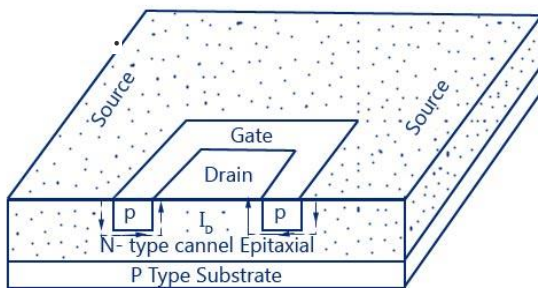


There are two types of FET: Junction FET (JFET) and metal oxide semiconductors (MOSFET) or insulated gate FET.

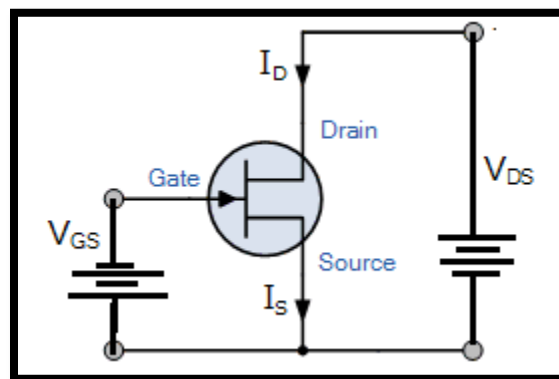
1- Junction Field Effect Transistor (JFET)

a) Structure

A JFET has a base semiconductor known as channel; it can be N or P-type. A voltage supply across the ends of the channel causes the flow of current. The two ends of the channel are known as source and drain. If the channel is N-type the current is due to the electrons, while if the channel is P-type the current is due to holes. The majority carriers enter the channel from source (S) and source current is denoted as I_S . The current entering the drain is denoted by I_D . The drain is positive with respect to the source (drain to source voltage V_{DS}). If the channel is N-type materials, two heavily doped p-type regions are found on opposite sides of the channel. Each of these p-regions is called a gate; they are connected together and appear with only one terminal. There is no current in the gate because of the reverse biasing on the gate –source junction. There is a junction between the gate and the channel so there exists a depletion layer on both sides of the junction in the channel and in the gate. The depletion layer on the gate side is of no importance but the depletion layer in the channel side is of importance in the operation of the JFET. Since the gate is on both sides of the channel the depletion layers on both sides of the channel will reduce the width of the channel thus reducing the current from the source to the drain.

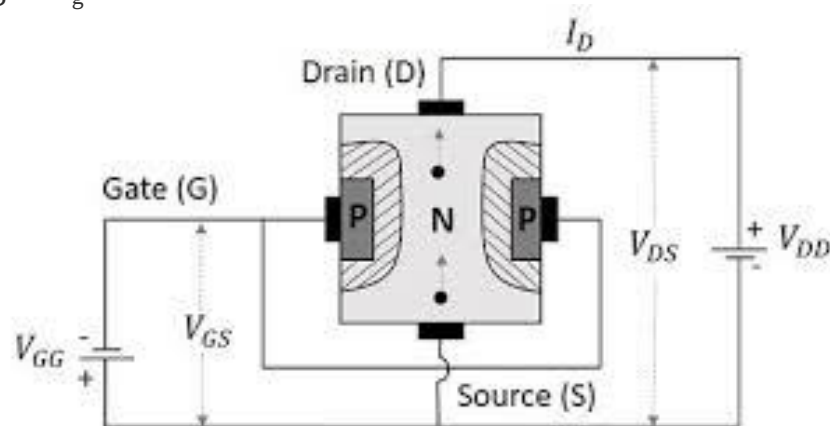


Geometry of JFET

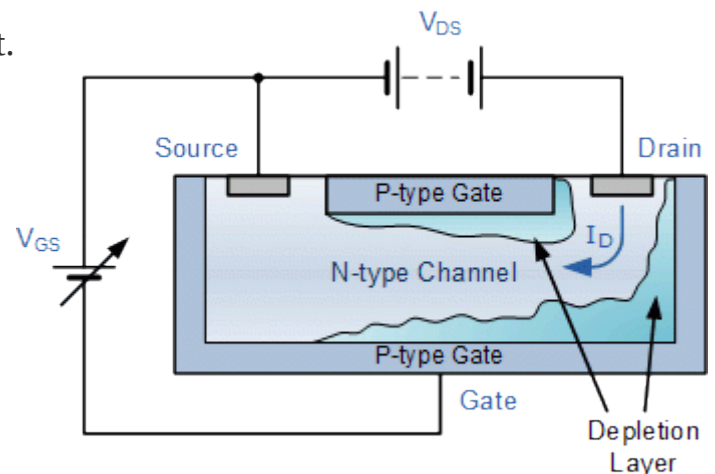


Operation

- ❖ Figure below shows the normal polarities for biasing an N-channel JFET. The idea is to apply a negative voltage between the gate and the source. Since the gate is reversed-biased, only a very small reverse current flows in the gate $I_g=0$.



- ❖ The name field effect is related to the depletion layers around each PN junction. Free electrons moving between the source and the drain must flow through the narrow channel between depletion layers. The size of these depletion layers determines the width of the conductive channel. The more negative the gate voltage the narrower the conductive channel becomes, i.e. the conductivity is reduced, because the depletion layers get closer to each other. Therefore, the gate voltage controls the current that flows between the source and drain. The more negative the gate voltage, the smaller the current.



b) Symbol and circuit configuration

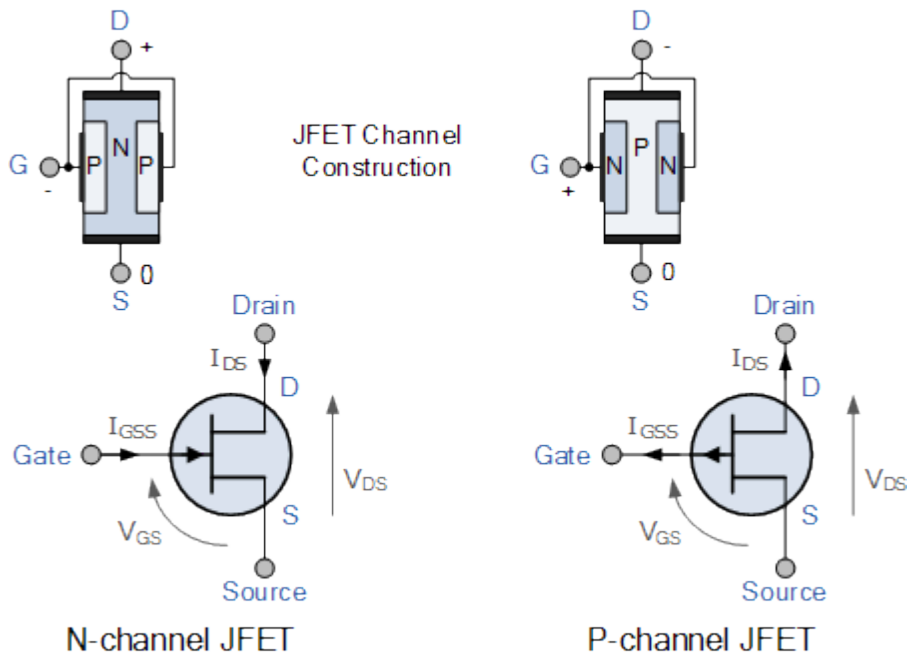


Fig. above shows the symbol and circuit for n- and p-channel JFET. For n-channel JFET, I_D and V_{DS} are positive and V_{GS} is negative.

c) JFET Characteristics

The curves plotted between the current value at the drain and the voltage applied in between drain and the source for different values of the gate-source voltage are the output characteristic curve that are also referred to as the **drain characteristics**.

The regions of the drain characteristics of the JFET are

1. **Ohmic Region:** When small values of V_{GS} the depletion layer of the channel is very small and the JFET acts like a voltage controlled resistor. It is also known as the linear region.

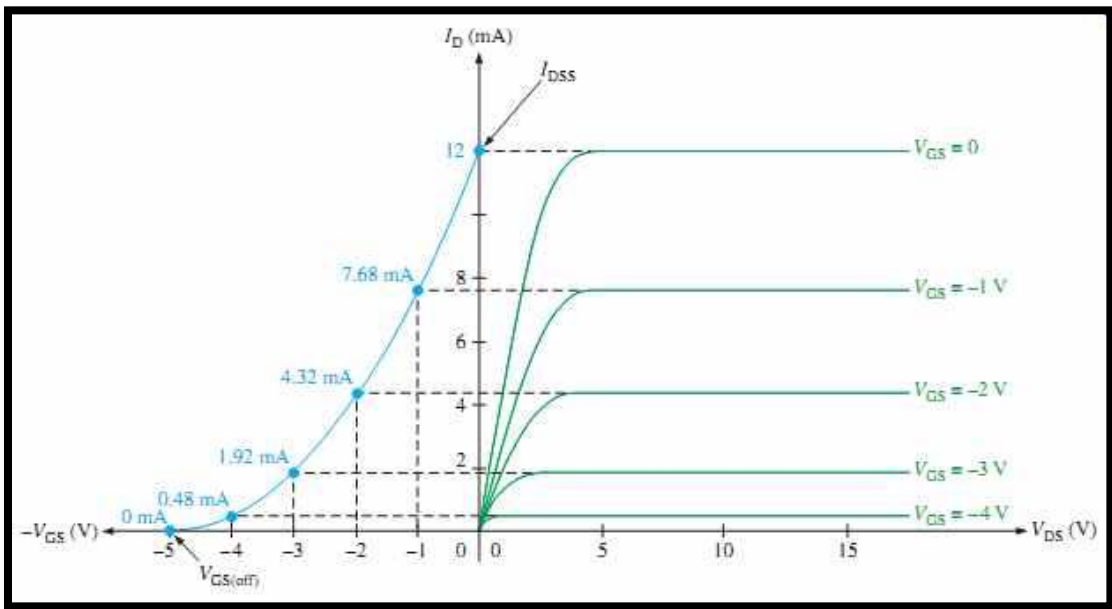
2. **Saturation or Active Region** – The JFET becomes a good conductor and is controlled by the Gate-Source voltage, (V_{GS}) while the Drain-Source voltage, (V_{DS}) has little or no effect.
3. **Breakdown Region** – The voltage between the Drain and the Source, (V_{DS}) is high enough to causes the JFET's resistive channel to break down and pass uncontrolled maximum current. But if the drain to source voltage is increased further then the device reaches the breakdown region in which the drain current increases indefinitely.

NOTES

- ❖ **Cut-off Region:** This is also known as the **pinch-off region** where the Gate voltage, V_{GS} is sufficiently negative so as to cause the channel to close making the drain current equal to zero. As the drain voltage is increased the channel tends to become narrower and narrower and current at the drain terminal gets smaller. At a particular drain to source voltage called the **pinch-off voltage** the drain current reaches the saturation level. [Cut off region is due to V_{GS} but the pinch off voltage is due to V_{DS}].
- ❖ **Pinch –off voltage V_p :** is the drain voltage above which the drain current become almost constant for shorted gate condition. When the drain voltage equals V_p , the conducting channel becomes extremely narrow and the depletion layers almost touch, further increase in drain voltage produce only the slightest increase in drain current.
- ❖ Now if a negative voltage is applied to the gate terminal then, in that case, the channel present at the gate reverse biases and the saturation current starts decreasing further. At a particular gate voltage the device stops conduction this is called the **cut-off-voltage**.

- ❖ **Shorted-gate condition**- if the gate voltage is reduced to zero, the gate is effectively shorted to the source. This is called the short-gate condition.
- ❖ I_{DSS} - it is the maximum drain current at shorted gate.

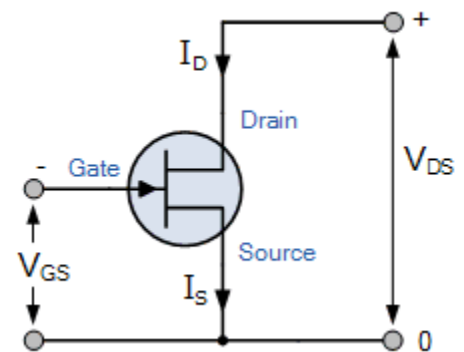
When $V_{GS} = V_{GS(off)}$ the depletion layers touch, cutting off the drain current. Since V_p is the drain voltage that pinches off current for the shorted-gate condition. $V_p = -V_{GS(off)}$



Transconductance

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

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



Drain-Source Channel Resistance

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

$$g_m = -\frac{2I_{DSS}}{V_{GS(OFF)}} \left[1 - \frac{V_{GS}}{V_{GS(OFF)}} \right]$$

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

d) JFET Biasing

1- Self-Bias

The idea is to use the voltage across the source resistor R_s to produce the gate-source reverse voltage.

From the above fig., since V_{GS} is reverse voltage negligible I_g flows through R_G . Therefore, the gate voltage with respect to ground is zero:

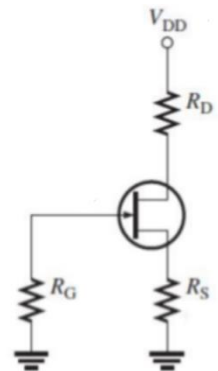
$$V_G = 0$$

The source voltage

$$V_s = I_d R_s$$

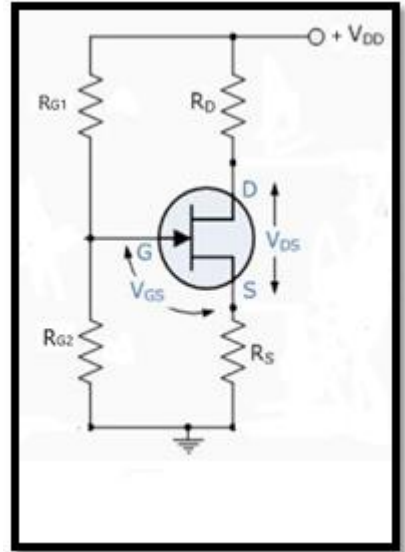
$$V_{GS} = V_G - V_s = -I_d R_s, \quad V_G = 0$$

R_s make a reverse bias between the gate and the source and maintain the value of I_d constant.



2- Voltage divider bias

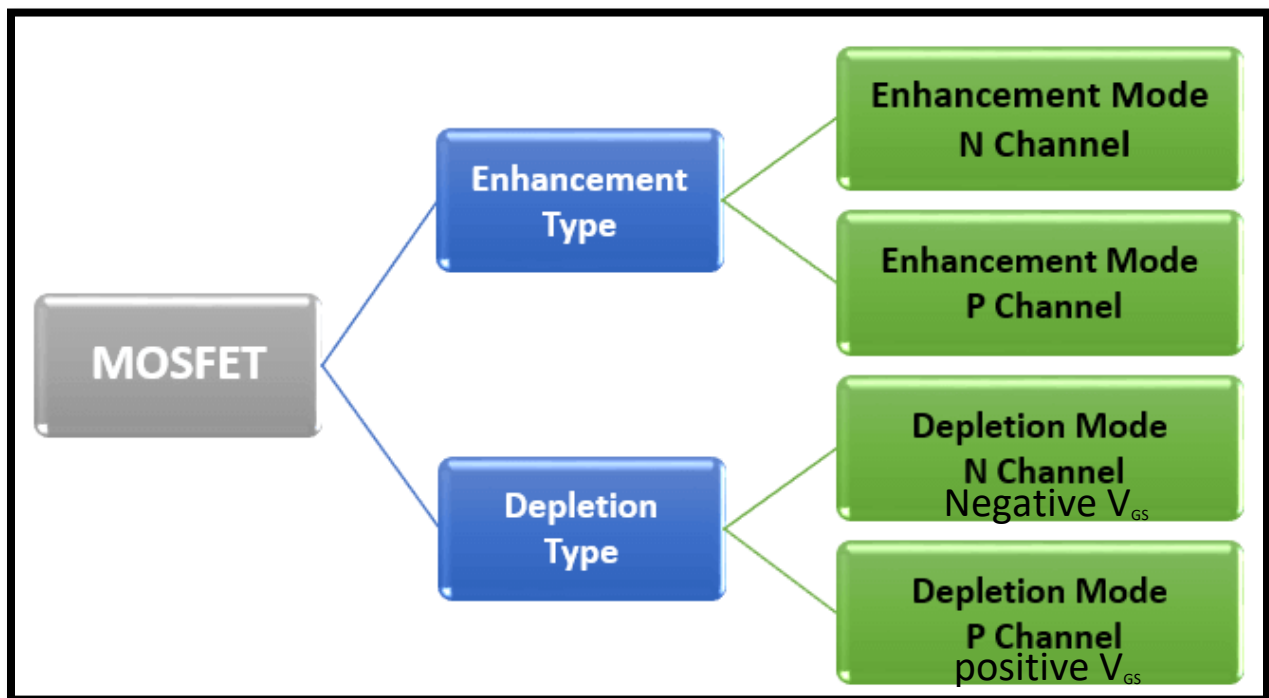
$V_{GS} = V_G - V_S$ for V_{GS} to be negative $V_S > V_G$



Lecture 10

2-Metal Oxide Semiconductor FET (MOSFET)

The MOS (metal-oxide- semiconductor) transistor (or MOSFET) is the basic building block of most computer chips, as well as of chips that include analog and digital circuits. The **MOSFET** is a voltage controlled field effect transistor that differs from a JFET in that it has a gate is insulated (by a metal oxide layer) from the channel which is electrically insulated from the main semiconductor n-channel or p-channel by a very thin layer of insulating material usually silicon dioxide.



1- Depletion Type MOSFEET

Because the gate is insulated from the channel, one can apply negative or positive voltages to the gate,

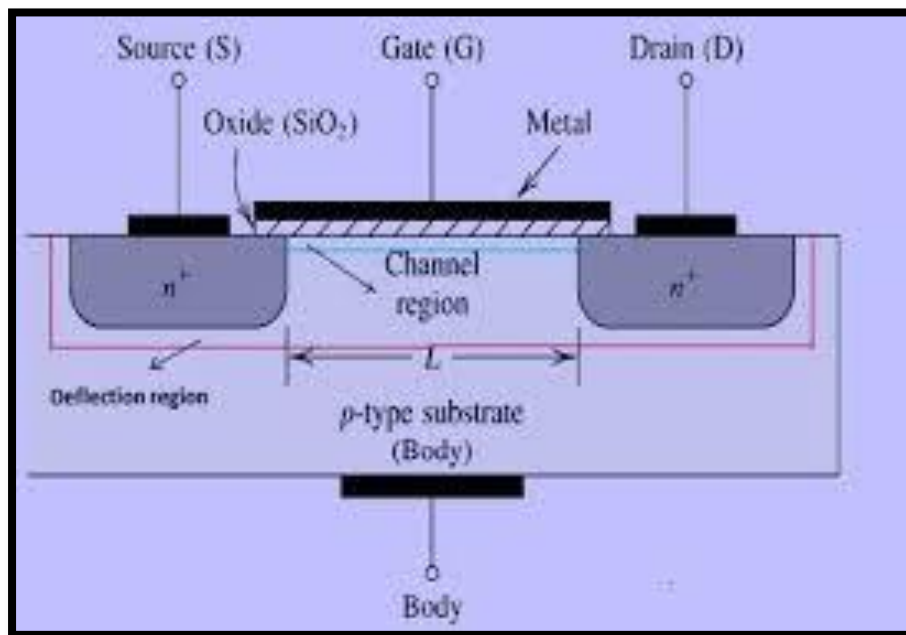
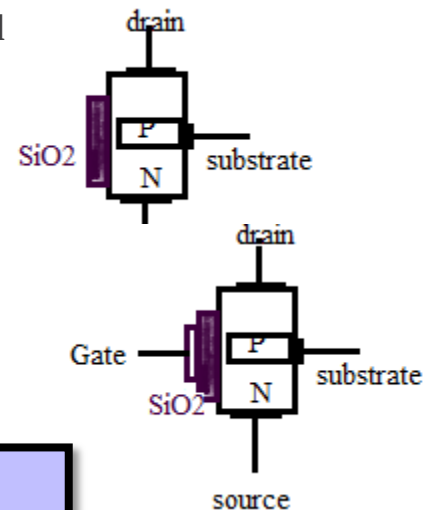
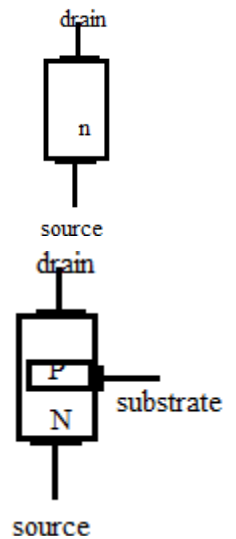
a/ IGFET region

To begin with, there is an N-region with a source and drain as in fig.a, as (+ve) voltage applied to the drain-source terminals forces free electrons to flow from the source to the drain.

The MOSFET has a region called the substrate. This P-region reduces the channel between the source and the drain so that only a small passage remains at the left side in fig.b. Free electrons flowing from the source to the drain must pass through thin narrow channel (lightly doped).

A thin layer of silicon dioxide (SiO_2), an insulator, is deposited over the left side, fig.c, of the channel.

Finally, a metallic gate is deposited on the insulator; Fig. d. because of the gate is insulated from the channel a MOSFET is also known as IGFET.



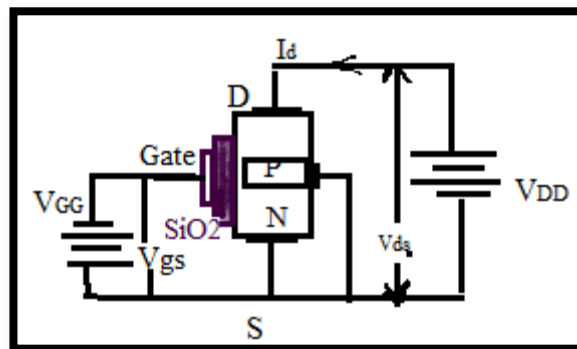
The depletion type MOSFET can be operated in two modes:

The depletion mode and the enhancement mode.

1- Depletion Mode:

This happens when negative voltage is applied on the gate .the operation here is similar to that of the JFET.

When V_{GS} is negative voltage, it depletes the N-channel of its electrons by inducing (+Ve) charge in it as shown in fig. a. The greater the negative voltage on the gate, the greater is the reduction in the number of electrons in the channel, and, consequently, the smaller it's the conductivity of the channel. So the more (-Ve) gate voltage is, the smaller the current through the channel. Enough negative voltage on the gate ($V_{GS(off)}$) cuts off the current between the source and the drain. Therefore, with (-Ve) gate voltage the action of a MOSFET is similar to that of a JFET.



2) Enhancement mode

In this mode the gate is made positive with respect to the source. The input gate capacitor is able to create free electrons in the channel which increasing I_d . this positive voltage will attracts the electrons into the channel opposite to the gate thus increasing the conductivity of the channel. A increasing the positive voltage

on the gate, increases the conductivity of the channel thus increasing the drain current

MOSFET Characteristics

This fig. shows typical drain curves for an N-channel D-MOSFET. For V_{GS} less than zero, we get depletion-mode operation, on the other hand, V_{GS} greater than zero given enhancement-mode of operation.

$V_{GS} < 0$ depletion

$V_{GS} > 0$ enhancement

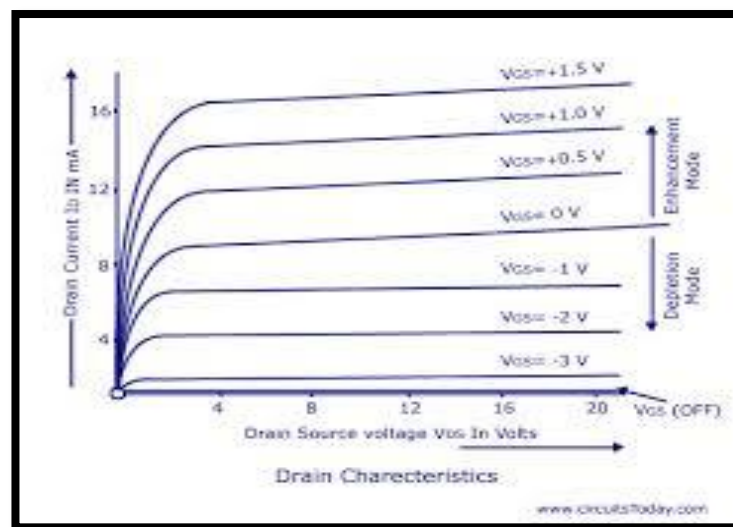
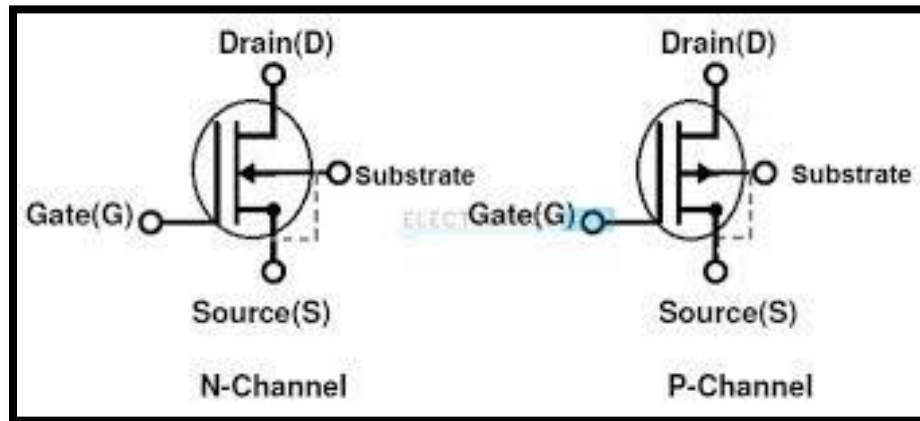


Fig.a Output characteristics

The transconductance curve, fig.b, I_{DDs} is no longer is the maximum possible current. As you can see, the transconductance curve extends to the right, so that I_d is greater than I_{DDs} in the enhancement mode. MOSFETs with a transconductance curve just like this fig. are easy to use since it does not require bias voltage. Q-value can be indicate by the vertical intercept where $I_d = I_{DDs} + V_{GS} = 0$. This means we do not have to provide any gate voltage at all, which simplifies biasing circuit.

Schematic Symbol



Biasing

D-MOSFET

1. The zero bias circuit.

$$V_{GS} = 0$$

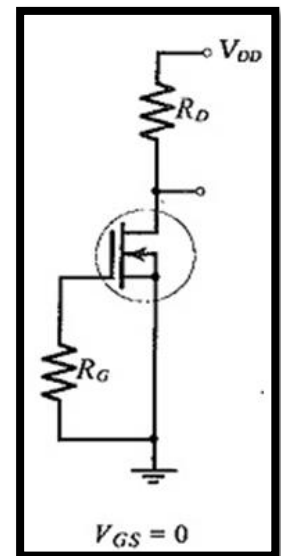
Being able to use zero V_{GS} is an advantage when it comes to biasing. It permits the unique biasing circuit in fig. b.

this simple circuit has no applied gate or source voltage.

Therefore, $V_{GS} = 0$, $I_d = I_{DSS}$ the drain voltage:

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

The zero bias of fig.a is unique with D-type MOSFET.



2. The voltage divider bias circuit

This circuit can be operated in the enhancement mode or the depletion mode according to the relation between the gate and the source voltage.

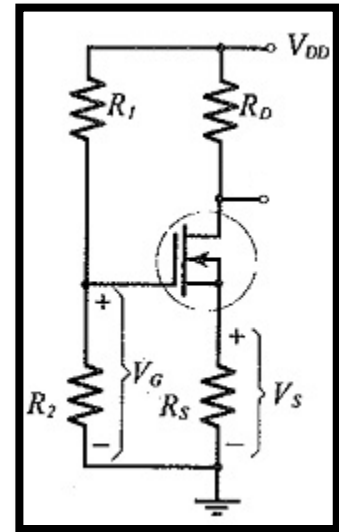
$$V_{GS} = V_G - V_S$$

$V_G > V_S$ enhancement mode

$V_S > V_G$ depletion mode.

Enhancement mode

$$V_{GS} = + V_G$$



2- Enhancement Type MOSFET (normally OFF mode)

E-MOSFET is a four-terminal device: source(S), gate (G), drain (D) and substrate (Sub.) terminals. The substrate of the MOSFET is in connection with the source terminal thus forming a three-terminal device.

Here, for N channel E-MOSFET V_G must be positive.

V_{GS} must be $\gg V_{GS(th)}$ for current flow.

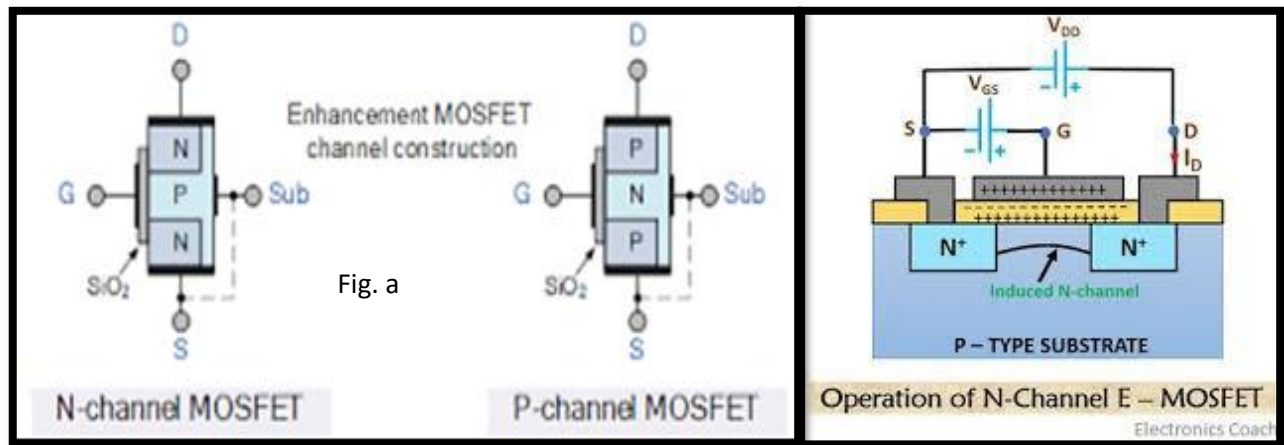
Transconductance equation is :

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

A/ Construction

Fig. a shows the construction of E-MOSFET. It shows the normal biasing polarities. As seen the substrate extends all the way to the insulating layer, there is no N-channel, no connection between the drain and the source. When gate voltage is zero, the V_{DD} supply tries to force free electrons to flow from the source to drain,

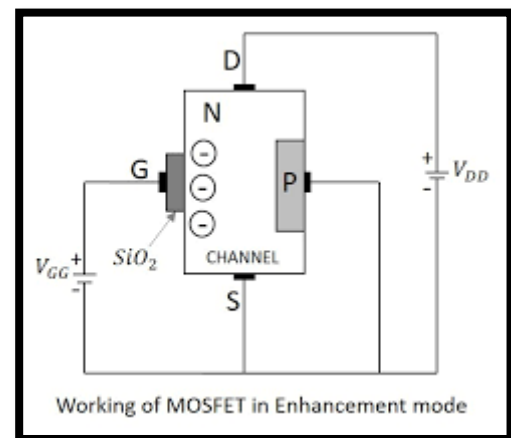
but the P substrate has only a few thermally produced free electrons. As a result, the drain current is almost zero. **The E- MOSFET is a normally OFF.**



B/ Operation

When V_{GS} is positive ($V_{GS} > 0$), the gate attracts free electrons into the P-substrate region. These free electrons recombine with the holes in the p-substrate near the silicon dioxide layer.

When V_{GS} is large enough, all these holes near the silicon dioxide layer are with electrons and the free electrons then acts as a channel between

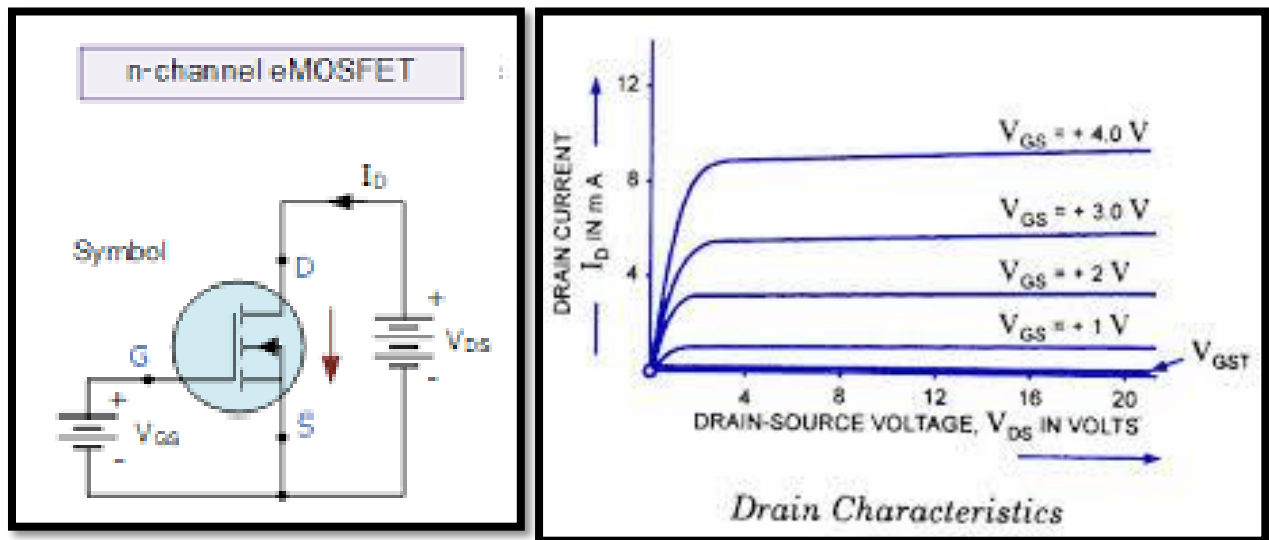


the drain and the source electrons start flowing from source and drain. This layer will become N-type because of the induced free electrons. The effect is as though an N-type layer between the source and drain has been created (a channel has been induced). **This N-type layer is known as n-type inversion layer or induced channel.** This E-MOSFET is called the N induced channel E-MOSFET. When this layer is formed, the terminal is turned on and free electrons can easily flow from the source

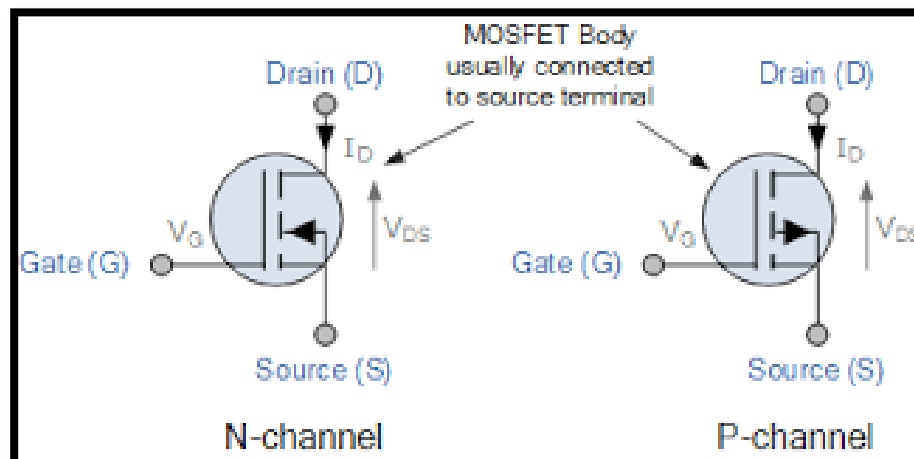
to drain. The minimum value of gate to source voltage V_{GS} which can create inversion layer is known as **gate threshold gate voltage** $V_{GS(th)}$.

Thus, the transistor is off when $V_{GS} < V_{GS(th)}$ and is on, when $V_{GS} > V_{GS(th)}$. Thus the conducting capability of enhancement type MOSFET depends on the action of N-type inversion layer.

The same can be said about the E-MOSFET where the drain and the source are of P-type semiconductor, the substrate will be of N-type semiconductor but of course the channel will a P-induced layer. **What must the polarity of the voltage on the gate be?**



C/Symbol



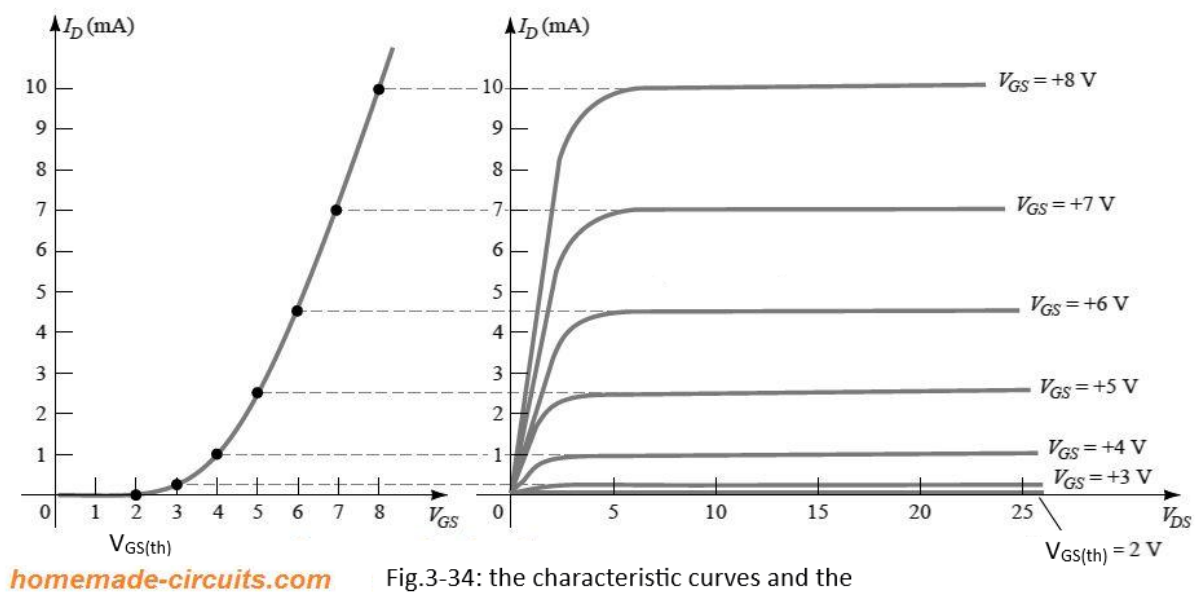
D/ Characteristics

For $V_{GS} < V_{GS(th)}$ $I_D = 0$.

The characteristic curves start for $V_{GS} > V_{GS(th)}$. As V_{GS} increases the drain current increases. For a certain value of V_{GS} , three regions are noticed: the ohmic region, the saturation region and the breakdown region (not drawn in the figure, which happens when the values of V_{DS} and V_{GS} exceeds the values that the MOSFET can stand. These voltages values will result in the destruction of the insulating layer).

From these curves, the transconductance curve can be drawn, which shows the relation between I_d and V_{GS} . The equation of this curve is:

$$I_d = k (V_{GS} - V_{GS(th)})^2 \text{ where } K \text{ is a constant}$$



homemade-circuits.com

Fig.3-34: the characteristic curves and the transconductance curve of the common drain E-MOSFET

E/ Biasing

1) Drain feedback bias, Fig. a, a type of bias that you can use only with enhancement type MOSFETs. When the MOSFET is conducting, it has a drain current $I_{d(on)}$ and a drain voltage of $V_{DS(on)}$. Since the gate current is approximately zero.

$$V_G = V_D$$

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

$$V_{DD} - V_G = I_D R_D$$

2- Voltage divider bias circuit

Fig,b

$$V_G = V_{R2}$$

