

UNIT I Overview of Electronic Components & Signals:

Active Components:

Active components are electronic component that deliver or supply power or produce energy in form of voltage or current . These components depends on the external power source to control or modify electrical signals . Active components are able to control and amplify flow of current in electric circuit . These components do not store power or energy , they provide energy to electric circuit . They are also responsible to perform function like Amplification , Rectification , Switching or we can say that they influence circuit actively .

Ex- Voltage Source, Current Source, Diode, BJT, Charges Capacitor, Charged Inductor

Passive Components

Passive components are electronic component that absorb or store the energy in circuit. These components do not depend on external power source for their operation . They are not able to control the flow of current rather it influence the current in circuit . As the name suggest "Passive" they perform operation like storing , releasing energy . They can only receive energy and can either store , absorb or release it in one or other form of energy. They cannot provide gain or amplification to signal rather they decrease the energy or strength of signal without distorting it .

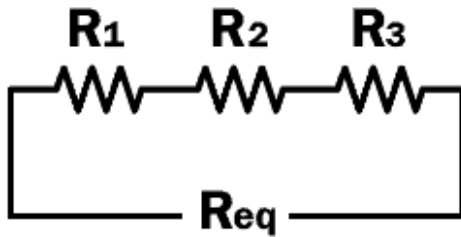
Ex-Resistor, Capacitor, Inductor

Elements Symbol	RESISTOR 	CAPACITOR 	INDUCTOR 
Denoted by	R	C	L
Equation	$R = \frac{V}{I}$	$C = \frac{Q}{V}$	$L = \frac{V_L}{(di/dt)}$
Series	$R_T = R_1 + R_2$	$\frac{1}{C_T} = \frac{1}{C_1} + \frac{1}{C_2}$	$L_T = L_1 + L_2$
Parallel	$\frac{1}{R_T} = \frac{1}{R_1} + \frac{1}{R_2}$	$C_T = C_1 + C_2$	$\frac{1}{L_T} = \frac{1}{L_1} + \frac{1}{L_2}$

Resistance in Series & Parallel Equations

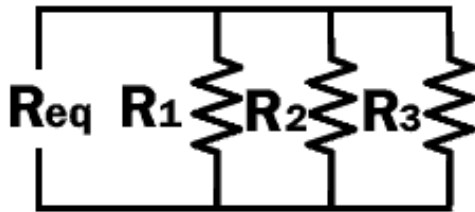
Resistance:

The total equivalent resistance of resistors connected in series or parallel configuration is given the following formulas:

Resistance In Series:

When two or more than two resistors are connected in series as shown in figure their equivalent resistance is calculated by:

$$R_{Eq} = R_1 + R_2 + R_3 + \dots R_n$$

Resistance In Parallel:

when the resistors are in parallel configuration the equivalent resistance becomes:

$$\left(\frac{1}{R_{Eq}} \right) = \left(\frac{1}{R_1} \right) + \left(\frac{1}{R_2} \right) + \left(\frac{1}{R_3} \right) + \dots$$

Where

- R_{Eq} is the equivalent resistance of all resistors ($R_1, R_2, R_3 \dots R_n$)

P-N Junction

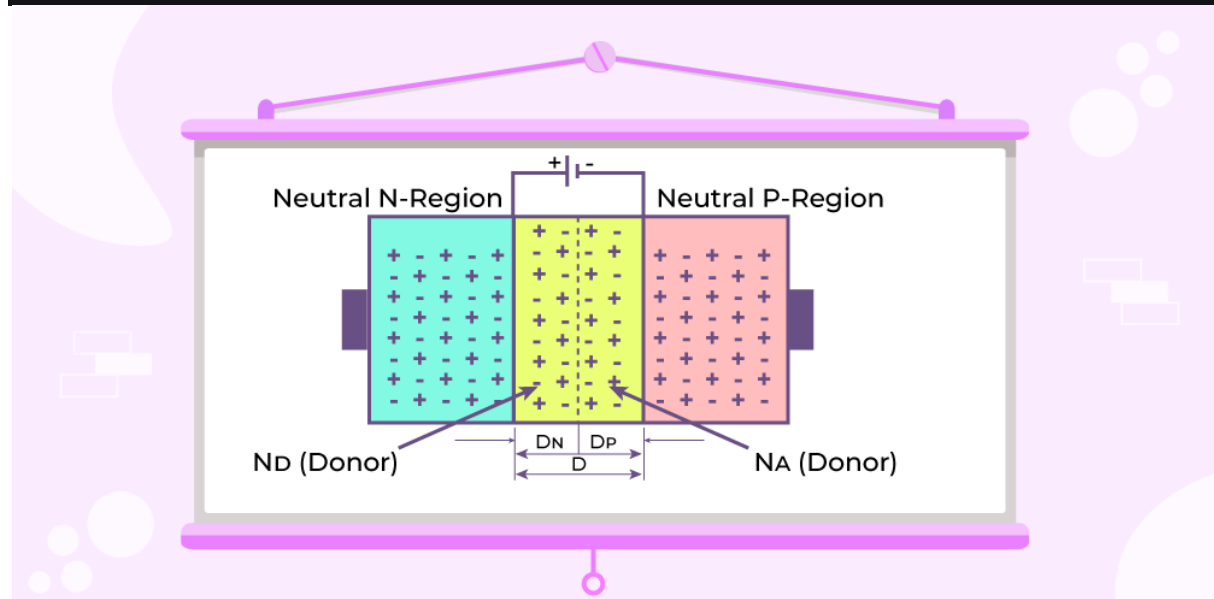
The semiconductor's p-side, or positive side, has an excess of holes, whereas the n-side, or negative side, has an excess of electrons. The doping process is used to produce the p-n junction in a semiconductor.

Formation of P-N Junction

When we utilize various semiconductor materials to form a p-n junction, there will be a grain boundary that will prevent electrons from moving from one side to the other by scattering electrons and holes, which is why we employ the doping procedure.

For example, Consider a p-type silicon semiconductor sheet that is very thin. A portion of the p-type Si will be changed to n-type silicon if a tiny quantity of pentavalent impurity is added. This sheet will now have both a p-type and an n-type area, as well as a junction between the two. Diffusion and drift are the two sorts of processes that occur following the creation of a p-n junction. As we all know, the concentration of holes and electrons on the two sides of a junction

differs, with holes from the p-side diffusing to the n-side and electrons from the n-side diffusing to the p-side. This causes a diffusion current to flow across the connection.



When an electron diffuses from the n-side to the p-side, it leaves an ionized donor on the n-side, which is stationary. On the n-side of the junction, a layer of positive charge develops as the process progresses. When a hole is moved from the p-side to the n-side, an ionized acceptor is left behind on the p-side, causing a layer of negative charges to develop on the p-side of the junction. The depletion area is defined as a region of positive and negative charge on each side of the junction. An electric field direction from a positive charge to a negative charge is generated due to this positive space charge area on each side of the junction. An electron on the p-side of the junction travels to the n-side of the junction due to the electric field. The drift is the name given to this motion. We can observe that the drift current runs in the opposite direction as the diffusion current.

Biasing Conditions for p-n Junction Diode

In a p-n junction diode, there are two operational regions:

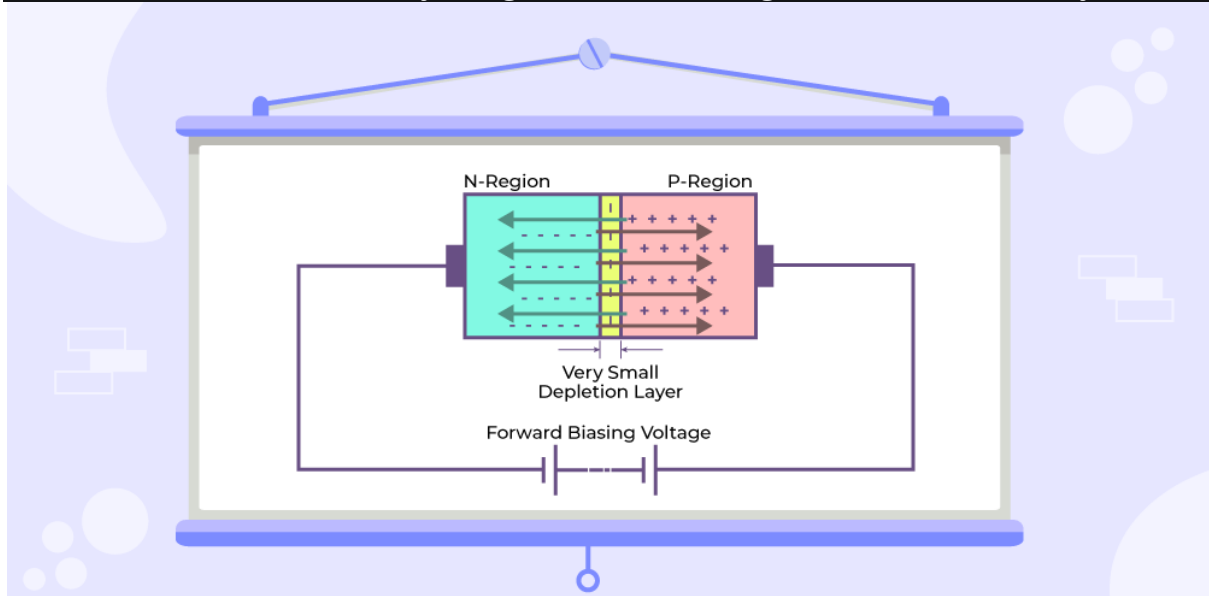
- p-type
- n-type

The voltage applied determines one of three biasing conditions for p-n junction diodes:

- There is no external voltage provided to the p-n junction diode while it is at **zero bias**.
- **Forward bias:** The p-type is linked to the positive terminal of the voltage potential, while the n-type is connected to the negative terminal.
- **Reverse bias:** The p-type is linked to the negative terminal of the voltage potential, while the n-type is connected to the positive terminal.

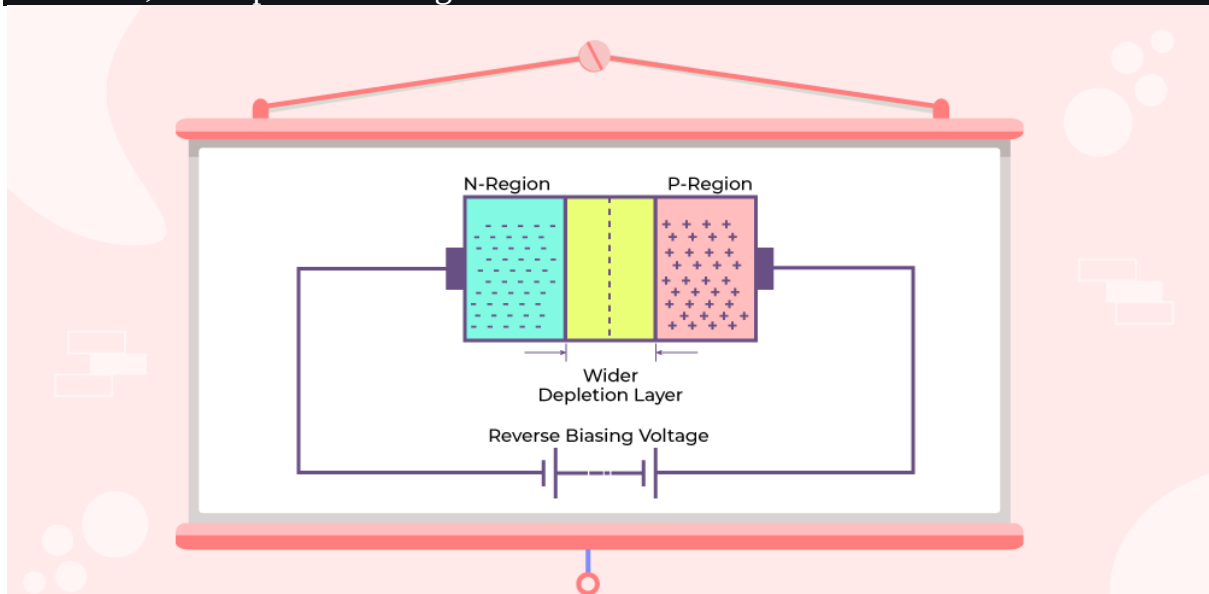
Forward Biased PN Junction

The resulting electric field is smaller than the built-in electric field when both electric fields are added together. As a result, the depletion area becomes less resistant and thinner. When the applied voltage is high, the resistance of the depletion zone becomes insignificant. At 0.6 V, the resistance of the depletion area in silicon becomes absolutely insignificant, allowing current to flow freely over it.



Reverse Biased PN Junction

The resultant electric field is in the same direction as the built-in electric field, resulting in a more resistive, thicker depletion zone. If the applied voltage is increased, the depletion area gets more resistant and thicker.



P-N Junction Formula

The p-n junction formula, which is based on the built-in potential difference generated by the electric field, is as follows:

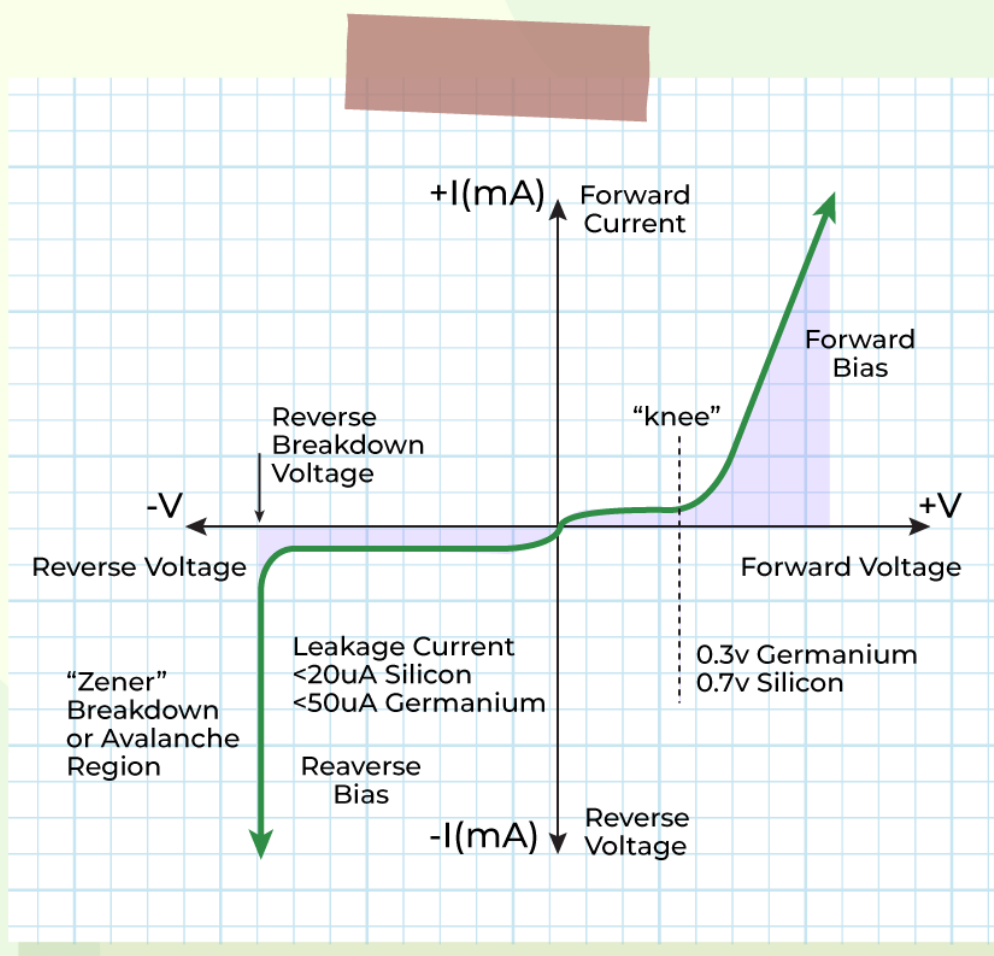
$$E_o = V_T \ln \left[\frac{N_D - N_A}{n_i^2} \right]$$

Current Flow in PN Junction diode

When the voltage is increased, electrons move from the n-side to the p-side of the junction. The migration of holes from the p-side to the n-side of the junction occurs in a similar manner as the voltage rises. As a result, a concentration gradient exists between the terminals on both sides.

There will be a movement of charge carriers from higher concentration regions to lower concentration regions as a result of the development of the concentration gradient. The current flow in the circuit is caused by the movement of charge carriers inside the p-n junction.

VI Characteristics of PN Junction Diode



A curve between the voltage and current across the circuit defines the V-I properties of p-n junction diodes. The x-axis represents voltage, while the y-axis

represents current. The V-I characteristics curve of the p-n junction diode is shown in the graph above. With the help of the curve, we can see that the diode works in three different areas, which are:

- Zero bias
- Forward bias
- Reverse bias

Zero Bias

There is no external voltage provided to the p-n junction diode while it is at zero bias, which implies the potential barrier at the junction prevents current passage.

Forward bias

When the p-n junction diode is in forwarding bias, the p-type is linked to the positive terminal of the external voltage, while the n-type is connected to the negative terminal. The potential barrier is reduced when the diode is placed in this fashion. When the voltage is 0.7 V for silicon diodes and 0.3 V for germanium diodes, the potential barriers fall, and current flows.

The current grows slowly while the diode is under forwarding bias, and the curve formed is non-linear as the voltage supplied to the diode overcomes the potential barrier. Once the diode has crossed the potential barrier, it functions normally, and the curve rises quickly as the external voltage rises, yielding a linear curve.

Reverse Bias

When the PN junction diode is under negative bias, the p-type is linked to the negative terminal of the external voltage, while the n-type is connected to the positive terminal. As a result, the potential barrier becomes higher. Because minority carriers are present at the junction, a reverse saturation current occurs at first.

When the applied voltage is raised, the kinetic energy of the minority charges increases, affecting the majority charges. This is the point at which the diode fails. The diode may be destroyed as a result of this.

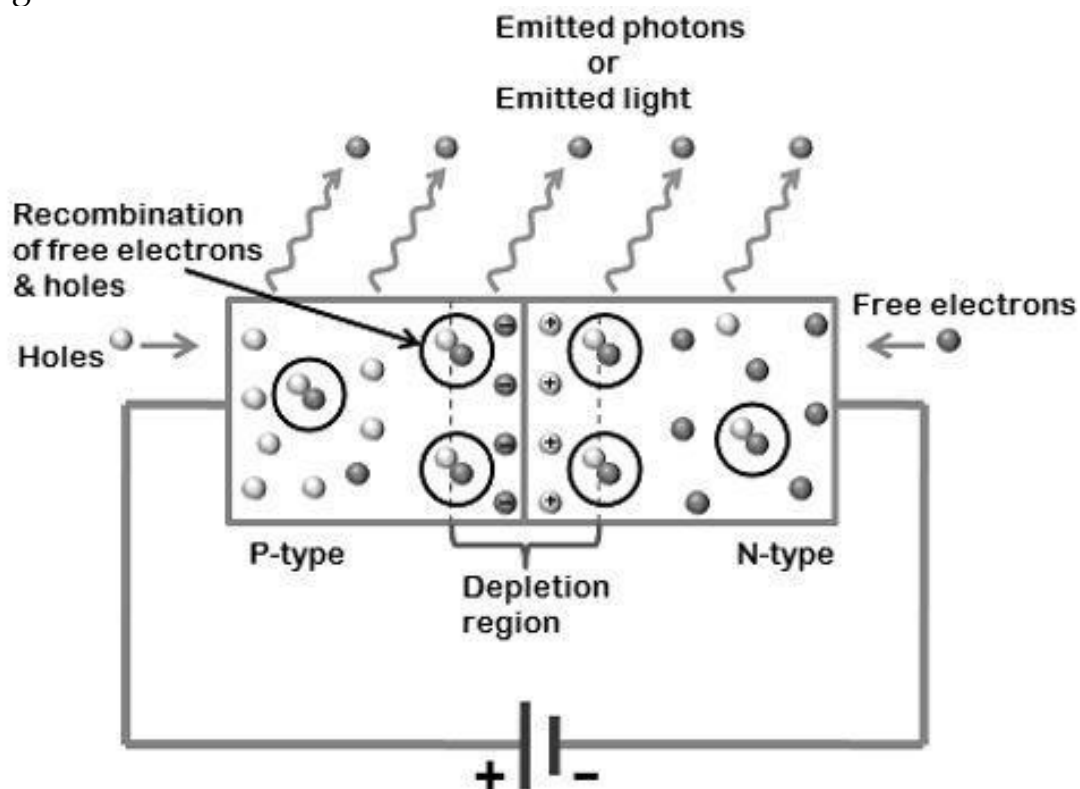
Applications of PN Junction Diode

There are various applications of PN Junction Diode in the field of electronics, some of those applications are listed as follows:

- A most common use case of a PN junction diode is as a rectifier which means converting AC current into DC current.
- Zener diode (which is a special type of PN junction diode) is used in circuits for voltage regulation.
- As Diode only conducts current in Forward bias, so in electrical circuits, it is used as a switch to turn on and off certain small circuits in a much more complex circuit.
- A reverse-biased p-n junction diode is utilized as a photodiode as it is sensitive to light.
- LED is also a special type of PN junction diode on a forward basis which emits light.

Working of LED

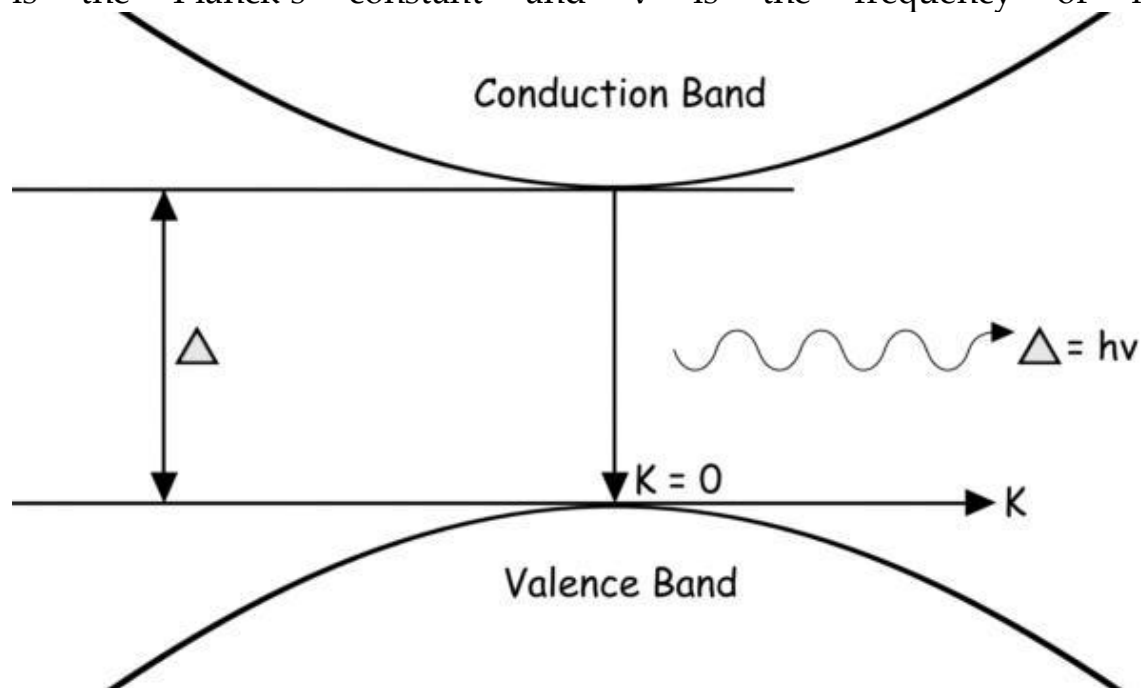
Like an ordinary diode, the LED diode works when it is forward biased. In this case, the n-type semiconductor is heavily doped than the p-type forming the p-n junction. When it is forward biased, the potential barrier gets reduced and the electrons and holes combine at the depletion layer (or active layer), light or photons are emitted or radiated in all directions. A typical figure blow showing light emission due electron-hole pair combining on forward biasing.



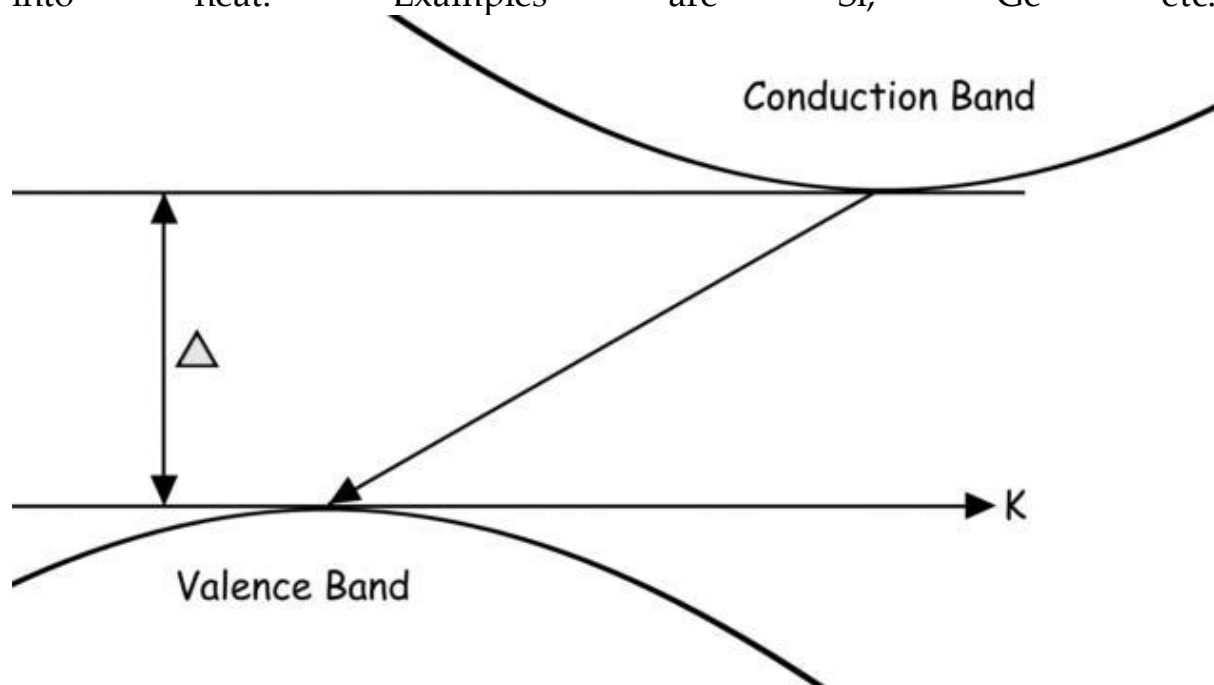
The emission of photons in an LED is explained by the energy band theory of solids, which dictates that light emission depends on the material's band gap being direct or indirect. Those semiconductor materials which have a direct band gap are the ones that emit photons. In a direct bandgap material, the bottom of the energy level of conduction band lies directly above the topmost energy level of the valence band on the Energy vs Momentum (wave vector ' k ') diagram.

When electrons and hole recombine, energy $E = h\nu$ corresponding to the energy gap Δ (eV) is escaped in the form of light energy or photons where h

is the Planck's constant and ν is the frequency of light.



Direct Band Gap
 Indirect band gap materials are non-radiative, as their conduction band's bottom does not align with the valence band's top, converting most energy into heat. Examples are Si, Ge etc.



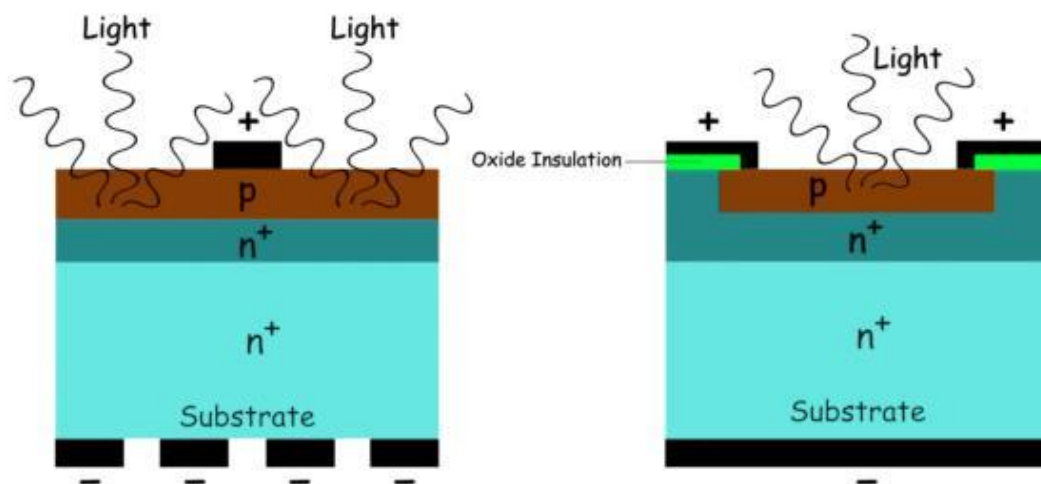
Indirect Band Gap
 Example of material which has direct band gap is Gallium Arsenide (GaAs), a compound semiconductor which is the material used in LEDs. Dopant

atoms are added to GaAs to give out a wide range of colors. Some of the materials used in LEDs are:

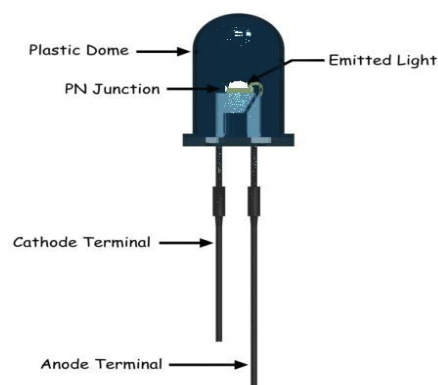
- Aluminium Gallium Arsenide (AlGaAs) – infrared.
- Gallium Arsenic Phosphide (GaAsP) – red, orange, yellow.
- Aluminium Gallium Phosphide (AlGaP) – green.
- Indium gallium nitride (InGaN) – blue, blue-green, near UV.
- Zinc Selenide (ZnSe) – blue.

Physical Structure of LED

LED is structured in such a way so that light emitted does not get reabsorbed into the material. So it is ensured that the electron-hole recombination takes place on the surface.



The above figure shows the two different ways of structuring LED p-n junction. The p-type layer is made thin and is grown on the n-type substrate. Metal electrodes attached on either side of the p-n junction serve as nodes for external electrical connection. The Light emitting diode p-n junction is encased in a dome-shaped transparent case so that light is emitted uniformly



in all directions and to takes place.

minimum internal reflection

What Are Zener Diodes

The **Zener Diode** or “Breakdown Diode”, as they are sometimes referred too, are basically the same as the standard **PN-junction diode** except that they are specially designed to have a low and specified **Reverse Breakdown Voltage** which takes advantage of any reverse voltage applied to it.

In the forward-biased direction, that is Anode is more positive with respect to its Cathode, a **zener diode** behaves like a normal junction diode when the forward voltage V_F across the diode exceeds 0.7 volts (silicon) causing the zener diode to conduct.

The forward current flowing through the conducting diode is at its maximum determined only by the connected load. Thus in the forward-bias direction, the zener behaves like a regular diode within its specified current and/or power limits and as such, the forward characteristics of a zener diode is generally of no interest.

However, unlike a conventional diode that blocks any flow of current through itself when reverse biased, that is the Cathode becomes more positive than the Anode, as soon as the reverse voltage reaches a pre-determined value, the zener diode begins to conduct in the reverse direction.

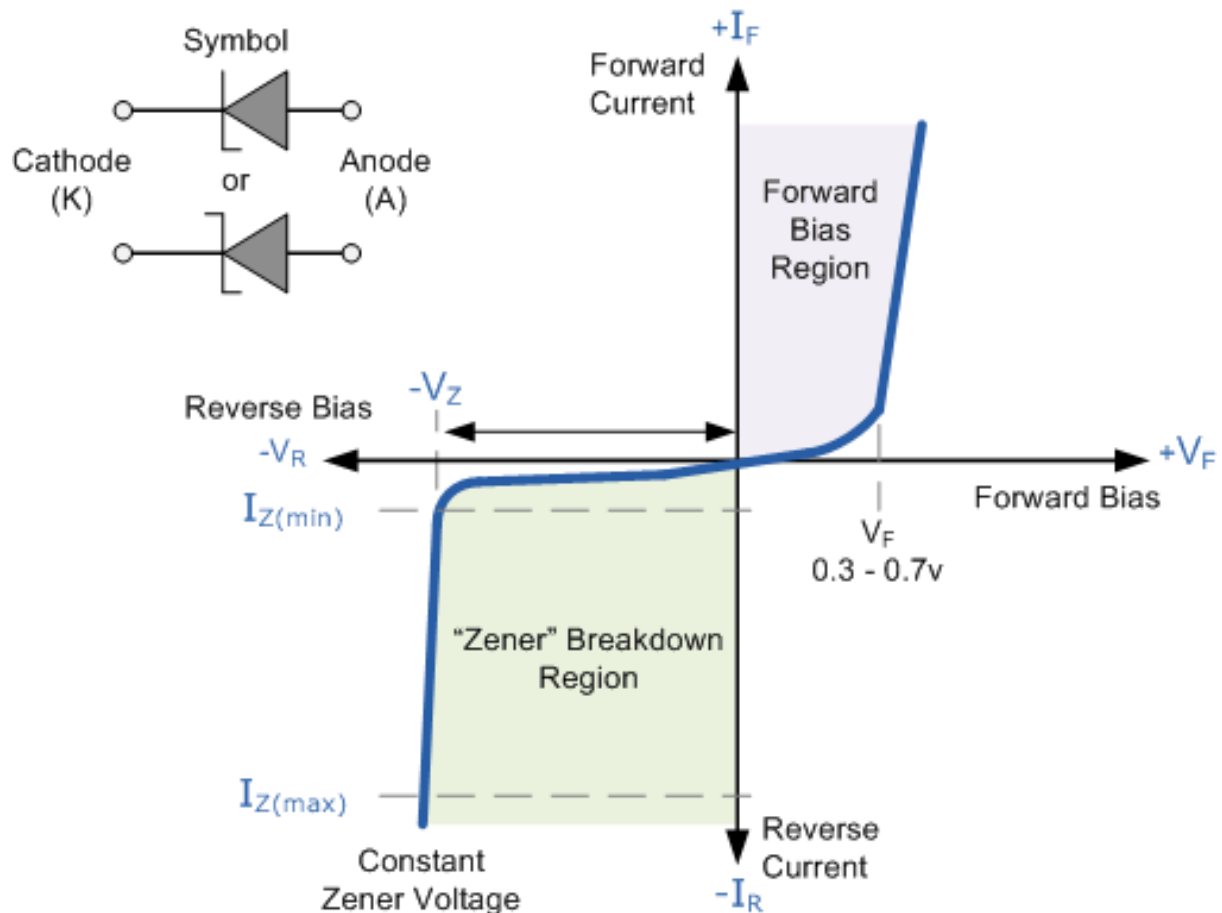
Since a zener diode is designed to work in the reverse breakdown region of its characteristic curve, they have a fixed breakdown voltage, V_Z value which is determined during manufacture. As the reverse voltage across the zener diode increases from 0 volts to its zener breakdown voltage, a small reverse or leakage current will flow through the diode which remains fairly constant as the reverse voltage increases.

Once the reverse voltage applied across the zener diode exceeds the rated voltage of the device, a process called *Zener Breakdown* occurs in the semiconductor depletion layer and a current starts to flow through the diode to limit this increase in voltage.

The current now flowing through the zener diode increases dramatically to its maximum circuit value (which is usually limited by a series resistor). Once zener breakdown occurs, the voltage drop across the diode remains fairly constant even though the current, I_Z through it can vary considerably. The voltage point at which the voltage across the diode becomes stable is called the “zener voltage”, (V_Z). For zener diodes this breakdown voltage value can range from a few volts up to a few hundred volts.

The point at which the zener voltage triggers the current to flow through the diode can be very accurately controlled (to less than 1% tolerance) in the doping stage of the diodes semiconductor construction giving the diode a specific *zener breakdown voltage*, (V_Z) for example, 4.3V or 7.5V. This zener breakdown voltage on the I-V curve is almost a vertical straight line.

Zener Diode I-V Characteristics



The **Zener Diode** is used in its “reverse bias” or reverse breakdown mode, i.e. the diode's anode connects to the negative supply. From the I-V characteristics curve above, we can see that the zener diode has a region in its reverse bias characteristics of almost a constant negative voltage regardless of the value of the current flowing through the diode.

This voltage remains almost constant even with large changes in current providing the zener diode's current remains between the breakdown current $I_{Z(min)}$ and its maximum current rating $I_{Z(max)}$.

This ability of the zener diode to control itself can be used to great effect to regulate or stabilise a voltage source against supply or load variations. The fact that the voltage across the diode in the breakdown region is almost constant turns out to be an important characteristic and one which can be used in the simplest types of voltage regulator applications.

The function of a voltage regulator is to provide a constant output voltage to a load connected in parallel with it in spite of the ripples in the supply voltage or variations in the load current. A zener diode will continue to regulate its voltage until the diode's holding current falls below the minimum $I_{Z(min)}$ value in the reverse breakdown region.

The Zener Diode Regulator

Zener Diodes can be used to produce a stabilised voltage output with low ripple under varying load current conditions. By passing a small current through the diode from a voltage source, via a suitable current limiting resistor (R_s), the zener diode will conduct sufficient current to maintain a voltage drop of V_{out} .

We remember from the previous tutorials that the DC output voltage from the half or full-wave rectifiers contains ripple superimposed onto the DC voltage and that as the load value changes so to does the average output voltage. By connecting a simple zener based stabiliser circuit as shown below across the output of the rectifier, a more stable output voltage can be produced.

What Are Bipolar Junction Transistors?

Unlike semiconductor diodes which are made up from two pieces of semiconductor material to form one simple pn-junction. The *bipolar transistor* uses one more layer of semiconductor material to produce a device with properties and characteristics of an amplifier.

If we join together two individual signal diodes back-to-back, this will give us two PN-junctions connected together in series which would share a common *Positive*, (P) or *Negative*, (N) terminal. The fusion of these two diodes produces a three layer, two junction, three terminal device forming the basis of a **Bipolar Junction Transistor**, or **BJT** for short.

Transistors are three terminal active devices made from different semiconductor materials that can act as either an insulator or a conductor by the application of a small signal voltage. The transistor's ability to change between these two states enables it to have two basic functions: "switching" (digital electronics) or "amplification" (analogue electronics). Then bipolar transistors have the ability to operate within three different regions:

- **Active Region** – the transistor operates as an amplifier and $I_c = \beta \cdot I_b$
- **Saturation** – the transistor is "Fully-ON" operating as a switch and $I_c = I(\text{saturation})$
- **Cut-off** – the transistor is "Fully-OFF" operating as a switch and $I_c = 0$



A

Typical

Bipolar Transistor

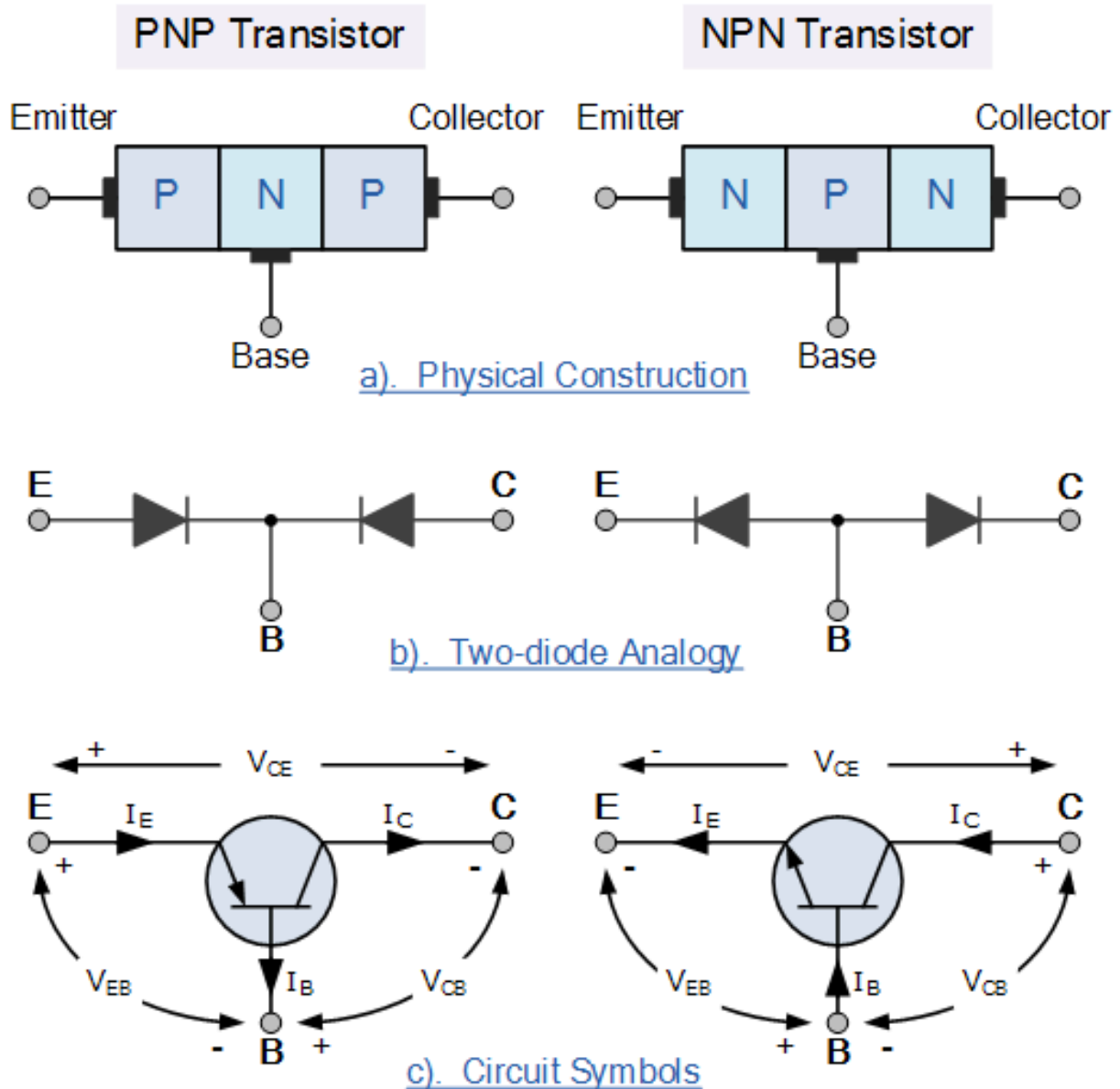
The word Transistor is a combination of the two words Transfer Varistor which describes their mode of operation way back in their early days of electronics development. There are two basic types of bipolar transistor construction, PNP and NPN, which basically describes the physical arrangement of the P-type and N-type semiconductor materials from which they are made.

The **Bipolar Transistor** basic construction consists of two PN-junctions producing three connecting terminals with each terminal being given a name to identify it from the other two. These three terminals are known and labelled as the Emitter (E), the Base (B) and the Collector (C) respectively.

Bipolar Transistors are current regulating devices that control the amount of current flowing through them from the Emitter to the Collector terminals in proportion to the amount of biasing voltage applied to their base terminal, thus acting like a current-controlled switch. As a small current flowing into the base terminal controls a much larger collector current forming the basis of transistor action.

The principle of operation of the two transistor types PNP and NPN, is exactly the same the only difference being in their biasing and the polarity of the power supply for each type.

Bipolar Transistor Construction



The construction and circuit symbols for both the PNP and NPN bipolar transistor are given above with the arrow in the circuit symbol always showing the direction of “conventional current flow” between the base terminal and its emitter terminal. The direction of the arrow always points from the positive P-type region to the negative N-type region for both transistor types, exactly the same as for the standard diode symbol.

Bipolar Transistor Configurations

As the **Bipolar Transistor** is a three terminal device, there are basically three possible ways to connect it within an electronic circuit with one terminal being common to both the input and output signals. Each method of connection responding differently to its input signal within a circuit as the static characteristics of the transistor vary with each circuit arrangement.

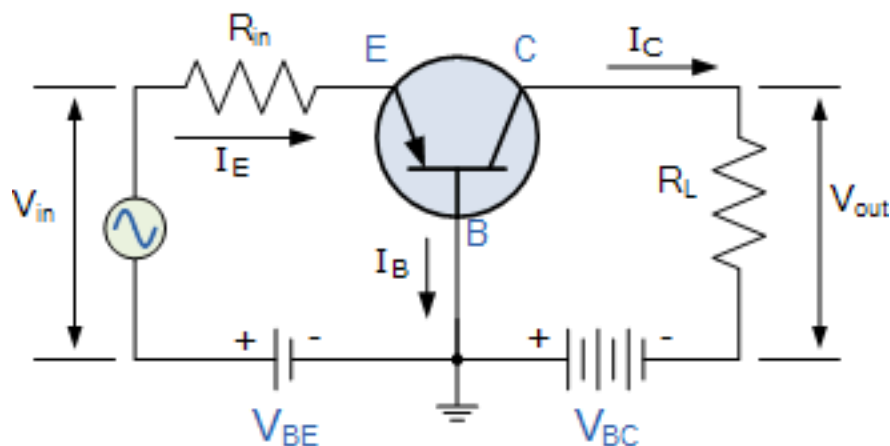
- Common Base Configuration – has Voltage Gain but no Current Gain.
- Common Emitter Configuration – has both Current and Voltage Gain.
- Common Collector Configuration – has Current Gain but no Voltage Gain.

The Common Base (CB) Configuration

As its name suggests, in the **Common Base** or grounded base configuration, the BASE connection is common to both the input signal AND the output signal. The input signal is applied between the transistors base and the emitter terminals, while the corresponding output signal is taken from between the base and the collector terminals as shown. The base terminal is grounded or can be connected to some fixed reference voltage point.

The input current flowing into the emitter is quite large as its the sum of both the base current and collector current respectively therefore, the collector current output is less than the emitter current input resulting in a current gain for this type of circuit of “1” (unity) or less, in other words the common base configuration “attenuates” the input signal.

The Common Base Transistor Circuit



This type of amplifier configuration is a non-inverting voltage amplifier circuit, in that the signal voltages V_{in} and V_{out} are “in-phase”. This type of transistor arrangement is not very common due to its unusually high voltage gain characteristics. Its input characteristics represent that of a forward biased diode while the output characteristics represent that of an illuminated photo-diode.

Also this type of bipolar transistor configuration has a high ratio of output to input resistance or more importantly “load” resistance (R_L) to “input” resistance (R_{in})

giving it a value of “Resistance Gain”. Then the voltage gain (A_v) for a common base configuration is therefore given as:

Common Base Voltage Gain

$$A_v = \frac{V_{out}}{V_{in}} = \frac{I_C \times R_L}{I_E \times R_{IN}}$$

Where: I_C/I_E is the current gain, alpha (α) and R_L/R_{in} is the resistance gain.

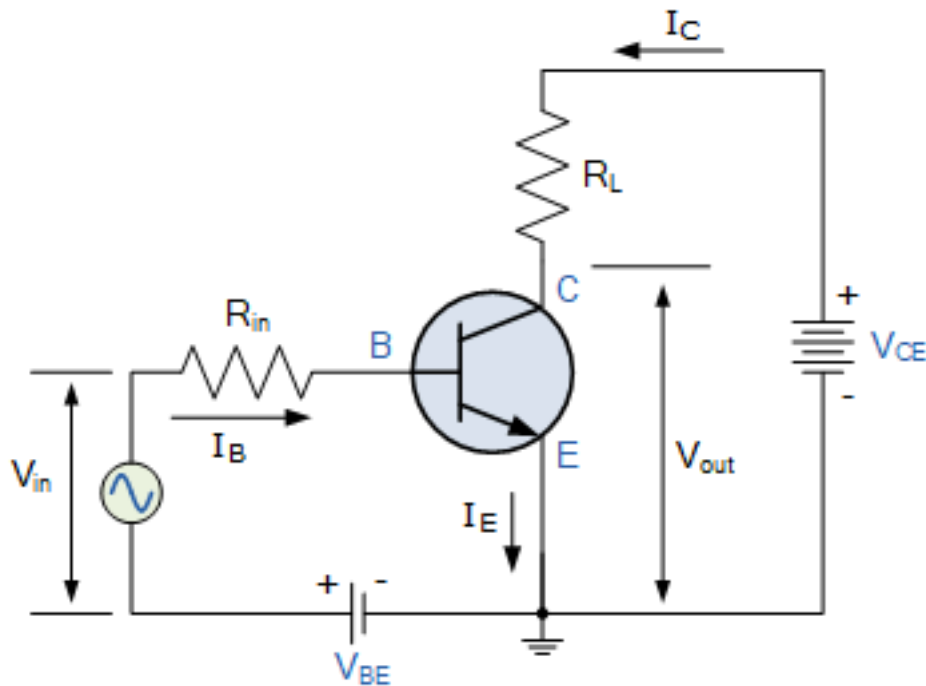
The common base circuit is generally only used in single stage amplifier circuits such as microphone pre-amplifier or radio frequency (R_f) amplifiers due to its very good high frequency response.

The Common Emitter (CE) Configuration

In the **Common Emitter** or grounded emitter configuration, the input signal is applied between the base and the emitter, while the output is taken from between the collector and the emitter as shown. This type of configuration is the most commonly used circuit for transistor based amplifiers and which represents the “normal” method of bipolar transistor connection.

The common emitter amplifier configuration produces the highest current and power gain of all the three bipolar transistor configurations. This is mainly because the input impedance is LOW as it is connected to a forward biased PN-junction, while the output impedance is HIGH as it is taken from a reverse biased PN-junction.

The Common Emitter Amplifier Circuit



In this type of configuration, the current flowing out of the transistor must be equal to the currents flowing into the transistor as the emitter current is given as $I_E = I_C + I_B$.

As the load resistance (R_L) is connected in series with the collector, the current gain of the common emitter transistor configuration is quite large as it is the ratio of I_C/I_B . A transistor's current gain is given the Greek symbol of Beta, (β).

As the emitter current for a common emitter configuration is defined as $I_E = I_C + I_B$, the ratio of I_C/I_E is called Alpha, given the Greek symbol of α . Note: that the value of Alpha will always be less than unity.

Since the electrical relationship between these three currents, I_B , I_C and I_E is determined by the physical construction of the transistor itself, any small change in the base current (I_B), will result in a much larger change in the collector current (I_C).

Then, small changes in current flowing in the base will thus control the current in the emitter-collector circuit. Typically, Beta has a value between 20 and 200 for most general purpose transistors. So if a transistor has a Beta value of say 100, then one electron will flow from the base terminal for every 100 electrons flowing between the emitter-collector terminal.

By combining the expressions for both Alpha, α and Beta, β the mathematical relationship between these parameters and therefore the current gain of the transistor can be given as:

$$\text{Alpha, } (\alpha) = \frac{I_C}{I_E} \quad \text{and} \quad \text{Beta, } (\beta) = \frac{I_C}{I_B}$$

$$\therefore I_C = \alpha \cdot I_E = \beta \cdot I_B$$

$$\text{as: } \alpha = \frac{\beta}{\beta + 1} \quad \beta = \frac{\alpha}{1 - \alpha}$$

$$I_E = I_C + I_B$$

Where: “I_c” is the current flowing into the collector terminal, “I_b” is the current flowing into the base terminal and “I_e” is the current flowing out of the emitter terminal.

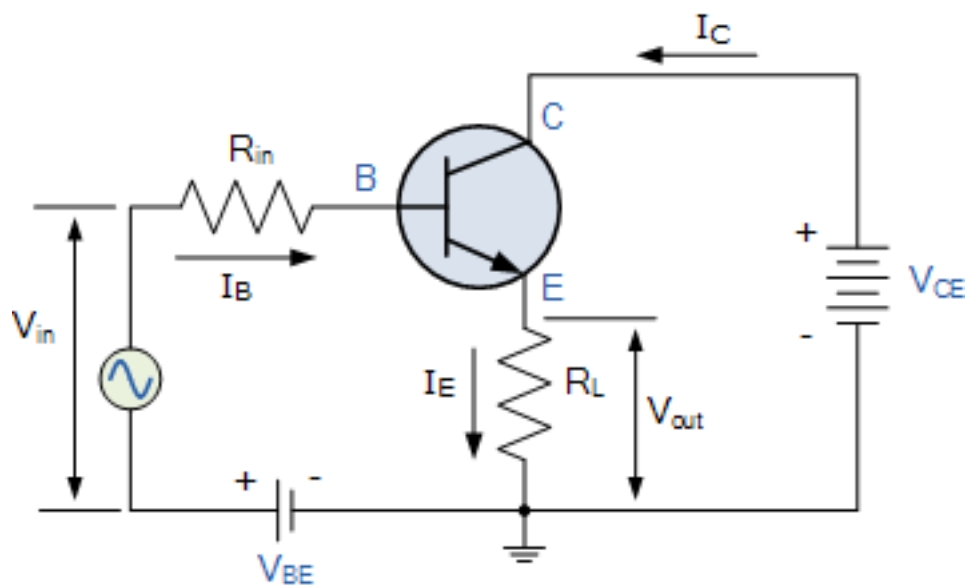
Then to summarise a little. This type of bipolar transistor configuration has a greater input impedance, current and power gain than that of the common base configuration but its voltage gain is much lower. The common emitter configuration is an inverting amplifier circuit. This means that the resulting output signal has a 180° phase-shift with regards to the input voltage signal.

The Common Collector (CC) Configuration

In the **Common Collector** or grounded collector configuration, the collector is connected to ground through the supply, thus the collector terminal is common to both the input and the output. The input signal is connected directly to the base terminal, while the output signal is taken from across the emitter load resistor as shown. This type of configuration is commonly known as a **Voltage Follower** or **Emitter Follower** circuit.

The common collector, or emitter follower configuration is very useful for impedance matching applications because of its very high input impedance, in the region of hundreds of thousands of Ohms while having a relatively low output impedance.

The Common Emitter Transistor Circuit



The common emitter configuration has a current gain approximately equal to the β value of the transistor itself. However in the common collector configuration, the load resistance is connected in series with the emitter terminal so its current is equal to that of the emitter current.

As the emitter current is the combination of the collector AND the base current combined, the load resistance in this type of transistor configuration also has both the collector current and the input current of the base flowing through it. Then the current gain of the circuit is given as:

The Common Collector Current Gain

$$I_E = I_C + I_B$$

$$A_i = \frac{I_E}{I_B} = \frac{I_C + I_B}{I_B}$$

$$A_i = \frac{I_C}{I_B} + 1$$

$$A_i = \beta + 1$$

This type of bipolar transistor configuration is a non-inverting circuit in that the signal voltages of V_{in} and V_{out} are “in-phase”. The common collector configuration has a voltage gain of about “1” (unity gain). Thus it can be considered as a voltage-buffer since the voltage gain is unity.

The load resistance of the common collector transistor receives both the base and collector currents giving a large current gain (as with the common emitter configuration) therefore, providing good current amplification with very little voltage gain.

Having looked at the three different types of bipolar transistor configurations, we can now summarise the various relationships between the transistors individual DC currents flowing through each leg and its DC current gains given above in the following table.

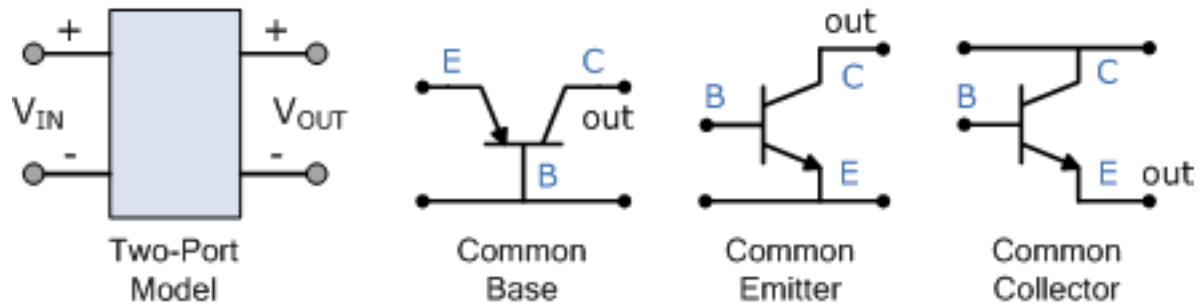
Relationship between DC Currents and Gains

$I_E = I_B + I_C$ $I_C = I_E - I_B$ $I_B = I_E - I_C$	$\alpha = \frac{I_C}{I_E} = \frac{\beta}{1 + \beta}$ $\beta = \frac{I_C}{I_B} = \frac{\alpha}{1 - \alpha}$
$I_B = \frac{I_C}{\beta} = \frac{I_E}{1 + \beta} = I_E (1 - \alpha)$	
$I_C = \beta \cdot I_B = \alpha \cdot I_E$	$I_E = \frac{I_C}{\alpha} = I_B (1 + \beta)$

Note that although we have looked at *NPN Bipolar Transistor* configurations here, PNP transistors are just as valid to use in each configuration as the calculations will

all be the same, as for the non-inverting of the amplified signal. The only difference will be in the voltage polarities and current directions.

Bipolar Transistor Configurations



with the generalised characteristics of the different transistor configurations given in the following table:

Characteristic	Common Base	Common Emitter	Common Collector
Input Impedance	Low	Medium	High
Output Impedance	Very High	High	Low
Phase Shift	0°	180°	0°
Voltage Gain	High	Medium	Low
Current Gain	Low	Medium	High
Power Gain	Low	Very High	Medium

A **signal** is said to be **periodic signal** if it has a definite pattern and repeats itself at a regular interval of time. Whereas, the signal which does not at the regular interval of time is known as an **aperiodic signal** or **non-periodic signal**.

Peak Value Definition: The maximum value attained by an alternating quantity during one cycle is called its Peak value. It is also known as the maximum value or amplitude or crest value.

Average Value

Definition: The average of all the instantaneous values of an alternating voltage and currents over one complete cycle is called **Average Value**.

If we consider symmetrical waves like sinusoidal current or voltage waveform, the positive half cycle will be exactly equal to the negative half cycle. Therefore, the average value over a complete cycle will be **zero**.

$$I_{av} = \frac{i_1 + i_2 + i_3 + \dots + i_n}{n} = \frac{\text{Area of alternation}}{\text{Base}}$$

R.M.S Value

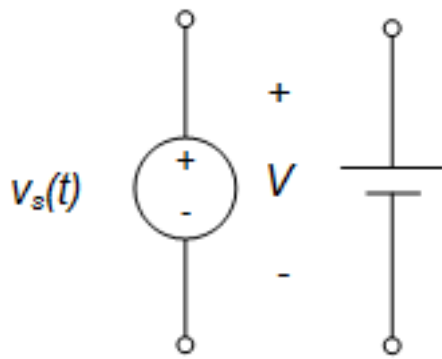
Definition: That steady current which, when flows through a resistor of known resistance for a given period of time than as a result the same quantity of heat is produced by the alternating current when flows through the same resistor for the same period of time is called **R.M.S** or effective value of the alternating current.

In other words, the R.M.S value is defined as the square root of means of squares of instantaneous values.

$$V_{RMS} = \sqrt{\frac{V_1^2 + V_2^2 + V_3^2 + V_4^2 + \dots + V_{11}^2 + V_{12}^2}{12}}$$

Ideal Independent Voltage Sources

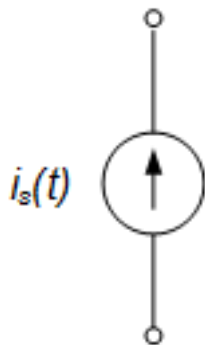
An independent voltage source maintains a specified voltage across its terminals. The symbol used to indicate a voltage source delivering a voltage $V_s(t)$ is shown in Figure. As indicated in Figure the voltage supplied by the source can be time varying or constant (a constant voltage is a special case of a time varying voltage).



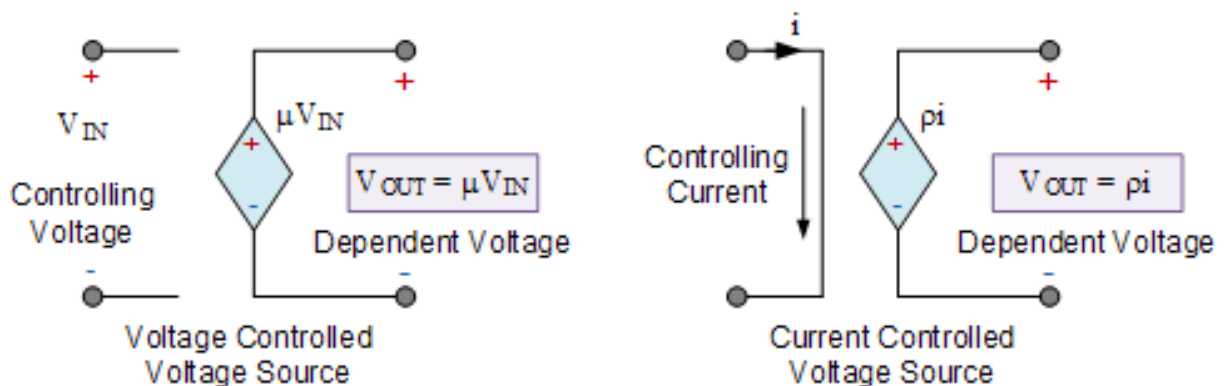
Ideal Independent Current Sources

An ideal independent current source maintains a specified current. The circuit symbol for an ideal independent source is shown in Figure. Note that the current value is listed next to the circuit symbol and that the direction of current flow in the source is provided on the source symbol—there is no need to assume a current direction. The current supplied by the source can be time-varying or constant.

The current provided on the circuit symbol is maintained *regardless* of the voltage difference across the terminals. Even the voltage polarity is unknown and must be determined (if necessary) from an analysis of the overall circuit.



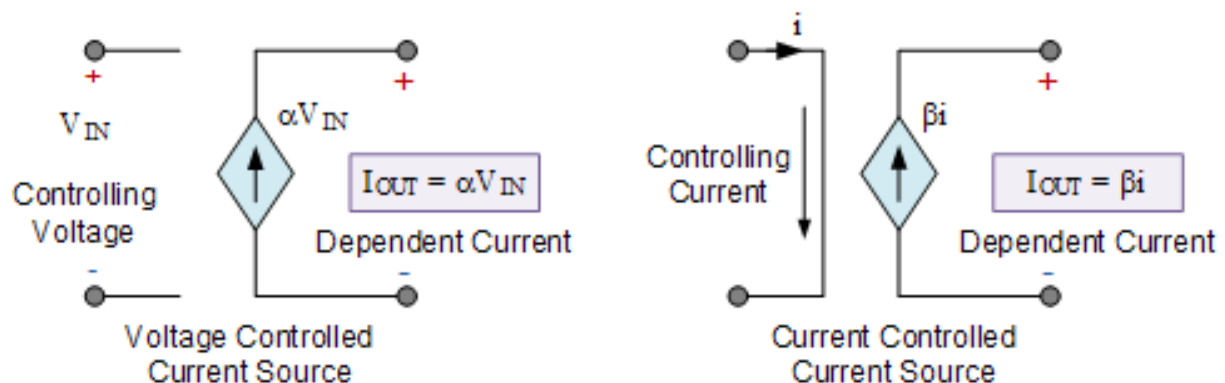
Dependent Voltage Source Symbols



Ideal dependent voltage-controlled voltage sources, VCVS, maintain an output voltage equal to some multiplying constant (basically an amplification factor) times the controlling voltage present elsewhere in the circuit. As the multiplying

constant is, well, a constant, the controlling voltage, V_{IN} will determine the magnitude of the output voltage, V_{OUT} . In other words, the output voltage “depends” on the value of input voltage making it a dependent voltage source and in many ways, an ideal transformer can be thought of as a VCVS device with the amplification factor being its turns ratio.

Dependent Current Source Symbols



An ideal dependent voltage-controlled current source, VCCS, maintains an output current, I_{OUT} that is proportional to the controlling input voltage, V_{IN} . In other words, the output current “depends” on the value of input voltage making it a dependent current source.

Then the VCCS output current is defined by the following equation: $I_{OUT} = \alpha V_{IN}$. This multiplying constant α (alpha) has the SI units of mhos, $\mathcal{U}$ (an inverted Ohms sign) because $\alpha = I_{OUT}/V_{IN}$, and its units will therefore be amperes/volt.

An ideal dependent current-controlled current source, CCCS, maintains an output current that is proportional to a controlling input current. Then the output current “depends” on the value of the input current, again making it a dependent current source.

FET AND MOSFET

Page

Field effect Transistor (FET)

- The operation of the device depends on electric field intensity produced in the channel.
- Voltage controlled device.
- High Ω/cm resistance device ($> 1 \text{ M}\Omega$).
- Lower dissipation is very small.
- Unipolar device.
- Majority Carrier device.
- No minority Carriers.
- Less noisy device than BJT, due to absence of minority Carriers.
- No leakage current & therefore temperature effect on the device is very less & therefore excellent thermal stability.
- Fabricated only with Si.
- When compare to BJT, FET is small in size & easier to fabricate.

Disadvantage :-

- Smaller gain.
- Smaller gain bandwidth product.

Regions of FET

Source :- It is the Source of majority Carrier.

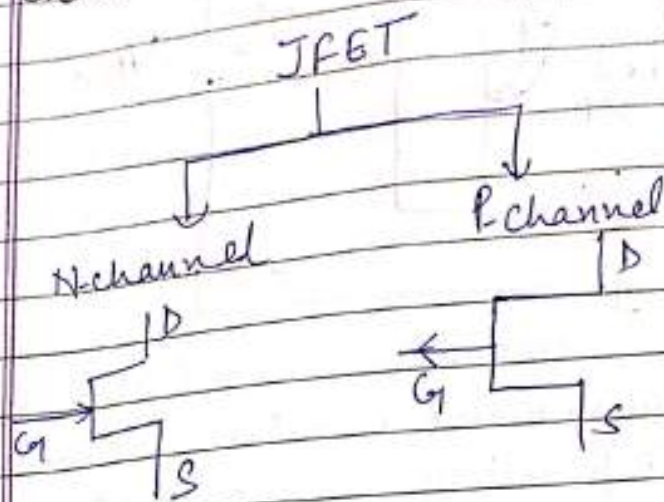
i.e. it is the ~~terminal~~ ^{point} by which majority Carrier will be entering

Drain :- It drains of majority Carrier ^{into drain}. It is the terminal by which majority Carrier will be leaving the device.

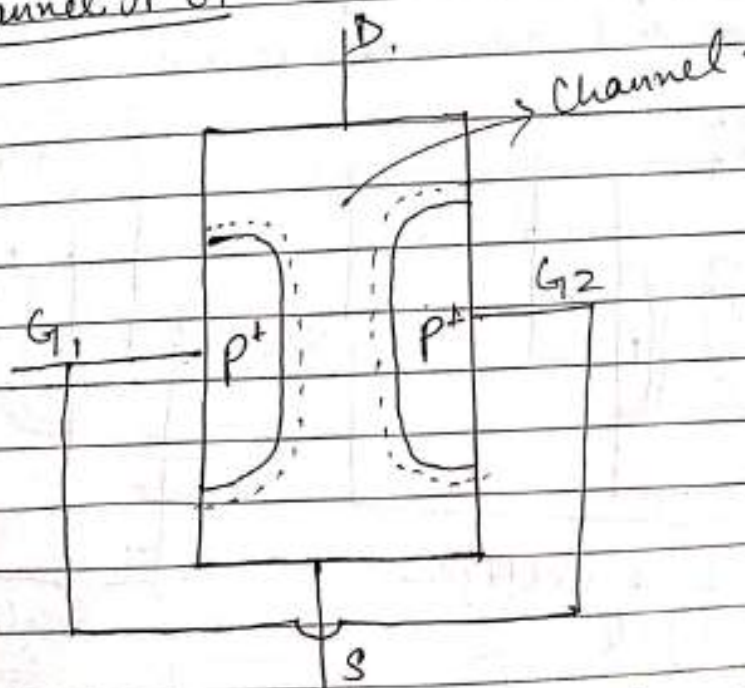
Gate :- It is terminal which control majority

Carriers moving from source to drain.

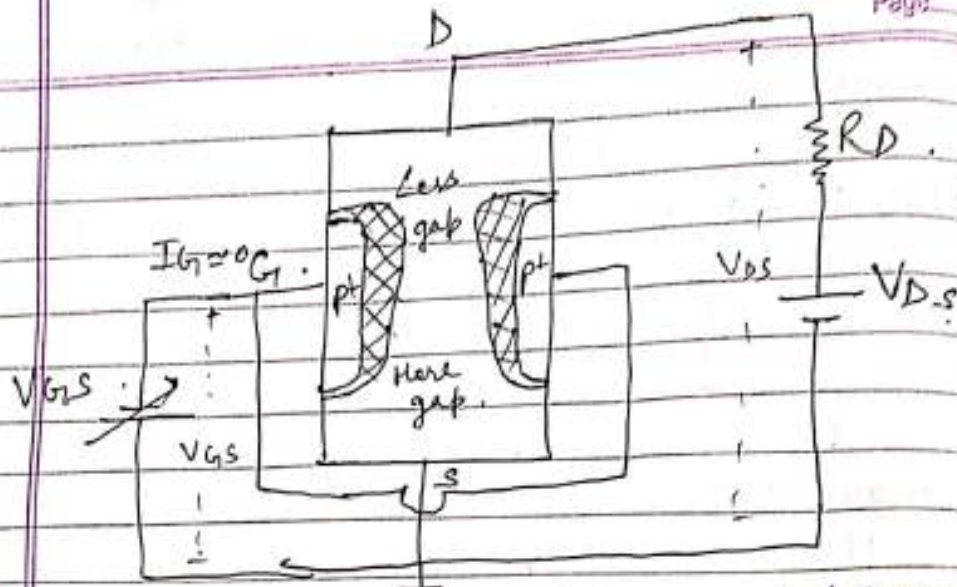
Channel :- It is the region b/w the two gates.



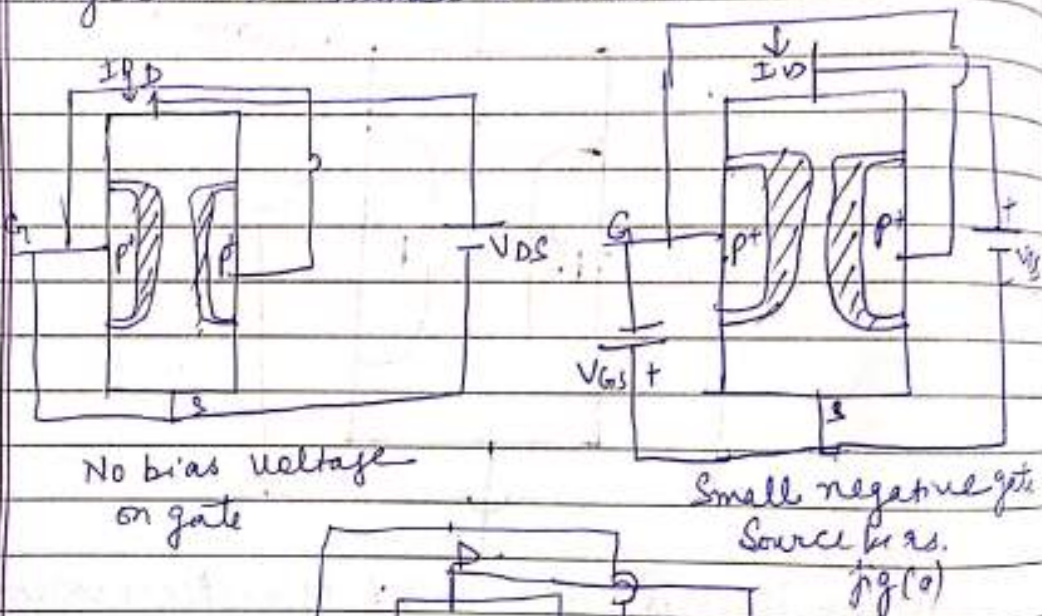
N-channel JFET



→ When JFET is open circuited, channel cross sectional area is maximum & therefore channel current density is minimum.

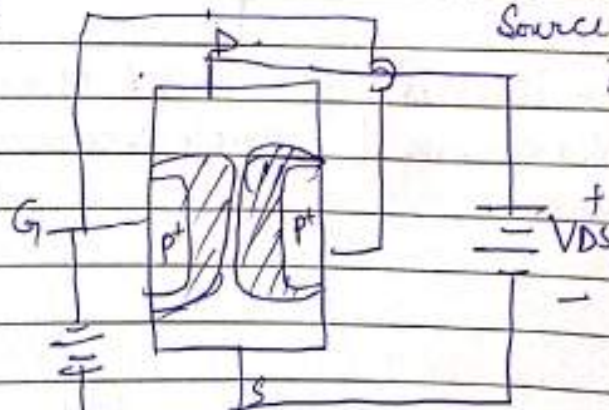


In JFET, the p-n junction between gate & source is always kept in reverse biased conditions. Since current in a reverse biased p-n junction is extremely small, practically zero, the gate current in JFET is often neglected & assumed to be zero.



No bias voltage on gate

Small negative gate source bias V_{GS}



Large negative gate source bias V_{GS}

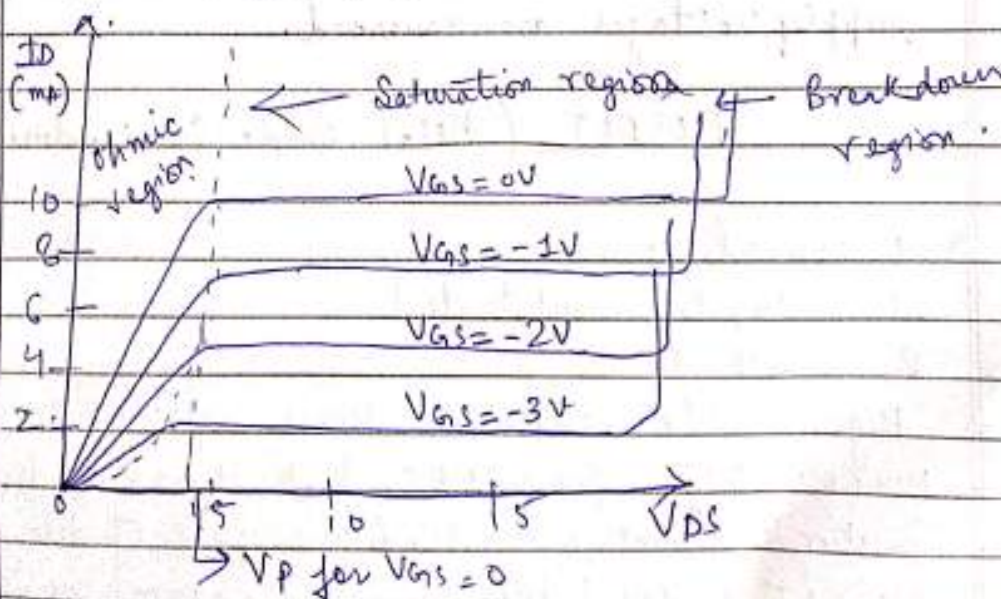
→ when gate to source voltage (V_{GS}) is applied by a d.c. supply (V_{GS}) and increased above zero as depicted in fig (a) the reverse bias voltage across the gate-source junction is increased. Because of this depletion regions are widened. This reduces the effective width of the channel & hence controls the flow of drain current through the channel.

→ when gate to source voltage is further increased, a stage is reached at which two depletion regions touch each other as shown in fig (b). This condition is called Pinch off.

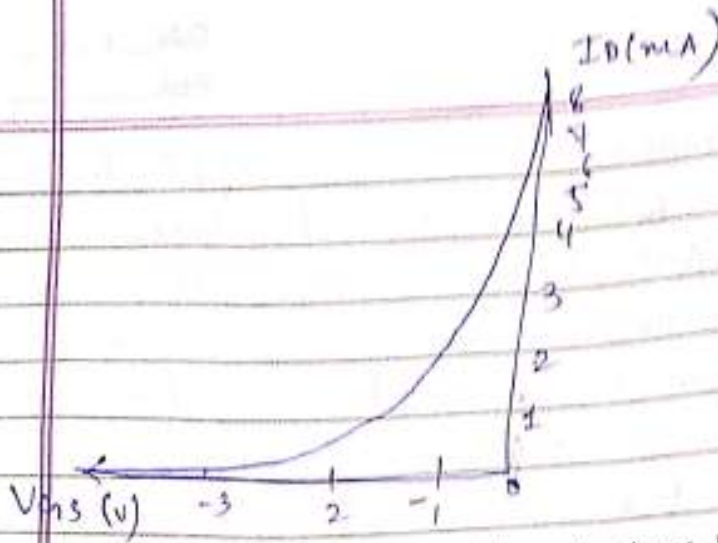
Pinch off Voltage (V_P). It is minimum value of V_{DS} required where I_D enters into Saturation for given value of V_{GS} .

$V_{GS \text{ off}}$ The cut off voltage is the value of V_{GS} at which drain current is zero.

$$V_{GS} = -|V_P|$$



Drain characteristics of n-channel JFET



Transfer characteristics of n channel JFET

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

The relationship b/w I_D & V_{GS} is non linear
this relationship is defined by Shockley's equation.

Note:- The operation of p-channel JFET is exactly similar to n-channel JFET except that the current carriers are holes & polarity of supply voltages are reversed.

MOSFET (Metal oxide Semiconductor FET)

- 4 channel devices
(S, G, D, Substrate)
- $R_i = 10^{10}$ to $10^{15} \Omega$
- Highest Ω/μ resistance device is MOSFET
- MOSFET differs from JFET in that it has no pn junction structure. Instead the gate of the MOSFET is insulated from the channel by SiO_2 layer.

MOSFET↓
Depletion MOSFET

→ there is a pre-existing channel.

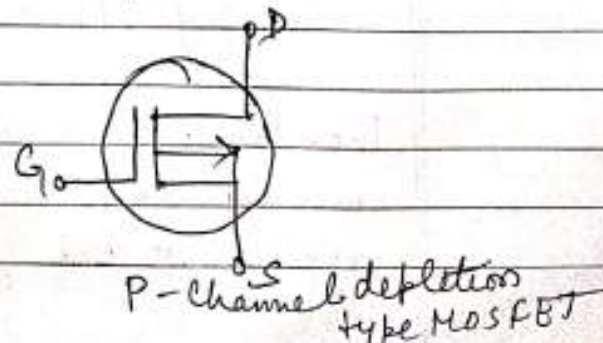
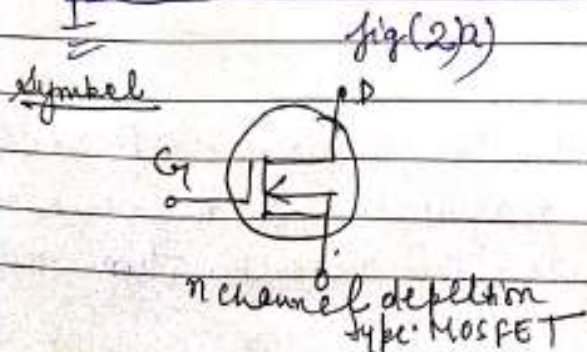
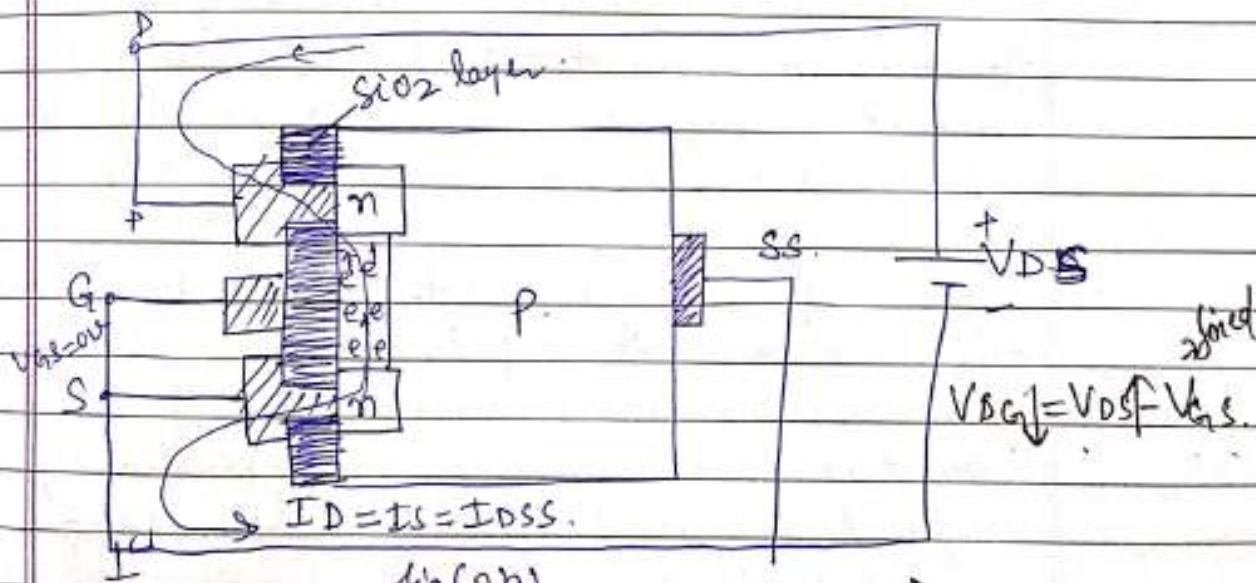
→ Suitable to operate in the depletion mode & enhancement mode.

↓
Enhancement MOSFET

→ no pre existing channel

→ Suitable to operate in the enhance mode only.

→ Input resistance of MOSFET > Input resistance of JFET
 → Lower dissipation is less < Lower dissipation in JFET.

Depletion MOSFET (D-MOSFET)1) n channel Depletion MOSFET.

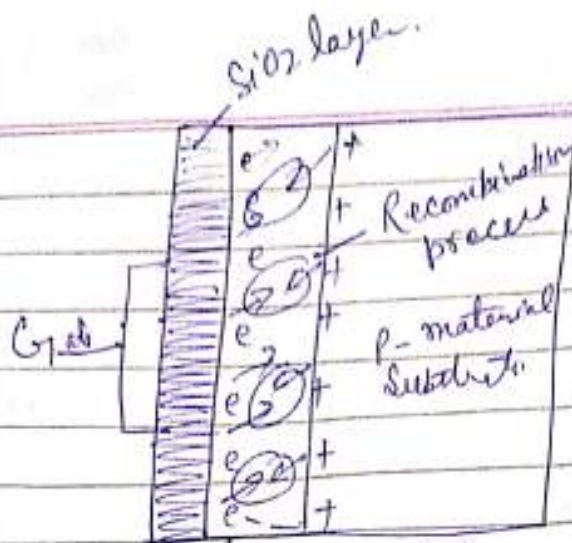


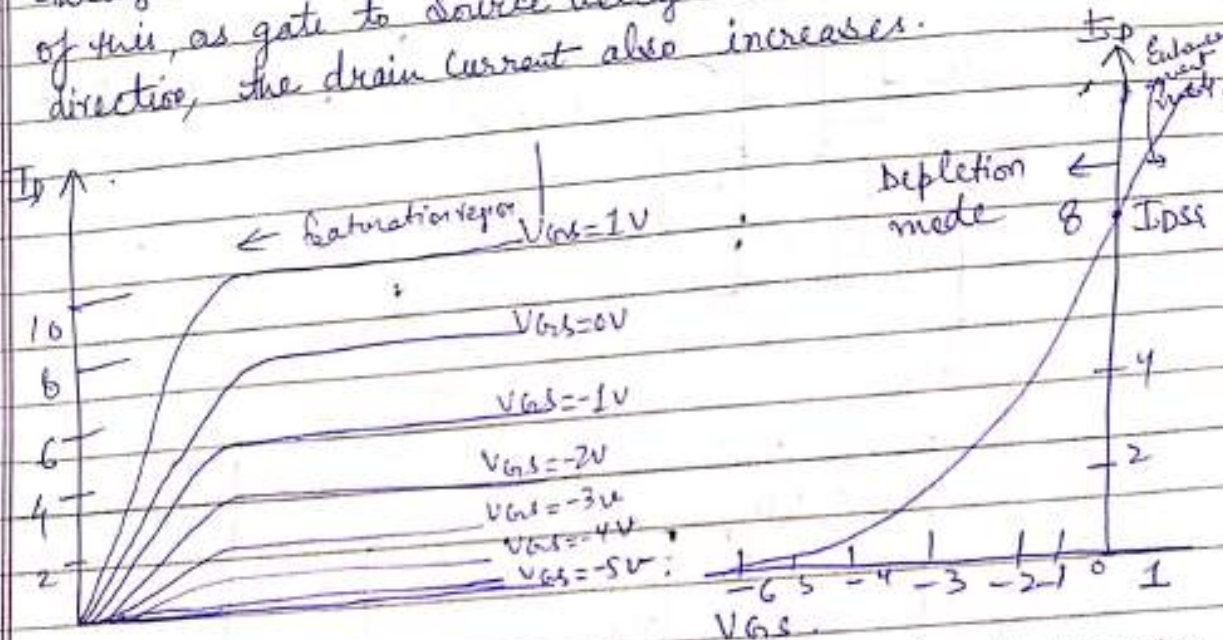
fig 2(b)

(Reduction in free electrons in the n-channel due to negative potential at the gate terminal)

- On the application of drain to Source Voltage, V_{DS} & keeping gate to Source Voltage to zero by directly connecting gate terminal to source terminal, free electrons from the n channel are attracted towards positive terminal of drain terminal. This establishes current through the channel to be denoted as I_{DSS} at $V_{GS} = 0V$ as shown in fig 2(a).
- if we apply negative gate voltage, the negative charges on the gate repel conduction electrons from the channel & attract holes from p-type substrate. This initiates recombination of repelled electrons & attracted holes as shown in fig 2(b).
- The level of recombination between electrons & holes depends on the magnitude of negative voltage applied to the gate. This recombination reduces

the number of free electrons in the n-channel for the conduction reducing the drain current.

For positive value of V_{GS} the positive gate will draw additional electrons from the p-type substrate due to reverse leakage current & establish new carriers through the collisions between accelerating particles. Because of this, as gate to source voltage increases in +ve direction, the drain current also increases.

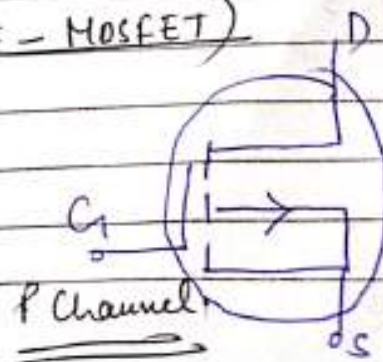
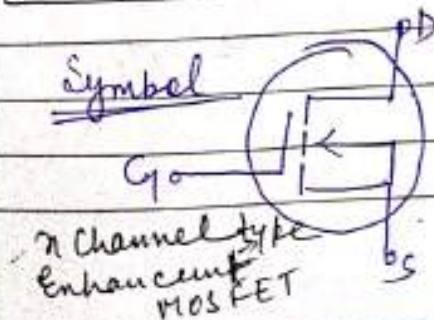


Drain characteristics for an n-channel depletion type MOSFET

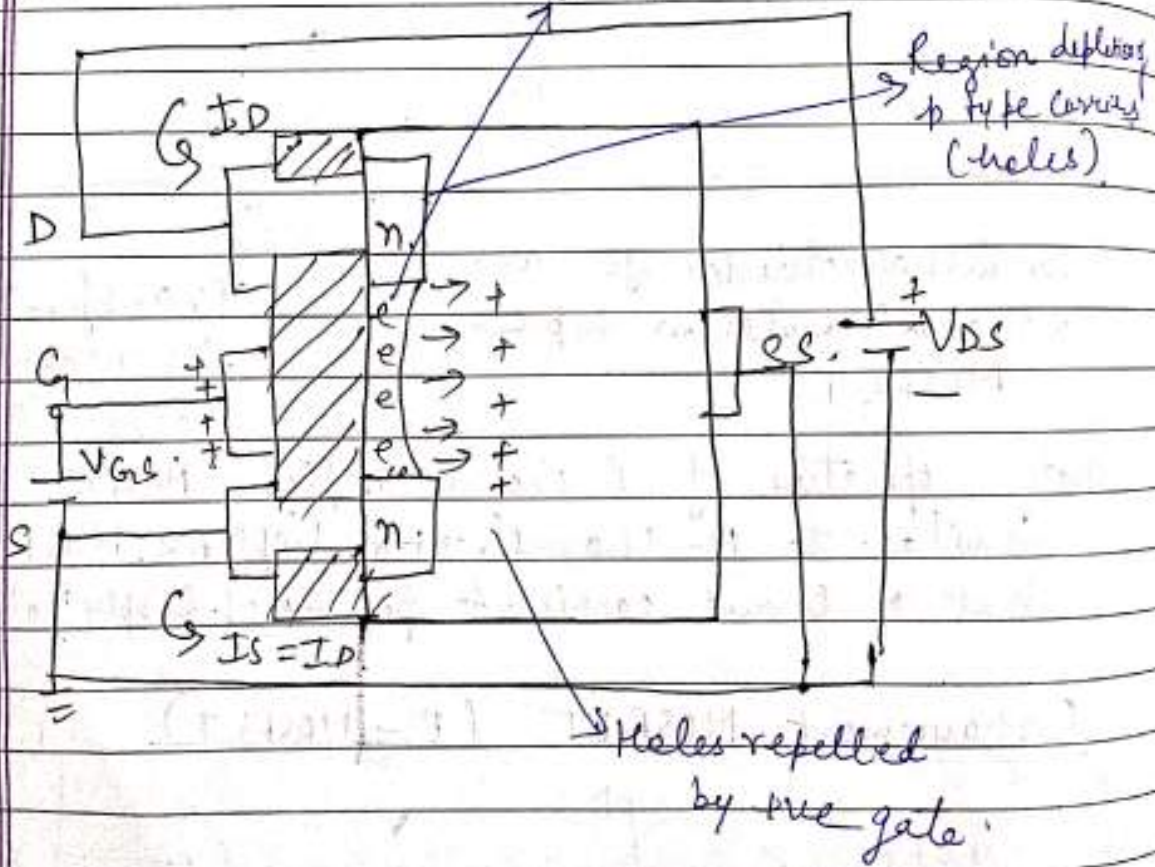
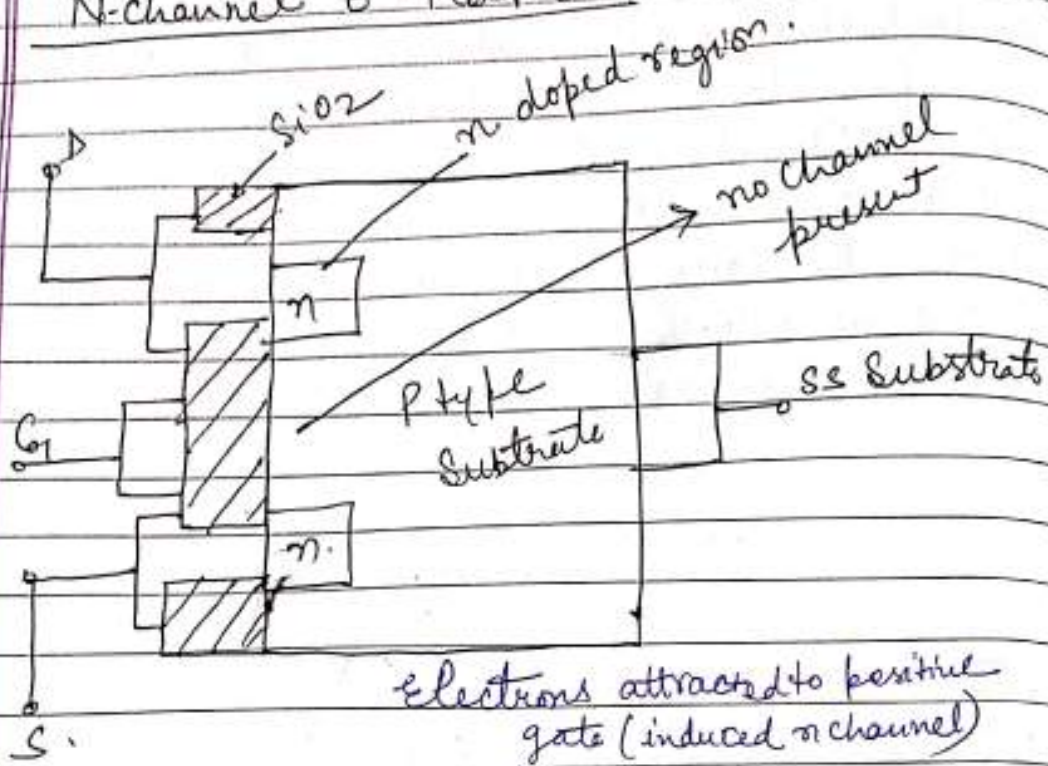
Transfer characteristics for n-channel depletion type MOSFET.

Note:- operation of p-channel Depletion MOSFET is exactly similar to n-channel MOSFET Depletion MOSFET except that of current carrier & polarity of supply voltage.

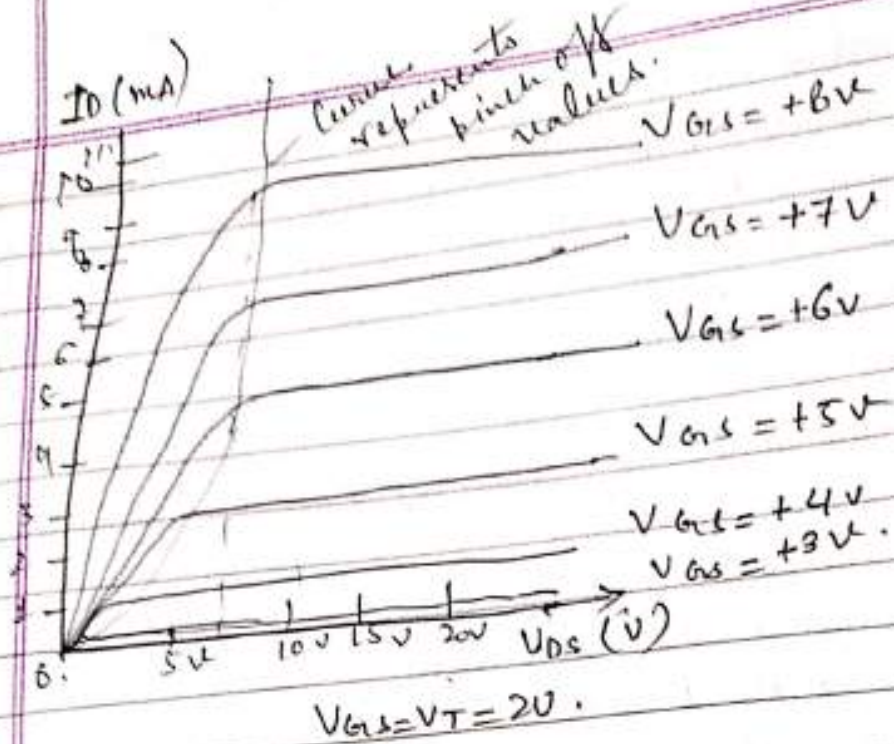
Enhancement MOSFET (E-MOSFET)



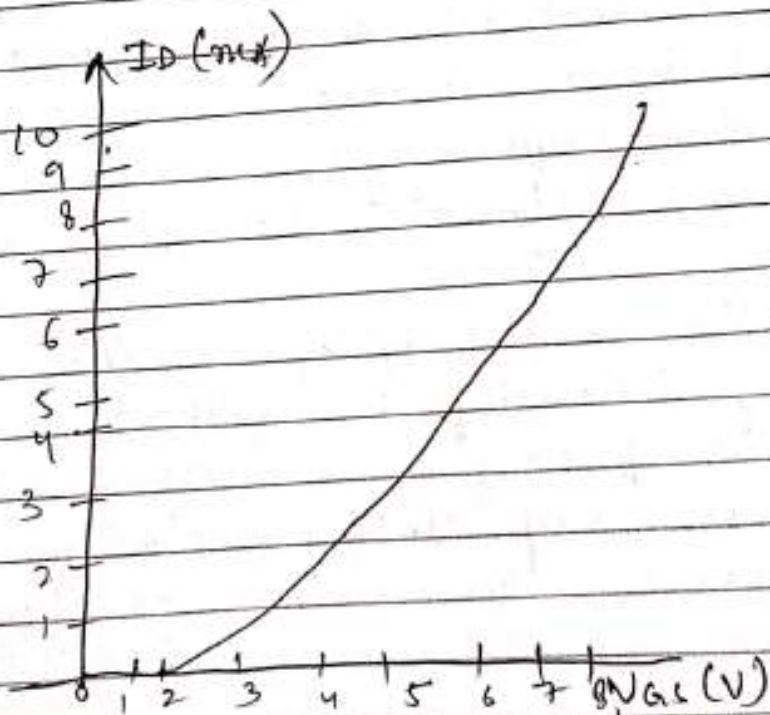
N-channel E-MOSFET



- On application of drain to source voltage V_{DS} & keeping gate to source voltage zero by directly connecting gate terminal to the source terminal, practically zero current flows quite different from the depletion type MOSFET & JFET.
- If we increase magnitude of V_{GS} in the positive direction, the concentration of electrons near SiO_2 surface increases. At a particular value of V_{GS} there is a measurable current flow between drain & source. This value of V_{GS} is called threshold voltage denoted by V_T .
- This we can say that in an enhancement type n channel MOSFET, a positive gate voltage above a threshold value induces a channel and hence the drain current by creating a thin layer of negative charges in substrate region adjacent to the SiO_2 layer.
- The conductivity of the channel is enhanced by increasing the gate to source voltage & thus pulling more electrons into the channel. For any voltage below the threshold value, there is no channel.



Drain characteristics of an n-channel enhancement type MOSFET.



Transfer characteristics for n channel Enhancement type MOSFET.

What is Electric Current?

Electric current is defined as a stream of charged particles—like electrons or ions—moving through a conductor or space. It measures how fast electric charge flows through a medium over time. The symbol for electric current in formulas is “I” or “i”. The unit for current is the ampere (A).

Mathematically, the flow rate of charge with respect to time can be expressed as,

$$I = dQ/dt$$

In other words, a stream of charged particles flowing through an electrical conductor or space is known as an electric current. The moving charged particles are called charge carriers which may be electrons, holes, ions, etc. The flow of current depends on the conductive medium.

For example: In the conductor, the flow of current is due to electrons. In semiconductors, the flow of current is due to electrons or holes.

In an electrolyte, the flow of current is due to ions and In plasma—an ionized gas, the flow of current is due to ions and electrons. When an electrical potential difference is applied between two points in a conductive medium, an electric current starts flowing from higher potential to lower potential.

The higher the voltage or potential difference, the more the current flows between two points. If two points in a circuit are at the same potential, then the current cannot flow. The magnitude of a current depends on the voltage or potential difference between two points. Hence, we can say the current is the effect of voltage.

Electromotive Force (EMF)

Electromotive force is defined as the energy provided by a power source, like a battery or generator, to make electric charge flow through a circuit. Despite its name, the electromotive force is the energy per unit charge or potential difference created by the source.

Basics of Electromagnetic Force EMF

Definition of EMF

- The term electromotive force, or EMF, describes the ability of a source (like a battery or generator) to push charges through a circuit.
- It is that energy delivered per unit charge, which is measured in joules per coulomb (J/C), hence equalling volts (V).

- Electromotive force refers to the voltage associated with the chemical reactions that take place within the battery, as it forces the charge through the circuit.

EMF can mathematically be represented as:

Where:

- ε is the EMF in volts,
- W is the work done in joules,
- Q is the charge in coulombs.

The **electromotive force unit** is the volt, with one volt being the amount of energy needed to move one coulomb of charge with one joule of energy. This relationship is used in understanding how much energy would be required for the flow of current.

Factors Affecting EMF

- **Chemical Reactions:** The type and concentration of chemicals in the battery's electrolyte and the electrodes' material decide the EMF produced.
- **The Formula For EMF In A Circuit**

In a circuit with external resistance R and internal resistance r , the EMF can be calculated using the formula:

Where:

ε - EMF in volts,

I - current in amperes,

R - external resistance in the circuit,

r - internal resistance of the power source.

Difference between EMF and Potential Difference

Electromotive Force (E.M.F)	Potential Difference
The difference in the potential of two electrodes of a battery.	Difference of potential between any two points on the circuit.
E.M.F is always greater than the potential difference between any points in the circuit.	This is always less than the E.M.F
Formula: $E = I(R + r)$	Formula: $V = IR$
This is caused by the electric, gravitational and magnetic fields.	This difference is only produced by the electric field.
The electromotive force is the amount of energy given to each coulomb of charge.	The potential difference is the amount of energy utilized by one coulomb of charge.
The electromotive force is independent of the circuit's internal resistance.	The potential difference is proportional to the circuit's resistance.
The electromotive force is responsible for transferring energy across the circuit.	The potential difference between any two places on the circuit is a measure of energy.
When the circuit is unchanged, the magnitude of the electromotive force is always larger than the potential difference.	When the circuit is completely charged, the size of the potential difference is equal to the circuit's emf.
The electromotive force is measured using an emf meter.	The potential difference is measured with a voltmeter.

Potential Difference : Potential difference is the difference in the amount of energy that charge carriers have between two points in a circuit. Potential difference is measured in volts and is also called voltage.

The work required to transfer a charge from one point to another is described as potential difference. For a charge to pass between two points, there must be a potential difference in current. The potential energy difference per unit charge or voltage is another term for it. The Volt is the unit of potential difference (V).

Subject – Fundamentals of Electrical & Electronics Engineering

Module4 Notes

The potential difference (which is the same as voltage) is equal to the amount of current multiplied by the resistance. A potential difference of one Volt is equal to one Joule of energy being used by one Coulomb of charge when it flows between two points in a circuit.

Electrical Energy:

Electric energy refers to the ability to do work by moving electric charges.

The formula to calculate electric energy is:

$$E = P * t$$

Where:

- E = Electric Energy
- P = Electric Power
- t = Time

Electrical Power:

Electric power is the rate at which an electric circuit transfers electric energy. It represents the amount of work done per unit of time.

The standard unit of electric power is the Watt (W).

Dt-13/09/2025

Magneto motive Force (MMF)- Definition: The current flowing in an electric circuit is due to the existence of electromotive force similarly magneto motive force (MMF) is required to drive the magnetic flux in the magnetic circuit. The magnetic pressure, which sets up the magnetic flux in a magnetic circuit is called Magneto motive Force. The SI unit of MMF is Ampere-turn (AT), and their CGS unit is G (Gilbert). T

he MMF for the inductive coil shown in the figure below is expressed as Where, N – numbers of turns of the inductive coil I – current

The strength of the MMF is equivalent to the product of the current around the turns and the number of turns of the coil.

$$F = N I$$

As per work law, the MMF is defined as the work done in moving the unit magnetic pole (1 weber) once around the magnetic circuit. The MMF is also known as the magnetic potential. It is the property of a material to give rise to the magnetic field. The magneto motive force is the product of the magnetic flux and the magnetic reluctance. The reluctance is the opposition offers by the magnetic field to set up the magnetic flux on it.

The MMF regarding reluctance and magnetic flux is given as

$$F = \Phi R$$

Where R – reluctance Φ – magnetic flux

Magnetic Field

The magnetic field is defined as the region around the magnet where its poles and the electrical charges experience the force of attraction or repulsion. The presence of the field is determined through the needle. In actual practice, the magnetic field has no real existence, and they are purely imaginary.

The magnetic field is created because of the magnetic line of force which possesses the following properties.

1. The magnetic line of force forms the closed loop.
2. The direction of magnetic line force is from north to south. But inside the magnetic line of force their direction is from south to north.
3. The magnetic lines never intersect each other.
4. The magnetic lines repel each other when they are parallel and in the same direction. 5. They are not affected by the non-magnetic materials.

The magnetic flux is defined as the total number of magnetic lines of force produced by the magnet. It is measured in Weber. The one Weber is equal to the 10^8 line of forces or the Maxwell. The Maxwell is the CGS unit of magnetic flux. The magnetic flux is similar to the electric current.

Inductance

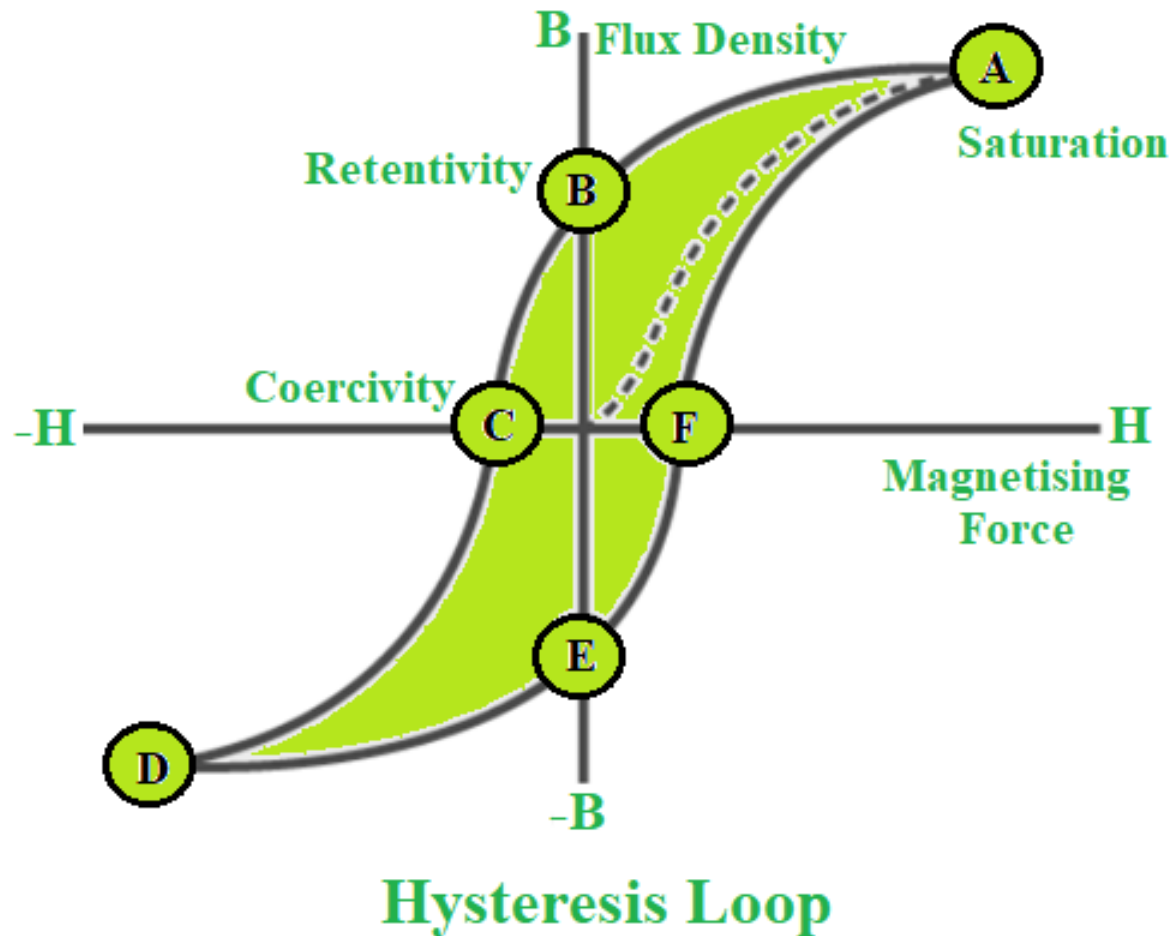
Inductance is defined as the ratio of the induced voltage to the rate of change of current causing it. In the SI system, the unit of inductance is the Henry (H), which is the amount of inductance that causes a voltage of one volt, when the current is changing at a rate of one ampere per second.

Permeability

Permeability is the property of medium or material which measures the easiness offered to pass the magnetic flux when an external magnetic field is applied.

Hysteresis:

Hysteresis Loop- In a system with a magnetic field, hysteresis occurs. Ferromagnetic materials have a common characteristic called hysteresis. The hysteresis effect is a phenomenon that occurs when the magnetization of ferromagnetic materials lags behind the magnetic field. The word hysteresis means "lagging." Magnetic flux density (B) lags after magnetic field strength, resulting in hysteresis (H).



When the magnetic field strength (H) is increased from zero, the magnetic flux density (B) increases.

As the magnetic field is increased, the value of magnetism rises until it hits point A, which is known as the saturation point, where B remains constant.

With a drop in the value of the magnetic field, there is a decrease in the value of magnetism.

However, when B and H are both zero, the substance or material retains some magnetic, which is known as retentivity or residual magnetism.

When there is a reduction in the magnetic field towards the negative side, magnetism likewise decreases.

The material is entirely demagnetized at point C. Coercive force is the amount of force necessary to eliminate a material's retentivity (C).

The cycle is repeated in the opposite direction, with the saturation point D, retentivity point E, and coercive force F.

The cycle is complete due to the forward and opposite direction processes, and this cycle is known as the hysteresis loop.

Retentivity and Coercivity

After an external magnetising field is used to magnetise a ferromagnetic material, the material will not relax back to its zero magnetization state when the external magnetising field is removed.

Retentivity - It is the amount of magnetism that remains after the external magnetising field is withdrawn. It refers to a material's capacity to maintain certain magnetic properties after an external magnetising force has been withdrawn. The value of flux density at the hysteresis loop's point B is the retentivity.

Coercivity - The coercivity of substance is defined as the amount of reverse(-ve H) external magnetising field necessary to totally demagnetize the substance. The value of H at the hysteresis loop's point C is the coercivity.

Difference between the soft magnets and hard magnets

In comparison to hard magnets, soft magnets magnetise and demagnetizes readily.

The retentivity of soft magnets is higher than that of hard magnets.

Hard magnets have a coercivity that is higher than that of soft magnets.

Soft magnets lose less energy than hard magnets due to their tiny surface area.

In the case of soft magnets, the loop area is less than that of hard magnets.

Soft magnets have higher magnetic permeability than hard magnets.

In soft magnets, I and χ are both high, but in hard magnets, they are both low.

Soft magnets are temporary magnets while hard magnets are permanent magnets. Ferrous-nickel alloy, Ferrites, etc. are examples of soft magnets while carbon steel, steel, tungsten, chromium steel, etc. are examples of hard magnets.

Faraday's Laws of Electromagnetic Induction

Faraday's Laws of Electromagnetic Induction describe the relationship between changing magnetic fields and the generation of electromotive force (EMF) or voltage. These laws provide the foundation for understanding how electrical generators, transformers, and inductors work.

Faraday's First Law: When the magnetic flux passing through a coil changes, an EMF is induced in the coil. This change in magnetic flux can be caused by moving a magnet towards or away from the coil, changing the strength of the magnetic field, or changing the orientation of the coil relative to the magnetic field.

Faraday's Second Law: The magnitude of the induced EMF is directly proportional to the rate of change of magnetic flux. In other words, the faster the magnetic flux changes, the greater the induced EMF.

Subject – Fundamentals of Electrical & Electronics Engineering

Module4 Notes

These laws have numerous applications in electrical engineering and technology. For example, they are used in the design and operation of electrical generators, which convert mechanical energy into electrical energy by rotating a coil within a magnetic field. Transformers, which change the voltage of an alternating current (AC) electrical signal, also rely on Faraday's Laws of Electromagnetic Induction.

Faraday's First Law of Electromagnetic Induction:

This law states that whenever there is a change in the magnetic flux passing through a coil of wire, an electromotive force (EMF) is induced in the coil. The magnitude of the induced EMF is directly proportional to the rate of change of magnetic flux.

Mathematically, it can be expressed as:

Where:

- is the electromotive force induced in volts
- is the magnetic flux in webers
- is time in seconds

The negative sign indicates that the induced EMF opposes the change in magnetic flux, according to Lenz's law.

Examples of Faraday's First Law

There are many examples of Faraday's First Law in action. Some of the most common include:

- **Electric generators:** Electric generators convert mechanical energy into electrical energy by using Faraday's First Law. The generator spins a rotor inside a stator, which creates a changing magnetic field. This changing magnetic field induces an EMF in the stator windings, which causes a current to flow.
- **Electric motors:** Electric motors convert electrical energy into mechanical energy by using Faraday's First Law. The motor stator has a series of electromagnets that create a rotating magnetic field. This rotating magnetic field induces an EMF in the motor rotor, which causes a current to flow. The current in the rotor interacts with the magnetic field to create torque, which turns the rotor.
- **Transformers:** Transformers transfer electrical energy from one circuit to another by using Faraday's First Law. The transformer has two coils of wire, a primary coil and a secondary coil. The primary coil is connected to the power source, and the secondary coil is connected to the load. The alternating current in the primary coil creates a changing magnetic field, which induces an EMF in the secondary coil. This EMF causes a current to flow in the secondary coil, which is then transferred to the load.

Faraday's First Law is a fundamental principle of electromagnetism. It has many applications in our everyday lives, from electric generators to electric motors to transformers.

Faraday's Second Law of Electromagnetic Induction:

Subject – Fundamentals of Electrical & Electronics Engineering

Module4 Notes

This law states that the magnitude of the induced EMF is equal to the rate of change of magnetic flux linkage. Magnetic flux linkage (λ) is defined as the product of the number of turns in the coil (N) and the magnetic flux (Φ).

Mathematically, it can be expressed as:

Where:

- is the electromotive force induced in volts
- is the magnetic flux linkage in weber-turns
- is the number of turns in the coil
- is the magnetic flux in webers
- is time in seconds

Example:

Consider a transformer, which consists of two coils of wire, a primary coil, and a secondary coil. When an alternating current (AC) flows through the primary coil, it creates a changing magnetic field. This changing magnetic field induces an EMF in the secondary coil, causing a flow of electric current. The number of turns in the primary and secondary coils determines the voltage transformation ratio of the transformer.

Faraday's Laws of Electromagnetic Induction have revolutionized the field of electrical engineering and have numerous applications in various devices and systems, including:

- Electric generators: Convert mechanical energy into electrical energy by utilizing Faraday's laws.
- Electric motors: Convert electrical energy into mechanical energy by utilizing Faraday's laws.
- Transformers: Change the voltage levels of AC power by utilizing Faraday's laws.
- Inductors: Store electrical energy in a magnetic field by utilizing Faraday's laws.

These laws continue to be fundamental principles in the study and application of electromagnetism, playing a crucial role in the development of modern electrical technologies.

Changing the Magnetic Field Intensity in a Closed Loop

Changing the magnetic field intensity in a closed loop is a fundamental concept in electromagnetism that has numerous applications in various fields. It involves manipulating the strength or direction of the magnetic field within a closed conducting loop, typically achieved by varying the current flowing through the loop or by moving the loop in the presence of an external magnetic field.

1. Faraday's Law of Induction: The key principle governing the change in magnetic field intensity is Faraday's law of induction, which states that a changing magnetic field induces an electromotive force (EMF) or voltage in a closed loop. Mathematically, it can be expressed as:

$$\text{EMF} = -d\Phi/dt$$

Subject – Fundamentals of Electrical & Electronics Engineering

Module4 Notes

where EMF is the electromotive force, Φ is the magnetic flux (the amount of magnetic field passing through the loop), and t represents time. The negative sign indicates that the induced EMF opposes the change in magnetic flux.

2. Lenz's Law: Lenz's law provides an additional rule to determine the direction of the induced EMF and the resulting current. It states that the induced current flows in a direction that opposes the change in magnetic flux. In other words, the induced magnetic field created by the current opposes the original change in magnetic field.

3. Applications:

a. **Electric Generators:** Electric generators convert mechanical energy into electrical energy by utilizing the principle of changing magnetic field intensity. As a rotating loop of wire (the armature) moves within a stationary magnetic field (the stator), the changing magnetic flux induces an EMF in the loop, causing an electric current to flow.

b. **Electric Motors:** Electric motors operate on the reverse principle. By supplying an electric current to a coil of wire (the stator), a magnetic field is generated. When a conducting loop (the rotor) is placed within this magnetic field, the changing magnetic flux induces an EMF in the loop, causing it to rotate.

c. **Transformers:** Transformers transfer electrical energy from one circuit to another through electromagnetic induction. They consist of two coils of wire (primary and secondary) wound around a shared iron core. When an alternating current flows through the primary coil, it creates a changing magnetic field that induces an EMF in the secondary coil, resulting in voltage transformation.

d. **Magnetic Levitation (Maglev) Trains:** Maglev trains use the principle of changing magnetic field intensity to achieve high-speed transportation. Powerful electromagnets on the track generate a changing magnetic field that induces currents in conducting loops on the underside of the train. These currents create opposing magnetic fields that levitate the train above the track, reducing friction and enabling incredibly fast speeds.

In summary, changing the magnetic field intensity in a closed loop is a fundamental concept in electromagnetism with numerous practical applications. By understanding and manipulating the relationship between changing magnetic fields and induced currents, we can harness this phenomenon to generate electricity, power motors, transform voltage, and even levitate trains.

Examples of Faraday's Law:

- **A bar magnet is moved towards a coil of wire.** As the magnet approaches the coil, the magnetic flux through the coil increases. This induces an EMF in the coil, which causes a current to flow. The direction of the current is such that it opposes the increase in magnetic flux.
- **A conducting loop is rotated in a magnetic field.** As the loop rotates, the magnetic flux through the loop changes. This induces an EMF in the loop, which causes a current to flow. The direction of the current is such that it opposes the change in magnetic flux.
- **A transformer is used to step up or step down the voltage of an alternating current (AC) power supply.** The transformer consists of two coils of wire, a primary

Subject – Fundamentals of Electrical & Electronics Engineering

Module4 Notes

coil and a secondary coil. The primary coil is connected to the AC power supply, and the secondary coil is connected to the load. As the AC current flows through the primary coil, it creates a changing magnetic field. This changing magnetic field induces an EMF in the secondary coil, which causes a current to flow in the load. The voltage of the current in the secondary coil is proportional to the number of turns in the primary and secondary coils.

Faraday's Law of Electromagnetic Induction is a fundamental principle of electromagnetism. It has many applications in electrical engineering, such as the design of generators, transformers, and motors.

Lenz's Law

Lenz's law is a fundamental principle of electromagnetism that describes the direction of the electromotive force (EMF) induced in a conductor when it is exposed to a changing magnetic field. It states that the EMF induced in a conductor is always such that it opposes the change in magnetic flux through the conductor. In other words, Lenz's law predicts the direction of the current that will flow in a conductor when it is exposed to a changing magnetic field.

Mathematical Formulation

Lenz's law can be mathematically expressed as follows:

where:

- is the EMF induced in the conductor (in volts)
- is the magnetic flux through the conductor (in webers)
- is time (in seconds)

The negative sign in the equation indicates that the EMF induced in the conductor opposes the change in magnetic flux.

Examples

Here are a few examples of Lenz's law in action:

- **A bar magnet is moved towards a coil of wire.** As the magnet approaches the coil, the magnetic flux through the coil increases. This induces an EMF in the coil that causes a current to flow in the opposite direction to the motion of the magnet. This current creates a magnetic field that opposes the motion of the magnet, slowing it down.
- **A conducting rod is moved through a magnetic field.** As the rod moves through the magnetic field, the magnetic flux through the rod changes. This induces an EMF in the rod that causes a current to flow in the opposite direction to the motion of the rod. This current creates a magnetic field that opposes the motion of the rod, slowing it down.
- **A solenoid is connected to a battery.** When the battery is turned on, the current flowing through the solenoid creates a magnetic field. This magnetic field induces an EMF in the solenoid that causes a current to flow in the opposite direction to the current from the battery. This current creates a magnetic field that opposes the magnetic field from the battery, reducing the overall magnetic field strength.

Applications

Lenz's law has a wide range of applications in electrical engineering and physics. Some of the most common applications include:

- **Electric motors:** Lenz's law is used to explain the operation of electric motors. When a current is passed through a coil of wire in a magnetic field, the coil experiences a force due to Lenz's law. This force causes the coil to rotate, which in turn drives the motor.
- **Generators:** Lenz's law is also used to explain the operation of generators. When a conductor is moved through a magnetic field, an EMF is induced in the conductor. This EMF can be used to generate electricity.
- **Magnetic brakes:** Lenz's law is used to design magnetic brakes. When a metal disk is rotated in a magnetic field, the disk experiences a force due to Lenz's law. This force opposes the motion of the disk, slowing it down.

Lenz's law is a fundamental principle of electromagnetism that has a wide range of applications in electrical engineering and physics. It is a powerful tool for understanding the behavior of electromagnetic systems.

Faraday's Law Derivation

Faraday's Law of Induction states that a changing magnetic field induces an electromotive force (EMF) in a conductor. This EMF is proportional to the rate of change of the magnetic flux through the conductor.

Derivation of Faraday's Law

Consider a conducting loop of wire placed in a magnetic field. The magnetic flux through the loop is given by:

where:

- is the magnetic flux (in webers, Wb)
- is the magnetic field (in teslas, T)
- is a differential area vector (in square meters, m^2)

If the magnetic field is changing, then the magnetic flux through the loop will also change. This changing magnetic flux will induce an EMF in the loop. The EMF is given by:

EMF

where:

- EMF is the electromotive force (in volts, V)
- is time (in seconds, s)

The negative sign in the equation indicates that the EMF opposes the change in magnetic flux.

Example

Subject – Fundamentals of Electrical & Electronics Engineering

Module4 Notes

Consider a solenoid with a length of 1 m and a radius of 0.1 m. The solenoid is wound with 1000 turns of wire. The current in the solenoid is increasing at a rate of 1 A/s.

The magnetic field inside the solenoid is given by:

where:

- is the magnetic field (in teslas, T)
- is the permeability of free space ($4 \times 10^{-7} \text{ T m/A}$)
- is the number of turns per unit length (in turns/meter)
- is the current (in amperes, A)

In this case, turns/meter and A. Therefore, the magnetic field inside the solenoid is:

$$T \text{ m/A turns/m A T}$$

The magnetic flux through the solenoid is given by:

$$T \text{ m Wb}$$

The EMF induced in the solenoid is given by:

$$\text{EMF Wb V}$$

The negative sign indicates that the EMF opposes the increase in magnetic flux.

Faraday's Experiment: Relationship Between Induced EMF and Flux

Michael Faraday conducted a series of experiments in the early 19th century that laid the foundation for our understanding of electromagnetic induction. One of his most famous experiments involved a coil of wire connected to a galvanometer, a device that measures electric current. When Faraday moved a magnet in and out of the coil, he observed that the galvanometer needle deflected, indicating the presence of an electric current. This current was induced by the changing magnetic flux through the coil.

Magnetic Flux

Magnetic flux is a measure of the amount of magnetic field passing through a given area. It is defined as the dot product of the magnetic field vector and the area vector:

where:

- is the magnetic flux in webers (Wb)
- is the magnetic field vector in teslas (T)
- is the area vector in square meters (m^2)

The direction of the magnetic flux is given by the right-hand rule. If you point your right thumb in the direction of the magnetic field vector, and your fingers in the direction of the area vector, then your palm will point in the direction of the magnetic flux.

Induced EMF

Subject – Fundamentals of Electrical & Electronics Engineering

Module4 Notes

When the magnetic flux through a coil of wire changes, an electromotive force (EMF) is induced in the coil. This EMF is given by Faraday's law of induction:

where:

- is the induced EMF in volts (V)
- is the magnetic flux in webers (Wb)
- is the time in seconds (s)

The negative sign in Faraday's law indicates that the induced EMF opposes the change in magnetic flux. In other words, the induced EMF creates a magnetic field that opposes the change in the original magnetic field.

Examples

Here are a few examples of Faraday's law of induction in action:

- When you move a magnet in and out of a coil of wire, the galvanometer needle deflects, indicating the presence of an induced EMF.
- When you turn on a light switch, the electric current flows through the light bulb, creating a magnetic field. This magnetic field induces an EMF in the nearby wires, which causes the light bulb to glow.
- When you use a metal detector, the changing magnetic field created by the metal object induces an EMF in the detector's coil, which causes the detector to beep.

Faraday's law of induction is a fundamental principle of electromagnetism. It has many applications in everyday life, from the electric motors that power our appliances to the generators that produce the electricity that we use.

Applications of Faraday's Law

Here are some examples of applications of Faraday's law:

- **Electric generators:** Electric generators convert mechanical energy into electrical energy by using Faraday's law. A generator consists of a coil of wire that is rotated in a magnetic field. The rotating magnetic field induces an EMF in the coil of wire, which causes an electric current to flow.
- **Electric motors:** Electric motors convert electrical energy into mechanical energy by using Faraday's law. A motor consists of a coil of wire that is placed in a magnetic field. When an electric current flows through the coil of wire, it creates a magnetic field that interacts with the magnetic field of the motor. This interaction creates a force that causes the motor to rotate.
- **Transformers:** Transformers change the voltage of an alternating current (AC) electrical signal by using Faraday's law. A transformer consists of two coils of wire that are wrapped around a common iron core. When an AC electrical signal is applied to one coil of wire, it creates a changing magnetic field in the iron core. This changing magnetic field induces an EMF in the other coil of wire, which causes an AC electrical signal with a different voltage to flow.
- **Magnetic levitation (maglev) trains:** Maglev trains use Faraday's law to levitate above the tracks. Maglev trains have superconducting magnets that create a strong

magnetic field. This magnetic field induces an EMF in the conducting rails on the tracks, which creates a force that levitates the train.

Faraday's law is a fundamental principle of electromagnetism that has many important applications in our everyday lives.

18/09/2025

Induced EMF

When a magnetic flux linking a conductor or coil changes, an electromotive force (EMF) is induced in the conductor or coil, is known as *induced EMF*. Depending upon the way of bringing the change in magnetic flux, the induced EMF is of two types –

- Statically Induced EMF
- Dynamically Induced EMF

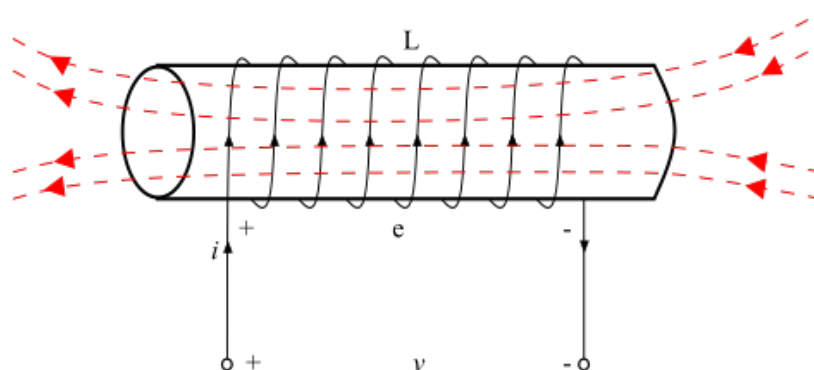
Statically Induced EMF

When the conductor is stationary and the magnetic field is changing, the induced EMF in such a way is known as *statically induced EMF* (as in a transformer). It is so called because the EMF is induced in a conductor which is stationary. The statically induced EMF can also be classified into two categories –

- Statically Induced EMF
- Mutually Induced EMF

Self-Induced EMF

When an EMF is induced in the coil due to the change of its own magnetic flux linked with it is known as *self-induced EMF*.



Explanation – When a current flows in a coil, a magnetic field produced by this current through the coil. If the current in the coil changes, then the magnetic field linking the coil also changes.

Therefore, according to Faraday's law of electromagnetic induction, an EMF being induced in the coil. The induced EMF in such a way is known as *self-induced EMF*.

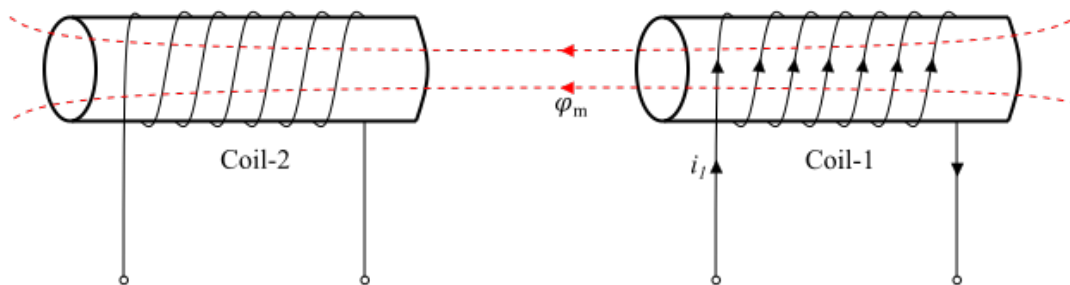
Mathematically, self-induced EMF is given by,

$$e = L \frac{di}{dt} \dots (1)$$

Where, L is the self-inductance of the coil.

Mutually Induced EMF

When an EMF is induced in a coil due to changing magnetic flux of neighbouring coil is known as *mutually induced EMF*.



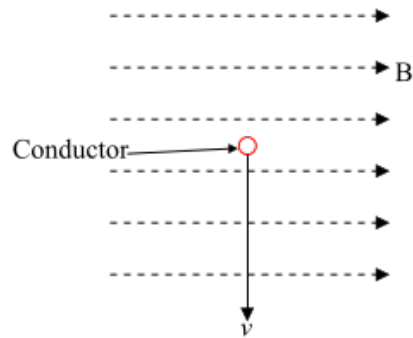
Explanation – Consider two coils coil-1 and coil-2 placed adjacent to each other (see the figure). A fraction of the magnetic flux produced by coil-1 links with the coil-2. This magnetic flux which is common to both the coils 1 and 2 is known as *mutual flux* (ϕ_m). Now, if the current in coil-1 changes, the mutual flux also changes and thus EMF being induced in both the coils. The EMF induced in coil-2 is known as *mutually induced EMF*, since it is induced due changing in flux which is produced by coil-1. Mathematically, the mutually induced EMF is given by,

$$e_m = M \frac{di_1}{dt} \dots (2)$$

Where, M is the mutual inductance between the coils.

Dynamically Induced EMF

When the conductor is moved in a stationary magnetic field so that the magnetic flux linking with it changes in magnitude, as the conductor is subjected to a changing magnetic, therefore an EMF will be induced in it. The EMF induced in this way is known as *dynamically induced EMF* (as in a DC or AC generator). It is so called because EMF is induced in a conductor which is moving (dynamic).

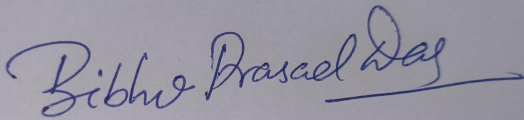


Comparison between Electrical Circuit and Magnetic Circuits

Parameter	Electric Circuits	Magnetic Circuits
Definition	A closed path followed by electric current is known as electric circuit.	A closed path followed by magnetic field line or magnetic flux is known as magnetic circuit.
Circuit quantities	In electric circuit, the major quantities associated with the circuit are EMF, voltage, current, resistance, capacitance and inductance, etc.	The fundamental quantities associated with magnetic circuits are MMF, magnetic flux, reluctance, etc.
Quantity flowing in the circuit	The electric current is the quantity which flows in an electric circuit.	The magnetic flux is the quantity which flows in a magnetic circuit.
Expression of Ohm's law	For electric circuit, Ohm's law is as ? Current , $I = \text{EMF} / \text{Resistance}$	For magnetic circuit, the expression of Ohm's law is ? Magnetic flux , $\phi = \text{MMF} / \text{Reluctance}$
Driving force	The EMF or electromotive force is the driving force in an electric circuit, which creates potential difference so that current can flow in the circuit.	In case of a magnetic circuit, the MMF or Magnetomotive force is the driving force, which causes the flow of magnetic flux in the core of the magnetic circuit.
Opposition	The resistance of the electric circuit opposes the electric current flowing in the circuit.	The reluctance of the core of magnetic circuit opposes the flow of magnetic flux in the circuit.
Measure of opposition	In an electric circuit, the resistance opposes the flow of current, which is given by, $R = \rho l / a$ Ω	In a magnetic circuit, the reluctance produces the opposition in the flow of magnetic flux, which is given by, $S = l / \mu a$ AT/Wb

Subject – Fundamentals of Electrical & Electronics Engineering
Module4 Notes

Reciprocal of opposition	In an electric circuit, the resistance produces the opposition whose reciprocal is known as conductance and is given by, Conductance, $G=1/R$	In a magnetic circuit, the opposition to the magnetic flux is reluctance whose reciprocal is known as permeance, and is given by, Permeance $=1/\text{Reluctance}$
Density	The electric current flows in an electric circuit whose density (called current density) is given by, $J=I/a \text{ A/m}^2$	Magnetic flux flows in a magnetic circuit whose density, called magnetic flux density, is given by, $B=\phi/a \text{ Wb/m}^2$
Field intensity	In an electric circuit, electric field exists whose field intensity is given by, $E=V/d$	In a magnetic circuit, there is a magnetic field whose intensity is $H=NI/l$



Bibhu Prasad Das