

Reading 25

Ron Bertrand VK2DQ
<http://www.res.net.au>

SEMICONDUCTORS Part 2

TRANSISTOR CONFIGURATIONS

There are different ways of connecting a transistor so that it will amplify. The principle of small changes in base current controlling larger changes in collector current remains the same. However, the different configurations we are going to discuss give the transistor amplifier different characteristics, which makes some configurations best for certain purposes.

I am not going to go very deep into this subject, as there is no need to for examination purposes. For the exam all you really need to be able to do is identify which configuration is illustrated, and remember some basic characteristics of each configuration.

The three configurations are:

- Common-emitter
- Common-base, and
- Common-collector (also called emitter-follower)

I will give you the basic circuits of each. For the exam, you need to identify which circuit is which, and that is all. Please do remember that with all transistor amplifier configurations, it is still the small base current controlling the large collector current which results in amplification. At the end of this section I will also provide a table of the different characteristics of each configuration. A detailed description is not necessary as this would be rather lengthy, and contain much information which is not required. You can consult a handbook if you really want to know more.

COMMON-EMITTER CIRCUIT

The basic transistor amplifier circuit is the common-emitter. It is also the most common configuration you will see and is easy to remember. The circuit shown in figure 1 is of a basic common-emitter amplifier using an NPN transistor. If a PNP transistor were to be used the polarity of the batteries would be reversed (current flows against the arrow).

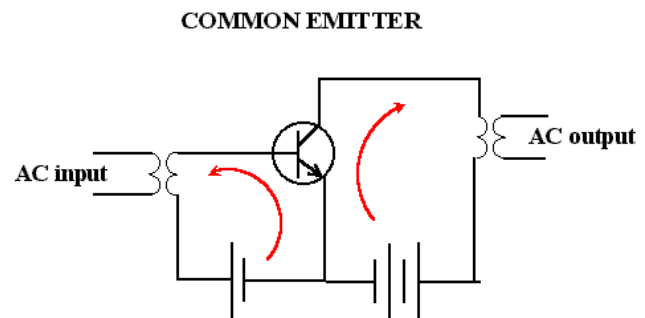


Figure 1.

This amplifier circuit uses transformer coupling at both the input and output. The transformers are air cores so this indicates that the circuit is used at radio frequencies, as

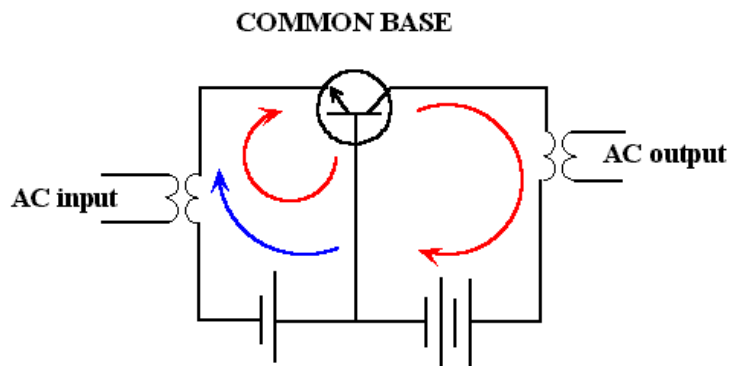
iron core transformers have too much loss in the laminated iron core at radio frequencies. The type of coupling does not determine the configuration.

The reason the configuration of figure 1 is called common-emitter is because the **emitter is the common leg (or lead) between input and output**. This is how you will identify the common-emitter circuit.

COMMON-BASE

The common base circuit is **identified by the common base leg between input and output**. In this circuit transformer coupling is once again used.

Figure 2.



from 1000 to 8 ohms.

Transformer coupling in these two circuits is for **impedance matching**. We have learnt that one of the functions of a transformer, if required, is to match two unequal impedances. For example, if you had an audio amplifier with an output impedance of 1000 ohms, it would not work well if connected to an 8 ohm speaker. An audio output transformer could be used to perform an impedance conversion

The controlling base current is shown flowing in the correct direction by the red curved line on the left. The other two lines, red on the right and blue on the left, are the collector and emitter currents respectively. The emitter current is the sum of the base bias current and the collector current. The difference is small, as the base current is small.

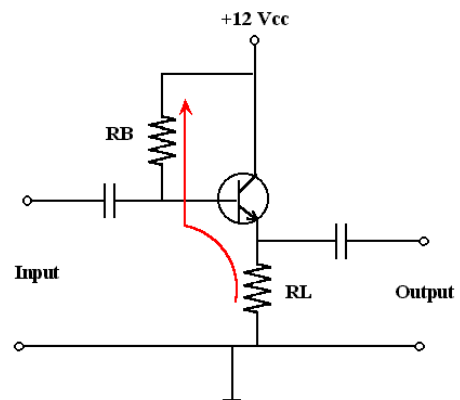
Just a note, in all the transistor circuits we have seen so far the bias voltage (between base and emitter) and the collector voltage (between collector and emitter) has used separate batteries. In practical circuits one source of supply is used for base bias and collector voltage.

EMITTER-FOLLOWER

Figure 3 shows an emitter-follower, also called a common collector configuration. I have shown how one source of supply can be used for both base-emitter forward bias and collector voltage. Unlike the previous two circuits this one is a little different. **The key to the identification of the common-collector (emitter-follower) is that the output is taken from across a load or resistor in the emitter circuit**. The output of this amplifier and all amplifiers of this configuration is taken from across a resistor shown R_L .

The only reason the different configurations are used is that the amplifiers provide different characteristics, a summary of which I will provide in table form shortly -

Figure 3.
EMITTER-FOLLOWER



they are all just amplifiers, relying on small changes in base current to control larger changes in collector current.

The coupling between the input and output of the circuit in figure 3 is called capacitive coupling. Coupling is simply the method by which input is fed to, and output taken from, a transistor amplifier.

Notice also how a **single resistor R_B is used to forward bias the base-emitter junction of the transistor**, doing away with having to use two power sources. I would also like to point out that a number of transistors could be connected via coupling to produce a multistage amplifier. I am just giving you here what is needed for the exam and nothing more, as the topic is too large, and frankly of no benefit to fill your head with unnecessary information at this stage. Though if you are going to do some construction I would strongly suggest further reading.

There you have it. The three basic transistor configurations and an illustration on biasing with a single voltage in the last circuit. We have been using BJT's here. We will be talking about other types of transistors soon, and we have already talked about the electron tube. All of these **active devices** can be connected in the above configurations.

Summary of the characteristics of each configuration:

Parameter	CE	CB	CC
Input impedance	1000 ohms	60 ohms	40,000 ohms
Output impedance	40,000 ohms	200,000 ohms	1000 ohms
Voltage gain	500	800	0.96
Current gain	20	0.95	50
Power gain	10,000	760	48
Phase in/out	180 deg	0 deg	0 deg

Table 1 – Summary of Configuration Characteristics

- A common-emitter has a low input impedance and a high output impedance and inverts the phase of the incoming signal.
- A common-base also has a low input impedance and has a high output impedance and does *not* create a phase inversion between input and output.
- A common-collector has high input impedance and a low output impedance and provides no phase change. (A common collector could drive a low impedance load such as a 50Ω antenna or an 8Ω speaker).

BIASING TRANSISTORS

Biasing transistors is another broad topic that much could be written about, however I am again going to confine this reading to what I think is the 'need to know' and even that may extend further than necessary.

Biasing a transistor means applying a fixed DC voltage between the base and emitter. By varying the amount of bias we can make the transistor amplify in the required class of operation i.e. A, B, AB or C.

If you have trouble understanding the need for bias (and many students do) imagine using a transistor without any bias, and feeding an AC signal to it to be amplified. Half of the AC signal would not forward bias the base-emitter junction and there would be no output from the transistor during this half cycle. Transistors (including FET's) and electron tubes are not AC devices – they work on varying DC. When we bias the input of an active device at a DC potential, we can then feed an AC signal to that input. The combination of the AC superimposed on top of the DC bias produces varying DC at the input of the transistor or other active device.

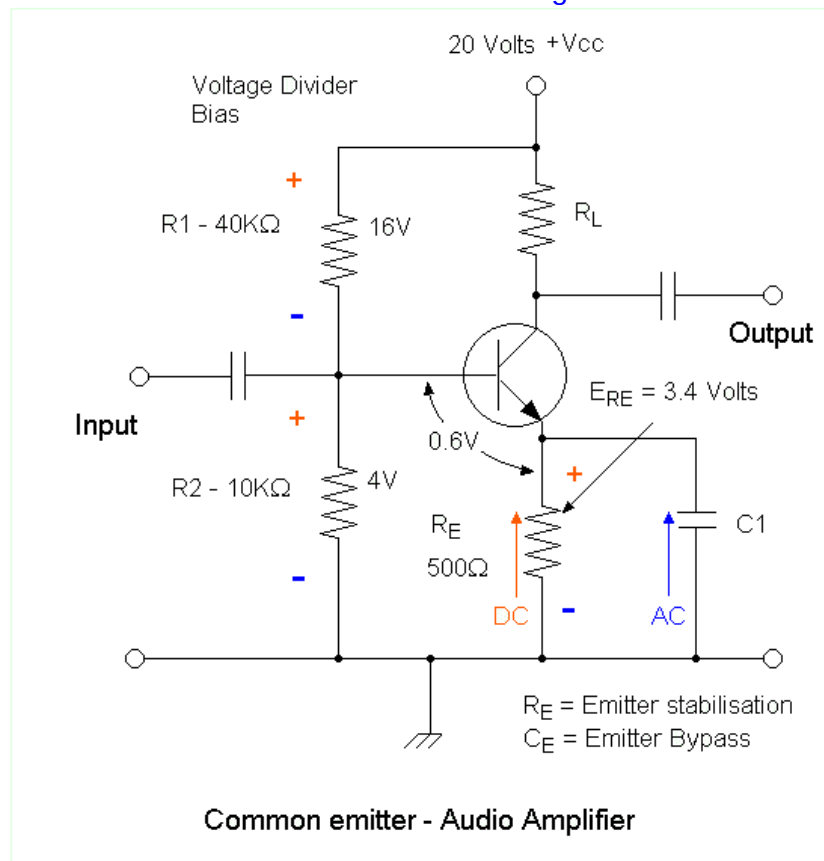
We have already seen one method of biasing using a single resistor (R_B) in the emitter-follower circuit shown earlier. This type of biasing is called **fixed bias**. I suggest you go back and have a look at it now.

Fixed bias is simple, however it is quite **unstable thermally**. If the transistor warms for any reason, due to a rise in ambient (surrounding) temperature or due to current flow through it, the collector current increases. The higher the current gain of the transistor the greater the instability of the circuit.

A far superior and the most common biasing arrangement for BJT's is shown in the schematic circuit of figure 4 and is called **voltage-divider bias**. I have shown typical values for a practical audio amplifier. You will not have to perform any calculations on the circuit but you may need to identify the circuit, not as a whole but in separate parts, so we will have a look at it to the necessary detail.

Figure 4.

Common-emitter circuit with voltage divider bias



The purpose of bias in any BJT (transistor) circuit is to **forward bias the base emitter junction** of the transistor. In the circuit of figure 4 this is achieved primarily by the series circuit consisting of R_1 and R_2 in series with the supply voltage (in this case 20 volts). The transistor is an NPN type which requires the base to be positive with respect to the emitter for forward bias. R_1 and R_2 in series, is across the 20 volt supply which is positive at the top (collector) and negative at the earth symbol. Now you could use ohms law to approximate the voltage across R_1 and R_2 knowing they are across a 20 volt supply and have the values shown. Their total resistance is 50K. 40/50 of 20 volts will be dropped across R_1 and 10/50 (1/5) of the voltage will be dropped across R_2 . One fifth of 20 volts is 4 volts. So the approximate voltage across R_2 is 4 volts. The top of R_2 is positive and the bottom negative. So neglecting R_E , there is about 4 volts applied between the base and the emitter of the transistor, positive at the base and negative at the emitter. The polarity is correct, so we have a forward biased base-emitter junction. This is how voltage divider bias works. R_1 and R_2 form a voltage divider.

Let's just have a look at some of the other components in the circuit and what they do. Resistor R_E you will notice is a relatively low value (500 ohms) and is connected between the emitter and ground. R_E provides thermal stability to the circuit. An increase in DC voltage through R_E caused by a reduction in the transistor collector-emitter resistance will cause an increase in voltage across R_E . **The voltage across R_E is in opposition to the bias voltage.** If the voltage across R_E was 1 volt then this opposes the voltage across R_2 and so the effective bias voltage is only 3 volts. The voltage across R_E will only increase if the transistor heats up and as a consequence its collector-emitter resistance decreases. Without R_E the transistor could go into **thermal runaway**.

Let's look at the situation without R_E . The transistor is operating. It heats up due to the current through it or an increase in the ambient (room or surrounding) temperature. If the transistor heats up, its emitter-collector current increases which causes it to heat up more, which causes its resistance to decrease more, which increases the current more, which causes it to heat up more etc. - this is thermal runaway and results in the transistor being destroyed.

R_E provides what is called **degeneration**. Degeneration is a single word, which explains the operation in the second last paragraph. If the current increases through R_E (due to the transistor heating up for whatever reason) the **voltage across R_E reduces the bias voltage, that decreases the emitter collector current**, allowing the transistor to cool and not go into thermal runaway.

Now, it is the **DC** current flowing through R_E that we are using to provide thermal stabilisation and prevent thermal runaway. There is an **AC** component in the circuit. This AC component is the signal we are trying to amplify. An AC voltage is applied to the input of the transistor via capacitive coupling and that AC signal will appear in the emitter collector circuit, amplified. **We do not want the AC signal to flow through R_E .** If the AC component (the signal to be amplified) did flow through R_E it would affect the bias of the transistor. So the simple approach is to provide an AC path around R_E for the AC component. This is accomplished by adding the parallel capacitor across R_E (C_1). So, the signal which is being amplified flows through C_1 and does not affect the biasing of the transistor, while the DC current, which is what affects the operating characteristic of the transistor and can cause thermal runaway, is allowed to flow through R_E .

If this is hard to understand, let me add the following points to refresh your memory. The circuit of the transistor has a current, this current without any input signal is just **DC**. When

an AC (or varying DC) signal is applied to the transistor, the current through the transistor now has **two** components (or parts): a DC component, which is responsible for bias; and an AC current which is the signal being amplified. The DC component is used to bias the transistor and prevent thermal runaway - the degenerative voltage across R_E must be treated differently by the transistor than the AC component, which is the signal we are trying to amplify.

If we were to allow the AC component of the current to flow through R_E , then any increase in the signal to be amplified would reduce the gain (amplification) of the transistor, defeating the whole purpose of amplification.

A capacitance will pass an AC current, but block a DC current. So the bias current flows through R_E (DC), while the signal to be amplified passes through C_1 . For this to work effectively the **capacitive reactance of C_1 must, as a rule of thumb, be about $1/10^{\text{th}}$ the value of R_E** . So in the circuit of figure 4 above, the X_C of C_1 would be about 50 ohms.

Since C_1 bypasses the AC component of the emitter-collector current around R_E , C_1 is called an **emitter bypass capacitor, or just a bypass capacitor**. R_E is called the emitter stabilisation resistor.

You do not have to explain how this works as I have done. But you do need to know the purpose of R_1 , R_2 , R_E , and C_1 (which may be labelled differently).

- **R_1 and R_2 provide voltage divider bias.**
- **R_E provides thermal stability.**
- **C_1 bypasses the AC component (ie. the signal to be amplified around R_E).**

R_L is just a load resistor. If it were not there, the output would be connected to the +20 DC rail and there would be no output other than +20DC - hardly a useful amplifier.

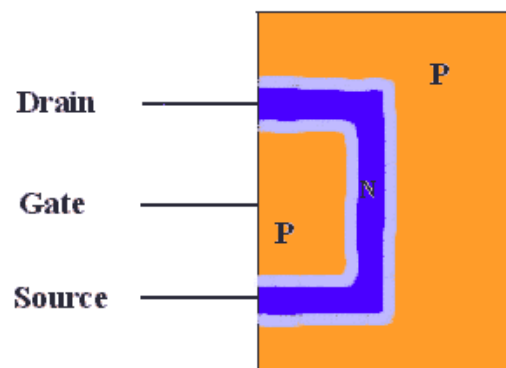
JUNCTION FIELD EFFECT TRANSISTOR (JFET)

Field-effect and bipolar junction transistors are entirely different devices. In fact, the FET's are solid-state amplifying devices that have operating characteristics very similar to those of triode electron tubes. There are three types of FET's, the junction (JFET), and two metal oxide semiconductor types known as MOSFET's. One is the depletion MOSFET, the other is an enhancement MOSFET.

The essential construction of a JFET is shown in figure 5. A block of P-material called the **substrate** has an N-type channel running through it. The figure shows a cross section. Think of the N-channel as a wormhole in the P substrate. As with any PN junction, electrons fall into holes at the barrier and a depletion zone is created. You can see the N channel running through the P material (substrate) - the depletion zone is drawn as light blue. The connections are made as shown. The drain can be thought of as similar to the anode, the source similar to the cathode, and the gate similar to the control grid of a triode valve. Though this device is fully manufactured from semiconductors, the purpose of the JFET is to amplify, as is the

Figure 5.

JFET



purpose of the BJT and the triode. If a supply voltage is connected between drain and source, there will be a current flow as there is no PN-junction either forward or reverse biased.

So we have connected a supply voltage between drain and source, and you can imagine a current flowing from source to drain through the N-channel.

Notice also that the depletion zone occupies part of the channel.

To refresh your memory, a depletion zone is called a depletion zone because it is depleted of **charge carriers**. A charge carrier is either an electron or a hole. A PN junction can be reverse biased increasing the size the depletion zone, and how much voltage is applied to a PN junction in the reverse bias condition will determine the **size** of the depletion zone.

If we wanted to reverse bias the JFET in the previous diagram we would apply a negative potential to the substrate that would attract the holes, and a positive potential to the channel that would attract electrons.

If we reverse bias the PN junctions on each side of the channel, we would **increase the size of the depletion zones**. Since the depletion zone would widen, they would **occupy more of the channel**. In fact, if we applied enough reverse bias the depletion zones could be made to occupy the entire channel, and no current would flow from source to drain. This is called **pinch off**.

We have a way of controlling the amount of current that flows in the channel by changing the amount of reverse bias between the gate connection and the channel (the substrate and the gate are connected together).

The gate is always biased negative in the JFET (n-channel), then the signal to be amplified can be applied to the gate. **The signal to be amplified would change the reverse bias and hence the size of the depletion zone. Since the depletion zone occupies part of the channel the resistance of the channel will alter with the input signal and in doing so control the source-drain current.**

So the JFET is an amplifier. Since the gate connection is always reverse biased (no current), the input impedance is very high. It is the changes in the negative potential on the gate which controls the much larger drain current.

This is, in essence, exactly how a triode works - by changing the negative potential of the control-grid of a triode we can control the much larger anode current. If the control-grid of a triode is made negative enough, no electrons will flow from the cathode to the anode - this is called **cut-off**.

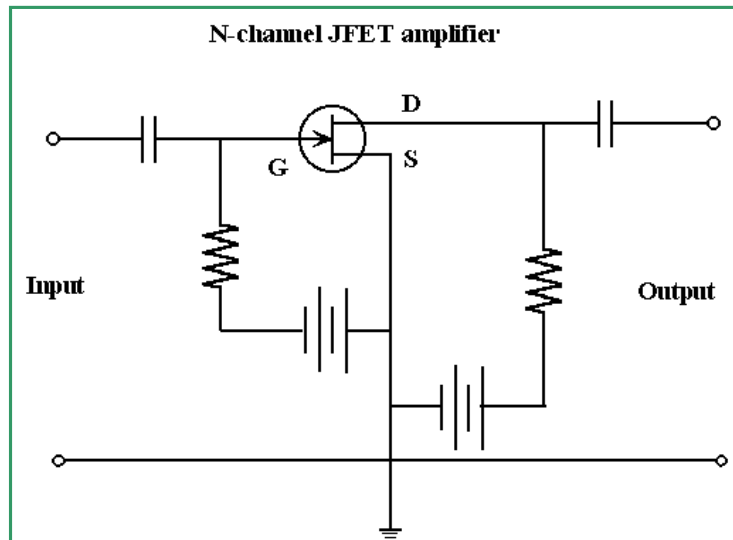
Similarly, if the negative potential of the gate of an N-channel JFET is large enough, there will be no source-drain current. This is called **pinch off**.

So, to run through the operation of an N-channel JEFT one more time, the gate is biased negative with respect to the source, causing the depletion zone to occupy part of the channel. How much depletion zone there is in the channel controls the channels resistance. So if a signal is applied between the gate and the source, it will control proportionately a much larger drain current. The drain current will be an amplified version of the signal applied to the gate.

The JFET is a **voltage amplifier** like a triode. By this we mean that changes of voltage **only** on the gate are required to control the drain current. Since voltage alone is needed to control the JFET, the input signal does not have to deliver any power. You can have all the voltage in the world, but without any current you have no power ($P=EI$). This is very different to the BJT which is a current amplifier. You can have voltage without current, but not current without voltage - think about that.

You could also have a P-channel JFET that would have an N-substrate. The operation would be the same except you would have to bias the gate positive.

Figure 6 - Schematic Diagram of Basic N-Channel JFET Amplifier.



Important points to remember: Symbol; High input impedance; Voltage Amplifier; the semiconductor version of a triode valve; why high input impedance can be an advantage.

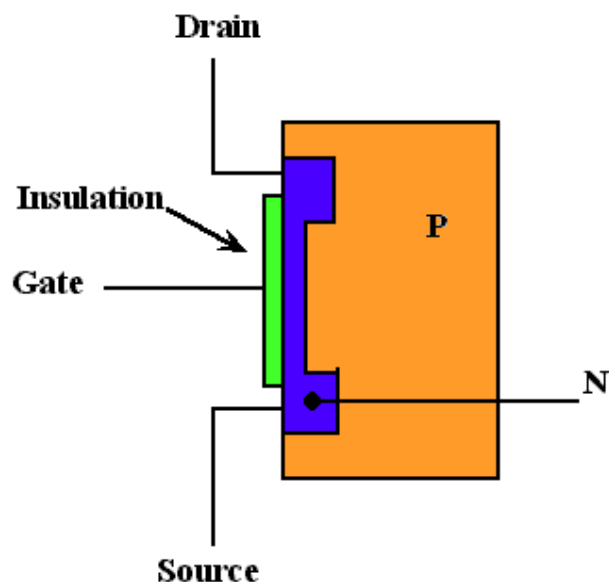
MOSFETS

Metal Oxide Semiconductor Field Effect Transistors.

Figure 7.

When a FET is constructed as shown in figure 7 with the **gate insulated** (silicon dioxide) from the very narrow N-channel, it is called an insulated gate FET, IGFET, or MOSFET, the latter being the most common term. Combinations of these transistors as a single unit are called CMOS, MOS, VMOS.

When a voltage is applied, **negative to gate and positive to source**, its electrostatic field extends into the N channel. This reduces its ability to carry current by repelling channel electrons, depleting the channel of its carriers. This is a depletion **MOSFET**.



Sufficient reverse bias can pinch off the drain current completely. Conversely, a forward bias can increase drain current up to a certain point. As a result, this device can operate with no bias at all, which simplifies circuitry. Unlike in vacuum tubes and JFET's, **a forward bias does not produce any input circuit current** due to the insulation between gate and channel. MOSFET's have extremely high input impedance in the order of 10 megohms.

The input impedance of MOSFETS (as much as 40 megohms) and similar devices are so high that they require special handling, as normal handling can result in a static charge being created which can destroy the device. You may find the some wrapped in metal foil to protect them from static voltages.

Figure 8.

N channel MOSFET



The MOSFET was a great breakthrough for semiconductor technology, and led to combinations of them being fabricated within a single package called an integrated circuit. Of course even BJT'S and JFET'S can be fabricated in combination with MOSFETS into a single package. It is a simple matter to create resistors in an integrated circuit by controlling the width and doping level of semiconductor strips. Capacitance in an integrated circuit can make use of the fact that a reverse biased PN junction acts like a capacitor. You are not likely to be

asked anything about MOSFETS in the examination, though you never know, so I have added this short description for completeness. Do know the symbol.

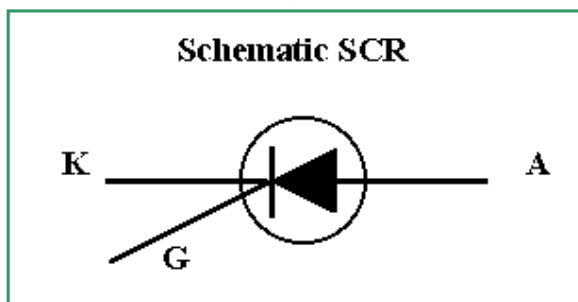
PROTECTED GATE MOSFET's

MOSFETS have very high input impedance, as high as 10 megohms. Because of this high input impedance even static electricity from handling them can build up a high voltage on the gate (or gates) and rupture the Metal Oxide insulator. Some MOSFETS are made with two zener diodes connected in parallel and in opposite directions between the gate and the substrate. Under normal operation neither of these zeners will conduct. However, if a high static voltage is created on the gate, one of the zeners will conduct and discharge the voltage. These are called protected gate MOSFETS. The symbol is that of an ordinary MOSFET, but you will see within the symbol the zener diode symbols. You may have to identify this.

SILICON CONTROLLED RECTIFIERS (SCR'S)

Exam purposes - **Identify the symbol - external operation and application.** For practicality reasons I will cover the internal operation of the SCR briefly. You need to concentrate your memory to the exam points I indicated above.

Figure 9.



The SCR construction shown in figure 10 illustrates one of the many NPNP, or multi-layer semiconductor devices. This one is called an SCR. "A" is the anode, "K" is the cathode and "G" the gate. If a voltage, even a large one, is place across anode and cathode (of any polarity), this device will not conduct because at least one junction will be reverse biased. The main current path of

an SCR is between cathode and anode.

Switch SW1 is normally open (NO). When SW1 is closed, the first P-Type region, acting like the base of an NPN BJT, is forward biased through R_G and R_L which allows current to flow through J2. Since J3 is forward biased already, current flows through the whole device from K to A and the load.

The gate switch SW1 can now be opened and current will continue to flow because J2 loses control of its current carriers. Even reverse biasing the gate will not stop cathode-anode current. It is necessary to reduce the applied voltage (V_{aa}) to almost zero or open SW2 to stop current flow. No current will flow if the source potential (V_{aa}) is reversed, even with SW1 closed.

Therefore, with an AC source rather than V_{aa} , which is DC, the SCR acts like a rectifier in that it allows current to flow in one direction only.

Figure 10.

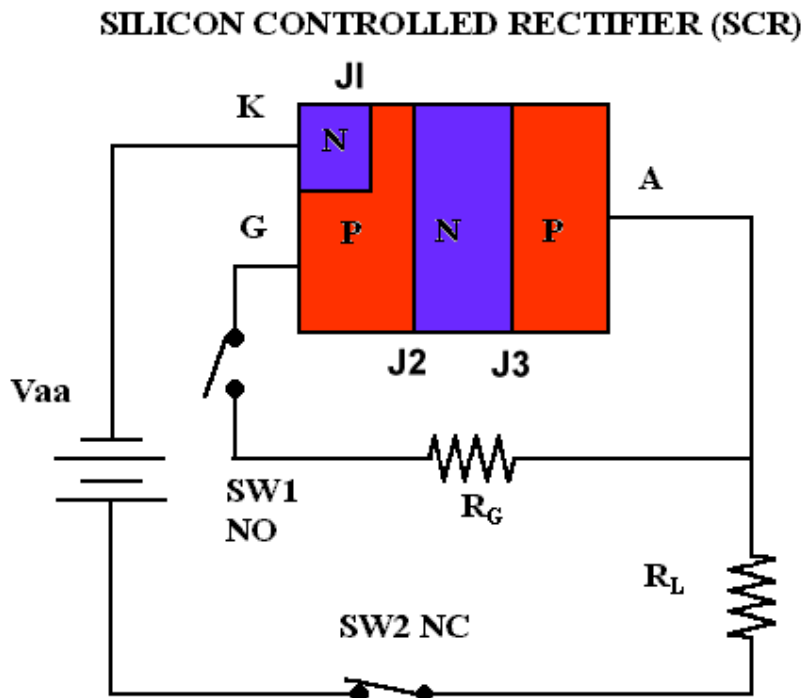
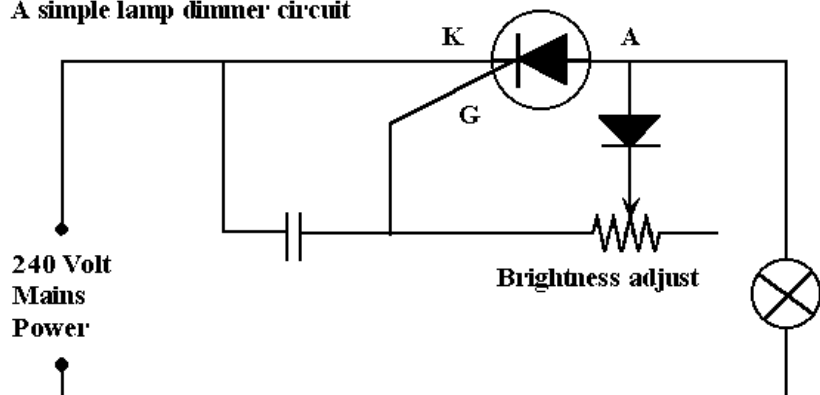


Figure 11.

So an SCR acts like a rectifier that can be switched ON by a control voltage on its gate and switched OFF by removing the Anode-Cathode supply voltage. SCR's can be made to handle medium to very high currents. A typical application could be a light dimmer or motor speed control circuit (though a TRIAC is better for the latter). The gate of the SCR is not switched (usually) with a mechanical switch, some of the main AC that the SCR is going to control is used to switch or 'fire' the SCR. An SCR is still a rectifier like any other rectifier, though when it starts to behave as a rectifier depends on the voltage applied to its gate. If you want to control both halves of an AC waveform you could use two SCR's. A device in one package which does this, and which has mostly replaced the SCR is called a **Triac**.

A simple lamp dimmer circuit



The gate of the SCR is not switched (usually) with a mechanical switch, some of the main AC that the SCR is going to control is used to switch or 'fire' the SCR. An SCR is still a rectifier like any other rectifier, though when it starts to behave as a rectifier depends on the voltage applied to its gate. If you want to control both halves of an AC waveform you could use two SCR's. A device in one package which does this, and which has mostly replaced the SCR is called a **Triac**.

Remember this: Once an SCR has been turned on the only way to turn it off is to reduce V_{aa} (the cathode-anode voltage) to a very low value or reduce it to zero (this has been asked in an exam – AOCP only).

TRIACS can be considered as back-to-back SCRs and perform as a full wave SCR for alternating current. An SCR is a unidirectional switch and TRIAC is a bidirectional switch. The gate control voltage of a TRIAC, unlike the SCR, can either be positive or negative in polarity.

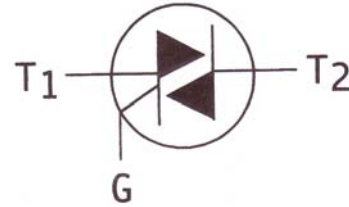
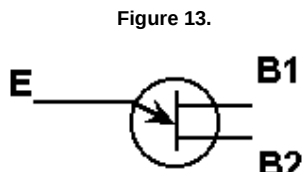


Figure 12

TRANSISTORS - UNIUNCTION TRANSISTOR

The only question about the unijunction transistor I have seen in the exam has been the identification of the symbol.



The unijunction transistor is a three terminal device. It has two bases B1 and B2 and an emitter. The unijunction does not have, to my knowledge, much of an application anymore, at least in radio. I believe all you really need to know is the symbol, however just in case, I will give you a quick tour of the UJT and explain its use.

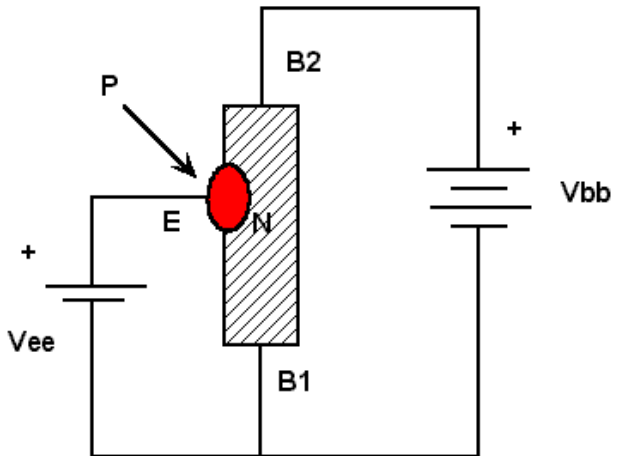
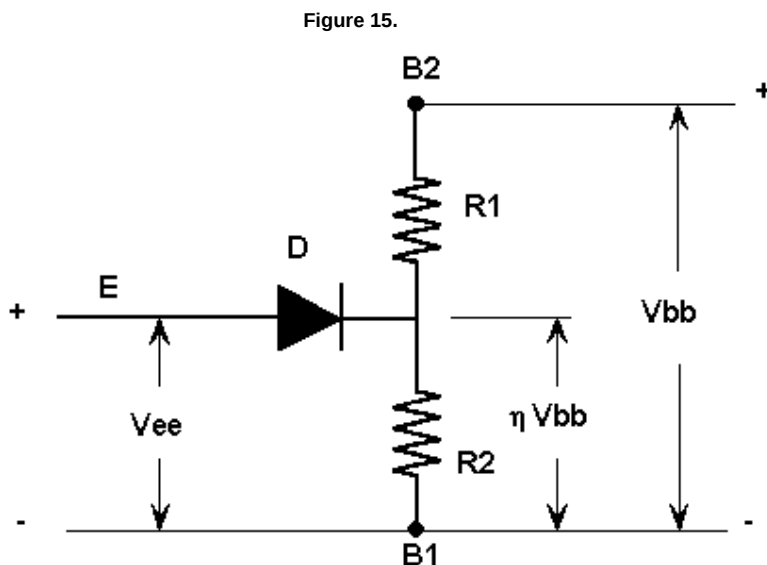


Figure 14.

The unijunction transistor is, in reality, a sort of double-based junction diode. Connections marked B1 and B2, are connected to the bottom and top, respectively, of a small bar of 'N' type silicon. A small region of 'P' type is fabricated to one side of the bar, as shown in figure 13. This 'P' region is called the emitter.

Let us disregard for a moment the emitter (P - region) and its circuit. A voltage V_{bb} is applied across both ends of the bar so that B2 is positive and B1 is negative. The bar acts like a voltage divider and since it is uniformly doped, the potential at any point along the bar is proportional to the location of that point between B2 and B1.



The positive potential of the point on the bar opposite the emitter is some definite fraction of V_{bb} , which we call η (the lower case Greek letter eta).

The equivalent circuit of a unijunction transistor is illustrated in figure 15. The silicon bar is represented by the voltage divider consisting of R_2 and R_1 .

Diode 'D' represents the junction diode formed by the emitter with the silicon bar. V_{bb} is the voltage applied between B2 and B1, and ηV_{bb} is the voltage drop in the silicon bar between the point opposite the emitter and B1.

If now the emitter is made positive by the application of V_{ee} , and this voltage is smaller than ηV_{bb} , the emitter diode is reverse biased. However, if V_{ee} is increased so the emitter becomes more positive than ηV_{bb} , then the junction becomes forward biased and a relatively large emitter current flows in the emitter circuit.

Forward biasing of the emitter with respect to B1 fills the 'N' region between the emitter and B1 with charge carriers (holes). In other words, the resistance between the emitter and B1 decreases dramatically. This resistance is represented by R_2 . So if you like, when the emitter is forward biased the resistance R_2 decreases dramatically. Also, the voltage across R_2 must decrease dramatically. Because emitter current flows, the current through R_2 increases.

(ADVANCED) CONCEPT OF NEGATIVE RESISTANCE

Note that the drop in voltage across R_2 coincides with a rise in current through it. Say that again! A drop in voltage coincided with a rise in current! Now this is not ohms law. Since, normally a rise in current is associated with a rise in voltage; we have, in the unijunction transistor, an example of **negative resistance**.

The unijunction gets its name from the fact that it contains only one P-N junction. It is used in special applications such as the **generator of a saw tooth voltage**.

It started with the BJT, JFET, then MOSFET

Then this.....

and this.....

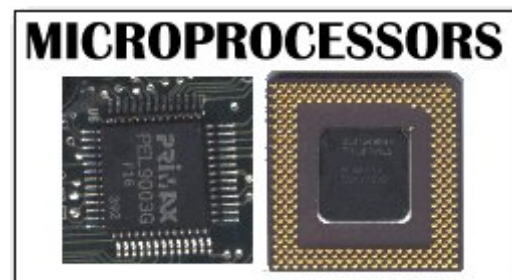
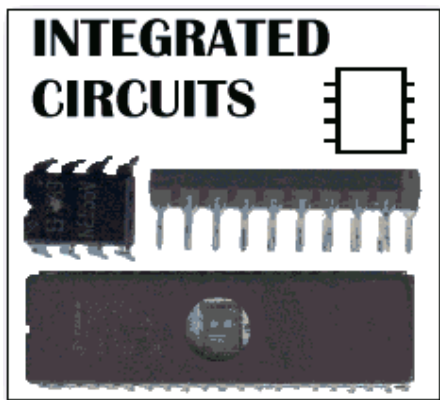


Figure 16 & 17.

*Note: For the AOCP(S) examination and syllabus the common emitter configuration only is required.

End of Reading 25.

Last revision: November 2013

Copyright © 1999-present Ron Bertrand

E-mail: admin@res.net.au

<http://www.res.net.au>

Free for non-commercial use with written permission