

Online Radio & Electronics Course

Reading 21

Ron Bertrand VK2DQ
<http://www.radioelectronicschool.net>

ELECTRON TUBES

One of the most significant developments of the early twentieth century was the invention of the electron tube. The British call it a valve. In the USA it is referred to as the "tube". Thomas Edison, the inventor of the electric light bulb (among other things) came extremely close to inventing the electron tube. When Edison was trying to develop a light bulb his most significant problem was to find a filament that did not burn out. He concluded early in his experiments that a filament would last longer if it were placed in a vacuum. Edison, being the entrepreneur that he was, would be kicking himself today if he knew just how close he had come to the discovery of the electron tube.

Edison tried literally hundreds of different types of filaments, without using much scientific method (Edison did not like scientists), more like a laborious trial and error approach. One of the things that frustrated Edison was that some of his filaments were giving off material which was deposited on the inside of the glass bulb causing the light to dim and overheat. He even inserted a metal electrode into some of his light bulbs to prevent these deposits from occurring. We now call this extra electrode an anode. It worked, but he dismissed the idea as impractical for a light bulb.

In amateur radio communications electron tubes are not used as much as they once were. The exception being many high power radio transmitters and amplifiers. Nevertheless, electron tubes are extensively used in industry and commercial radio communications. The television screen, computer monitor, and many other everyday devices still use electron tubes as their basic principle of operation. The examiner expects you to have a basic understanding of electron tube technology. Also, many of the principles of amplification, oscillation, rectification and the like, are the same with modern semiconductor devices as they are with electron tubes.

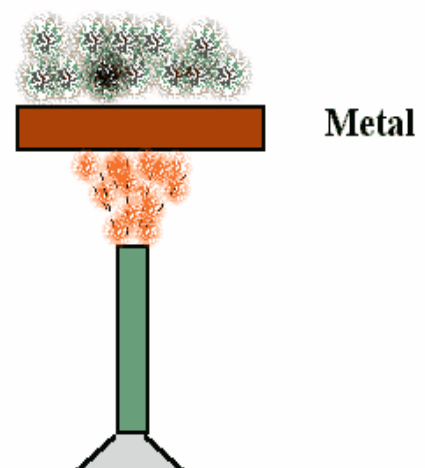
THERMIONIC EMISSION

When metals are heated to very high temperatures electrons are boiled off. These boiled off electrons have left the metal and form a cloud called a **space charge**. Since electrons have left the metal it becomes positive. The space charge, consisting of electrons, is negative. As long as the metal remains heated the space charge will exist, as it requires energy to keep the negative electrons away from the positive metal. Metals that are designed to enhance their ability to create a space charge are called **cathodes**. Since electrons have left the heated cathode, the cathode will be

Figure 1.

Thermionic Emission

Negative Space Charge



positive with respect to the space charge. The cathode will continually attract electrons back – however, electrons will be continually ejected by the thermionic emission.

CATHODES

A cathode is something capable of emitting electrons. For example, in a light bulb a piece of wire carrying an electric current becomes hot. If sufficient current flows, the wire becomes white-hot. Heat is 'really' the measure of the internal kinetic (moving) energy of atoms or molecules. Electrons in the outer orbits of the atoms move so rapidly due to increased temperature that some of the less tightly bound ones fly outward, and away from the wire into the surrounding space. This produces a cloud of free electrons around the wire, as long as it is hot enough. The wire has become a cathode. When the atoms in the filament wire lose electrons, the filament is left with a positive charge. The electrons expelled outwards are then attracted back toward these positively charged atoms in the filament (cathode); some do make it back, only to be expelled again. This results in a constant in and out, back and forth movement of electrons, in the area directly surrounding the cathode. This creation of the space charge around a hot filament is called the Edison effect or **thermionic emission**.

The space charge has no practical use in a light bulb. Filaments designed for the purposes of creating a space charge are called cathodes. There are a number of materials used for cathodes, however, by far the most popular is **thoriated tungsten**. Thorium is an excellent electron emitter, but this material is not strong enough for practical use. Mixing thorium with the more rugged tungsten produces a thoriated-tungsten cathode that will permit many electrons to be emitted at a relatively low temperature (1600 degrees centigrade). Cathodes of this type are used in many radio transmitting tubes.

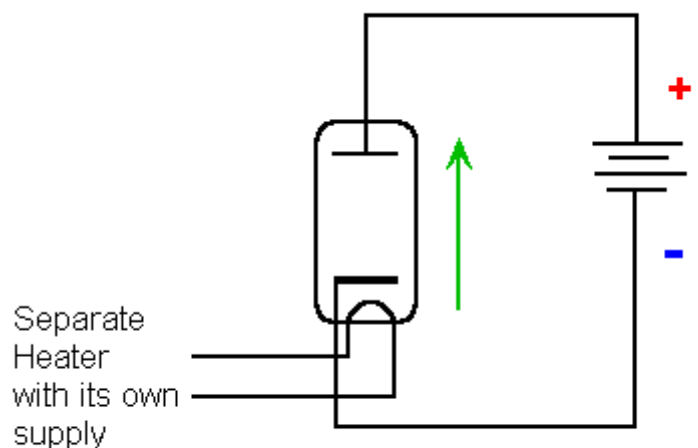
To enhance the space charge creating properties of the cathode, it is placed in a glass envelope and the air is pumped out. To improve the vacuum further, the entire assembly is heated to force out any additional gas molecules in the metal electrodes. Then a small amount of magnesium (**called the getter**) placed inside the envelope during manufacture, is ignited. The resultant chemical action of the vapourising magnesium with the gases released from the metal, removes the final traces of gas in the tube. The by-products of the vapourising magnesium condense on the inside of the electron tube forming the silvery coating commonly seen on electron tubes.

Figure 2.

INDIRECTLY AND DIRECTLY HEATED CATHODES

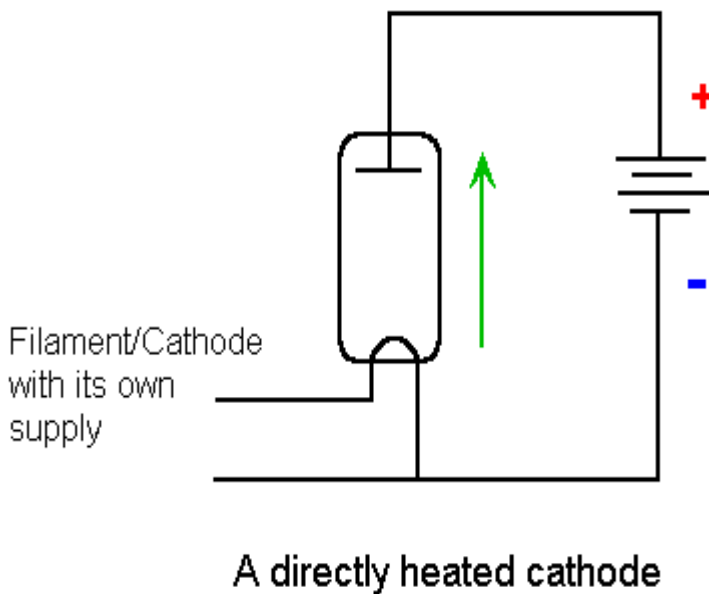
The cathode of an electron tube can be heated by a separate filament, in which case the filament is called a **heater**. Such an arrangement is called an **indirectly heated cathode**. Alternatively, the heater itself can be treated with thorium to act as a cathode. This arrangement is called a **directly heated cathode**.

In figure 2 showing an indirectly heated cathode, the heater is a third element connected to its own power supply.



An indirectly heated cathode

Figure 3.

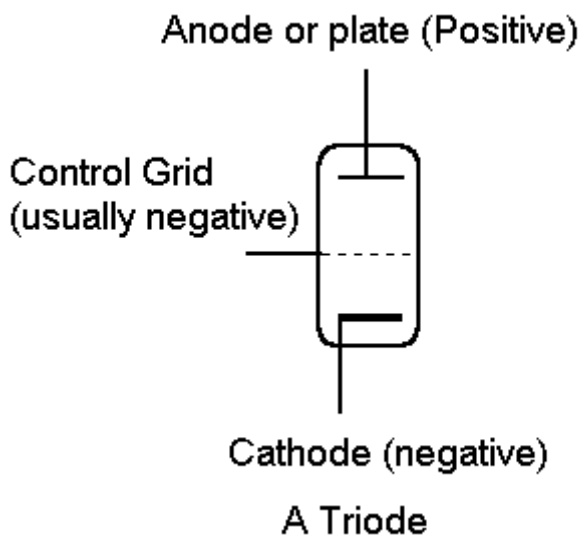


If a second electrode is inserted into the electron tube and a positive voltage is applied to it with respect to the cathode, electrons will flow from the space charge, through the vacuum, to the second electrode. This second electrode is called an **anode or plate**. This electron tube is called a diode (di = two). The anode does not have a space charge and therefore electron flow is only possible from the cathode to anode. So here we have an electric current without a conductor. Which is why we earlier in this course correctly defined an electric current as the ordered movement of electrons, and left the 'through a conductor' part out.

It should now be obvious that electric current does not have to occur in a conductor. In the last reading we discussed power supplies and rectification. The diode electron tube just described can be used in any of the power supply circuits. A diode electron tube will always conduct if its anode (also called a plate) is positive with respect to the cathode, and a space charge exists. Electron flow from anode to cathode is not possible.

THE TRIODE

Figure 4.



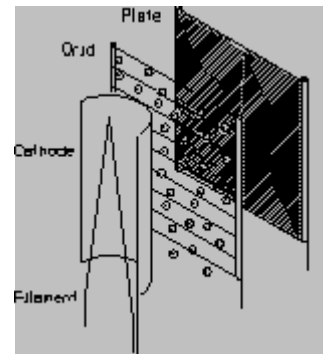
A circuit diagram of a Triode (tri=3) is shown in figure 4. In a diode, a cathode and anode are required to enable electrons to flow. In addition, a triode has a third electrode called the **control grid**. The grid is a fine metal wire, usually nickel, molybdenum, or iron. The anode and cathode are both solid electrodes. The control grid in the triode is a **fine spiral wire** placed between the cathode and the anode. All electrons attracted to the plate from the cathode go through the openings in the control grid. The grid is connected to the base of the electrode so that voltage can be applied to it. Imagine if you will a Triode operating without voltage applied to

its grid, it will act just like a diode. If we make the grid slightly **negative with respect to the cathode** the negative grid will inhibit (restrict) the flow of electrons from cathode to anode (like charges repel). The voltage between cathode and anode may be hundreds of volts. On the other hand, a **small negative voltage** on the control grid will have a substantial effect on the flow of current from the cathode to the anode. The significance of this is that we are able to control a large cathode-anode current with a very small negative grid

voltage. Making the grid more negative will decrease the anode current; making the grid less negative will increase the anode current.

Figure 5.

A pictorial diagram of the construction of a triode is shown in figure 5.



In passing, this is why the electron tube got its name - the valve. Think about how a water valve works. A minor input of energy from your hand can control a huge pressure of water. Likewise, a small voltage on the control grid can control a large anode/plate current.

Under normal circumstances the control grid is negative. Because of this, it does not draw any electrons from the cathode. It is the **negative voltage alone** on the control grid, that is able to control what can be a substantially large current from the cathode to anode (voltage without current is no power $P=EI$).

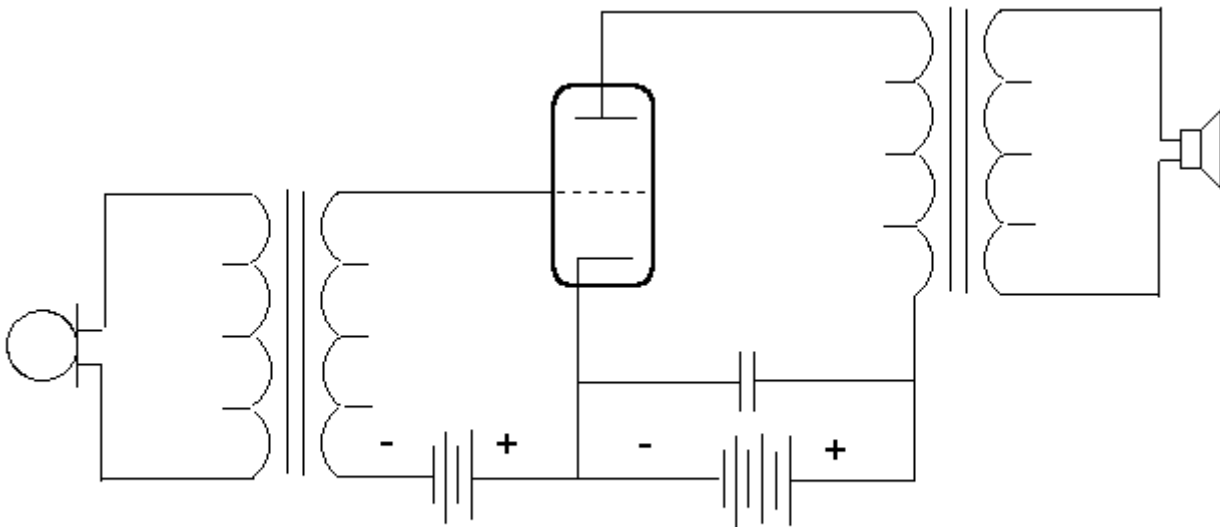
The control grid is placed **closer to the negatively charged cathode**, and farther from the highly positively charged anode. This placement allows it to function more effectively than if it were simply midway between the cathode and anode.

The cathode-anode circuit of the electron tube is the output circuit. The control grid-cathode circuit of the electron tube is the input circuit. No current flows in the control-grid cathode of the electron tube (there are special cases where current does flow but it is still insignificant). On the other hand, the voltage and current in the output circuit of the electron tube is **substantial**. And as we learnt, the product of current and voltage is power (watts). Therefore, with the triode we have the means to control large amounts of power by simply applying a small negative voltage to the control grid.

We could for example, connect a microphone to the control-grid of an electron tube, and use the very small voltage from this microphone to change the negative potential on the control grid, that in turn would control the anode-cathode current. The microphone produces small voltages corresponding to sound waves striking its diaphragm. These small voltages are then applied to the control-grid, which in turn controls the larger anode-cathode current. The anode-cathode current will contain all of the information (intelligence) that was initially fed into the microphone.

What we have just described is amplification. The Triode was the first electronic device ever to be able to produce amplification.

Figure 6.



A simple triode audio amplifier (PA System)

The triode amplifier in figure 6 uses a transformer to match the impedance of the microphone to the input of the triode, and another transformer to match the output impedance of the triode to the impedance of the speaker. The small battery on the left makes the control-grid negative at all times – this is called '**negative bias**'. The voice signal from the microphone on the secondary side of the transformer is AC. The AC will be superimposed on top of the DC bias to produce a varying DC negative voltage on the control-grid. The capacitor across the larger cathode-anode supply is to bypass the signal that has been amplified around the battery. In a practical amplifier, batteries would not be used.

Can you see how many devices use Faraday's law in this circuit?

DISADVANTAGE OF THE TRIODE

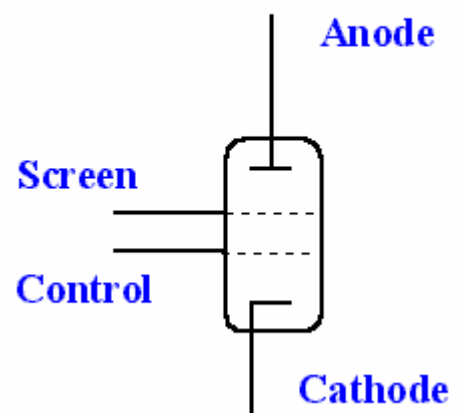
When we discussed capacitance earlier in this course we found that any two conductors separated by an insulator has capacitance. The problem with the triode is the capacitance between the **control-grid and anode**. (more on why this capacitance is bad shortly). We also learned in our study of capacitance that capacitors in series decrease the overall circuit capacitance. Inserting another grid called the **screen grid** between the control grid and the anode, decreases the capacitance between the control grid and the anode.

Why is the control-grid to anode capacitance bad?

The anode is in the output circuit of the triode. The control grid is in the input circuit of the triode. The unwanted capacitance we're talking about is between the control-grid and anode. This unwanted capacitance may well be inside the electron tube, however, it may as well be a physical capacitor connected externally from the anode to the control grid. This means that there is a connection through this capacitance from

Figure 7.

Tetrode



the output of the triode to the input of the triode, and this is undesirable. This is unwanted capacitance and it results in unwanted **feedback**.

You will be well aware of the problems with feedback from, say, a public address system. If sound from the speakers in a public address system get into the microphone of the same system, you will have feedback, and it will oscillate (squeal).

Now, if we want the triode to oscillate that would be just fine. However, if we want it to amplify, then this oscillation (because of feedback) is a problem. The inclusion of a screen grid between the control-grid and anode **reduces the control-grid to anode capacitance** and improves the stability of the triode as an amplifier. A triode with a screen grid is no longer a triode, it is called a Tetrode.

TETRODES

The screen grid of tetrodes (tet=four) is not used to control the plate current, but has a steady **positive** DC voltage on it to help accelerate electrons to be collected by the anode. The path for current inside a tetrode is from the cathode through the control grid, and through the **spaces** in the screen grid, to the anode.

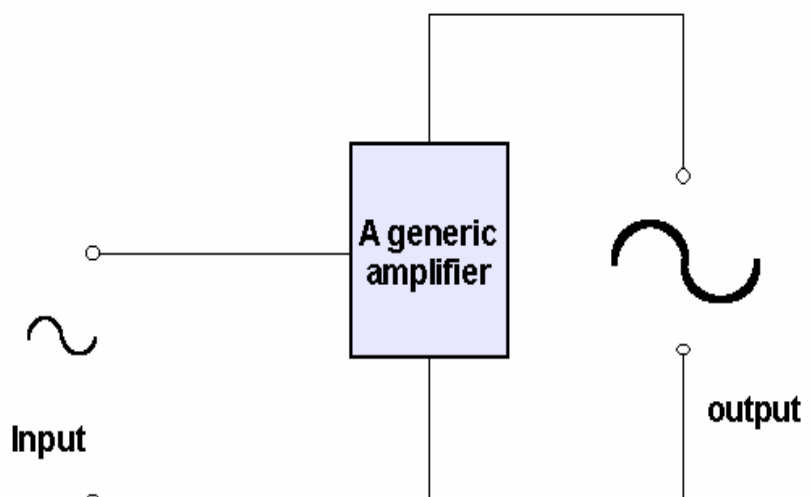
Since the screen grid is positive, it will collect some electrons. Most of the electrons arriving at the screen grid have enough momentum to continue to the anode. Electrons arriving at the screen grid pass through it, attracted by the much higher positive potential on the anode. The screen grid current is wasted current, since it is not used in the output circuit. The benefits of the screen grid (reduced control-grid to anode capacitance) outweigh this disadvantage.

MORE ABOUT THE TRIODE AND FEEDBACK

Even though we are talking about electron tubes in this reading, much of what you learn here can be carried over to other devices that amplify, such as transistors.

The diagram in figure 8 shows a generic amplifier. The signal to be amplified is applied to the input and a larger signal appears at the output.

All amplifiers work this way regardless of what the amplifying device is. The power for the amplified output is provided by the power supply that is not shown.



An amplifier with no feedback

Figure 8.

Figure 9 below shows a little more detail of a generic amplifier. Most amplifiers have, as a **natural consequence** of their construction, a small amount of **unwanted capacitance between their output and their input**. I have shown this unwanted capacitance on the amplifier below in red. Remember, this capacitance is not a real physical capacitor, rather an unwanted capacitance – I will explain how it is produced in the electron tube shortly. The consequence of this unwanted capacitance in any amplifier, between output and input, is that some of the output signal can return to the input. If you like, some of the output signal ‘leaks’ back to the input. This is called feedback. In the case of unwanted capacitance between output and input as shown below, the feedback is called **positive feedback**.

Positive feedback means the output signal is fed back to the input in such a way that it adds to, or is in phase with, the input signal. I have shown the feedback signal as a small sine wave. Can you see how this sine wave is in phase with the input signal? Positive feedback *adds* to the input signal – so it’s like taking some of the output and amplifying it all over again. This is called **regeneration**. The terms regeneration and positive feedback can be used interchangeably.

If there is enough positive feedback the amplifier will stop being an amplifier, and it will oscillate. Like a public address system where the microphone is placed too close to the speaker and it howls. An amplifier with too much positive feedback will oscillate (like the howl of a public address) but you may not hear it at all – it could oscillate on a radio frequency above human hearing. It certainly won’t be amplifying as it should. The amplifier is said to be **self-oscillating**, and amplifiers are supposed to amplify not self-oscillate. We will cover more of this in detail later. The point I am trying to make here is that **unwanted positive feedback is a bad thing and we must prevent it from happening**.

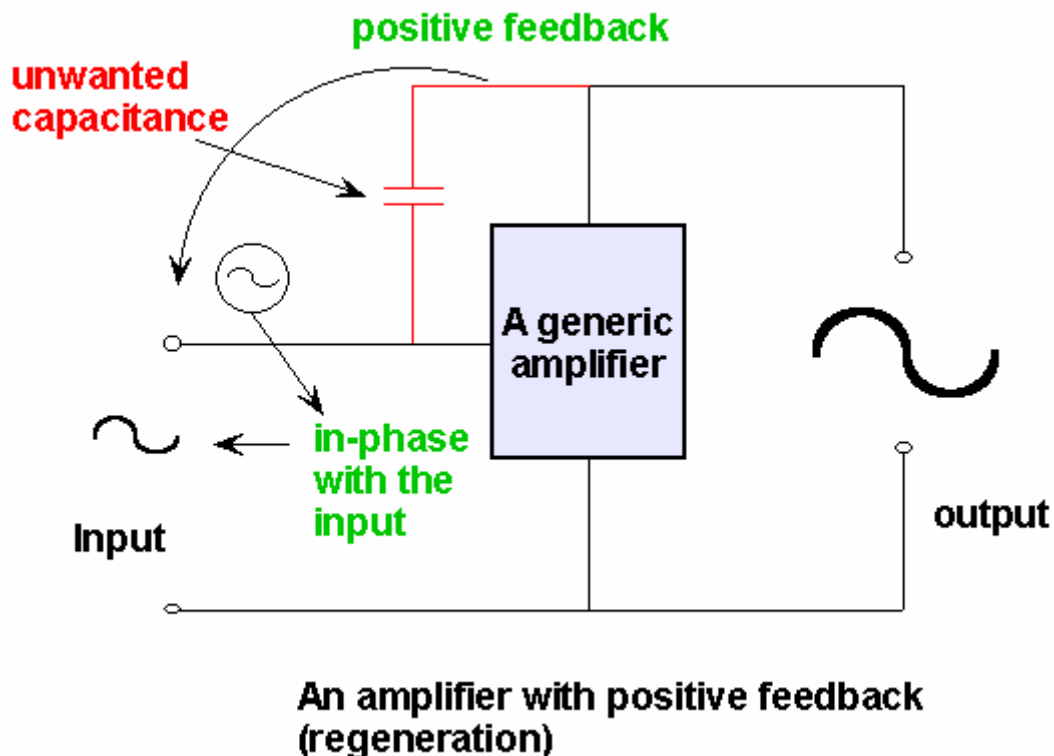


Figure 9.

THE TRIODE AMPLIFIER

Just to recap for a moment look at a basic triode circuit in figure 10.

The input signal is applied to the control grid and cathode. The output signal is taken from the anode and the cathode. (The cathode is common to both input and output). If any of the output signal in the anode were able to get back to the control grid, then this would be positive feedback and undesirable.

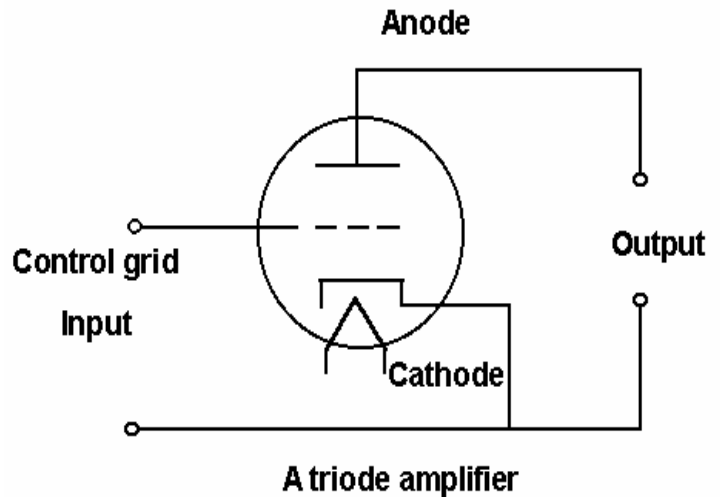


Figure 10.

UNWANTED CAPACITANCE IN THE TRIODE

We have an unwanted capacitance formed in a triode between the control grid and the anode. It is there because everything is in place to form a capacitor. The control grid and the anode are the two plates and the vacuum between them is the dielectric. There are other capacitances in the triode too. Between the cathode and the control-grid for example – and many others. However, it is the capacitance between the anode and the control grid that will allow positive feedback to occur. The reactance of a capacitor goes down with increasing frequency – so the higher the frequency a triode (and other amplifying devices) are operated at, the more positive feedback there is. Remember, enough positive feedback will cause the amplifier to oscillate and stop doing its proper job of amplifying. This is where the tetrode comes in.

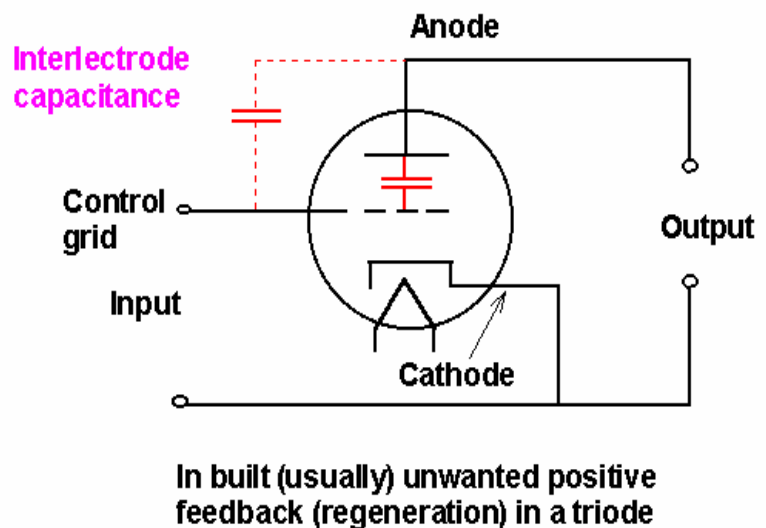


Figure 11.

HOW THE TETRODE HELPS

The tetrode breaks the capacitance between the control-grid and the anode into **two capacitors in series**. This is shown in the diagram of figure 12. Now what happens when you connect capacitors in series? Does the total capacitance go down or up?

I hope you said the total capacitance goes **down**. This is the **primary** function of the screen grid. The insertion of the screen grid reduces the unwanted interelectrode capacitance between the control-

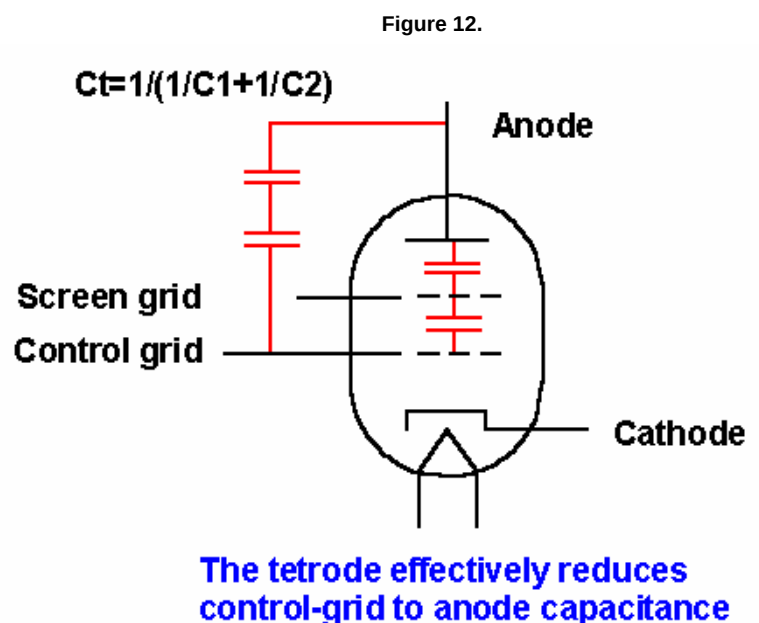


Figure 12.

grid and the anode, and in doing so reduces the likelihood of unwanted positive feedback, enabling the electron tube to be used at higher frequencies.

Just to reinforce things – this unwanted capacitance is not a real physical capacitor that anyone has deliberately added to the electron tube – it is the nature of the electron tube itself to have this unwanted capacitance.

When worse comes to worse, and we have positive feedback which is troublesome - the only thing we can do is to deliberately use components in the circuit to apply an equal amount of negative feedback to the electron tube. Negative feedback is out-of-phase feedback. Negative feedback is also called 'degeneration'. The process of deliberately applying negative feedback to cancel out unwanted positive feedback is called **neutralisation**. We will talk more about neutralisation as it is an important feature of amplifiers used at radio frequencies.

PENTODES

Fast moving electrons (relatively speaking), in an electron tube, may strike the anode with considerable force. Some of the electrons striking the anode will collide with electrons in the anode material and **bounce back**, or electrons arriving at the anode may dislodge electrons in the anode material **and** cause them to be **emitted from the anode**. This phenomenon is called **secondary emission**. **Secondary emission is undesirable**. Electrons participating in secondary emission are not performing any useful function, and in fact can build up a small cloud of electrons somewhat similar to the space charge around the cathode, though nowhere near the same size. This secondary space charge around the anode will inhibit or restrict the flow of electrons from cathode to anode.

Pentodes (pent = 5) have the same construction as tetrodes but with the addition of another grid, called the **suppressor grid**, in the space between the screen grid and anode. The suppressor grid has a negative voltage placed on it with respect to the cathode. Now you might think that a negative voltage on the suppressor grid would also inhibit cathode-to-anode current. In fact it does, but not to the extent of secondary emission. The negative voltage on the suppressor grid does not appreciably reduce the anode-cathode current. Also, by the time electrons come close to the anode they have built up a significant amount of momentum that will carry them through the negative electric field of the suppressor grid.

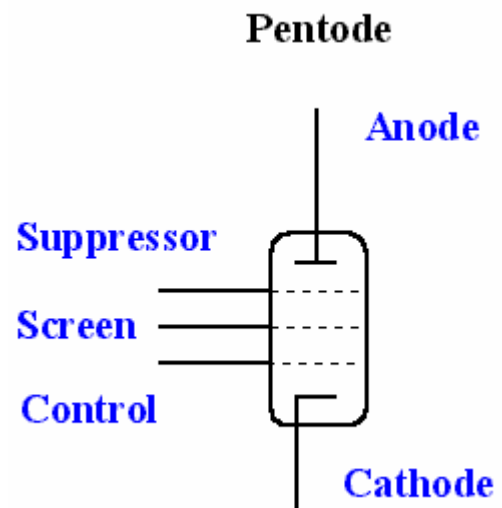


Figure 13.

The secondary electrons, the ones created by secondary emission, do not have a great deal of energy and **are repelled back to the anode by the negative potential on the suppressor grid**. The pentode has the **additional** advantage of having still less capacitance between the control grid and the anode. The suppressor grid, like the screen grid, acts like another series capacitor between the control grid and anode reducing the unwanted capacitance even further.

X-RAYS

Just as a matter of interest, if the cathode current is made very high by an extremely large voltage on the anode relative to the cathode, and if secondary emission is encouraged, a by-product of this is that X-ray radiation will be emitted from the anode. This was how the X-ray machine was developed. In an X-ray machine, the anode is specially shaped at an angle of 45 degrees to the approaching electron stream to enhance the emission of X-ray radiation.

For exam purposes, it is important for you to know the purpose of all the electrodes in electron tubes and their relative electric potential with respect to the cathode.

We have described all the basic types of electron tubes. There are others, and these are only modifications of the types of tubes that we have just discussed. An electron tube may have two diodes within one envelope. The two diodes may share one cathode, or they may have separate cathodes - this is called a duo-diode. Other types are duo-triodes and triode-pentodes.

SOFT TUBE

An electron tube can become faulty and have unwanted "gas" or "air" in its vacuum. This means the vacuum is not as "hard" as it should be. Such a fault with electron tubes is called a **gassy** or **soft** tube.

MICROPHONICS

Another problem which can occur with electron tubes is microphonics. An electron tube is a mechanical device as well as an electronic device, and if it is tapped with a screwdriver, or flicked with a finger, it will vibrate. Vibration of the grid will modify the cathode-to-anode current and this vibration will 'modulate' the output.

Modulate or modulation has not been discussed yet. Modulation means imposing the information of something onto something else. In the discussion above the vibrating grids would transfer this vibration 'message' to the output current.

If an audio (sound) system was constructed from electron tubes, and some of the audio output sound wave energy from that system was able to reach the electron tubes and caused them to vibrate, then the system would produce microphonics, or perhaps even feedback. In some old 'wireless' sets, if you hit them on the cabinet, you would hear a 'thud' sound from the speaker.

By the way, don't you just love that term 'wire-less'? I do, but then perhaps I am a bit strange!

ANOTHER LOOK AT FEEDBACK

Feedback is not a problem specific to electron tubes. Feedback is a sub-topic of its own. However, since we have mentioned it in this reading, we may as well discuss it a little further.

Feedback is only a problem if you don't want it. If you want an electron tube to oscillate, then you may take advantage of the capacitance between the control grid and the anode.

Alternatively, you may use external components, usually a capacitor or an inductor, to take some of the electron tubes output and feed-it-back to its input.

An oscillator is a device for producing electronically, AC voltage. The frequency of the AC voltage may be in the audio range or high up in the microwave region.

There are two types of feedback, called positive and negative - this has nothing to do with polarity.

When some of the output of an amplifier is fed back to its input, in such a way that the voltage is **in phase** with the input voltage, this is called **positive feedback**. This is the type of feedback we require if we want oscillation to take place.

If some of the output of an amplifier is fed back to its input so that it is **out of phase** with the input, this is called **negative feedback**. Negative feedback reduces the amplification (gain) of an amplifier, however it also has benefits. A perfect amplifier (no such thing) would do nothing to the input signal except increase its amplitude. In practice, all amplifiers produce some distortion. The distortion simply means that what we are getting out of the amplifier is a larger signal, however, there is **always** unwanted information in the output created by the amplifier itself. Negative feedback can actually reduce the distortion of an amplifier. The ability of an amplifier to amplify with **minimum** distortion is called **linearity**. Negative feedback does improve the linearity of an amplifier.

Hey! Remember our old friend **direct proportion**? (perhaps an enemy!). If an amplifier's output is directly proportional to its input, it is linear.

THINGS TO DO

It is not very hard at all to get your hands on an old electron tube. If you can please do so. Place it in a cloth bag or towel and smash it, and have a look at the construction. You will see that the heater is in the centre with a cylindrical cathode around it; cylindrical spirals around the cathode for the various grids, and finally a cylindrical anode. Any repair shop will have lots of old ones lying around. It really does pay to have a look.

End of Reading 21.

Last revision: March 2007

Copyright © 1999-2002 Ron Bertrand

E-mail: manager@radioelectronicschool.net

<http://www.radioelectronicschool.net>

Free for non-commercial use with permission