

Reading 41

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

MISCELLANEOUS

The purpose of this reading is to catch anything that may have slipped through the previous forty readings or just does not fit anywhere else. In time, the material in this reading will be incorporated into earlier readings. In a live 'upfront' class things don't tend to slip through, as someone will usually ask a question prompting an answer for an area that may otherwise have been missed. This reading will catch those things that were missed. Please be aware that some of the material here may have already been moved to earlier readings and therefore be duplicated. This has been done to ensure that all people have all the material. This reading may be re-posted while revision is going on with the online course. If this happens you will be advised.

POWER SUPPLIES – PASS TRANSISTORS

You may recall when we studied zener regulators in power supplies that the zener diode cannot handle a significant amount of current, usually no more than a few hundred milliamps. Also, a zener diode is not a perfect regulator. There is a way of improving the current handling capabilities of a zener regulator circuit and improve regulation at the same time. This is done by adding a **pass transistor**.

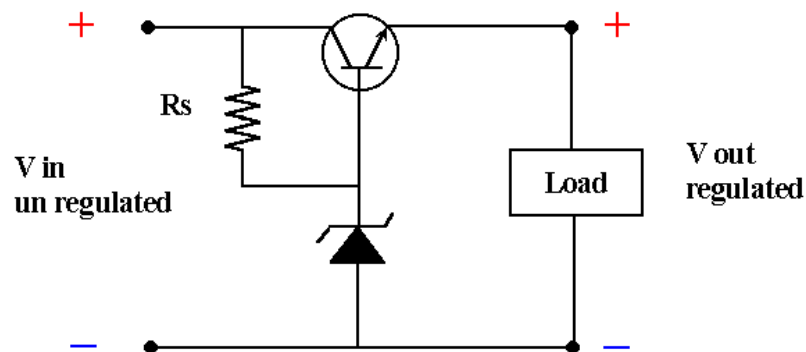


Figure 1.

Figure 1 is the schematic diagram of a zener regulator with a pass transistor added. You will recognise R_s as the normal current limiting resistance that prevents the zener from being destroyed. The zener is reverse biased, as shown by a positive voltage being applied to its cathode.

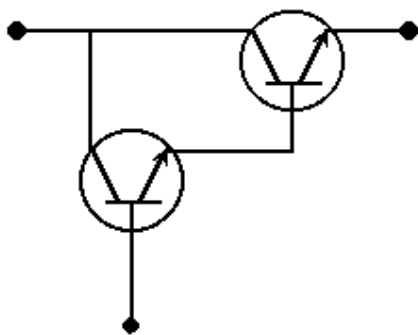
The zener regulated voltage drives the base of a transistor connected as an emitter follower (common collector). The output of the emitter follower goes to the load. Because of the current gain (β) of the emitter follower, any change in load current is reduced by a factor of β when viewed from the base. This means voltage regulation is improved by a factor of β , which is a significant improvement.

Use of the pass transistor increases the maximum load current the regulator can handle. You need to be able to identify this circuit as a zener regulator with a pass transistor. The purpose again – to improve regulation and the current handling capabilities of the zener.

As a matter of interest, a pass transistor can be added to a three terminal voltage regulator in much the same way and for the same reasons.

THE DARLINGTON PAIR

Two transistors can be connected in a configuration known as the Darlington pair. As far as I know you do not even have to know what a Darlington pair is. However, they have appeared in circuits in AOCP exams and it helps to eliminate wrong answers if you know what they are.



A Darlington pair using npn transistors

The higher the β the higher the input impedance of an emitter follower. One way to increase current gain or β is with a Darlington pair, two transistors of the one type connected as shown in figure 2.

The overall current gain is equal to the product of the two individual betas.

Transistor manufacturers can put a Darlington pair inside a single transistor package. Notice how the external connections to the two transistors are a base, emitter and collector.

Figure 2.

One application you may see of a Darlington pair (particularly in the exam) is its use as an improved pass transistor configuration, as discussed in the previous section.

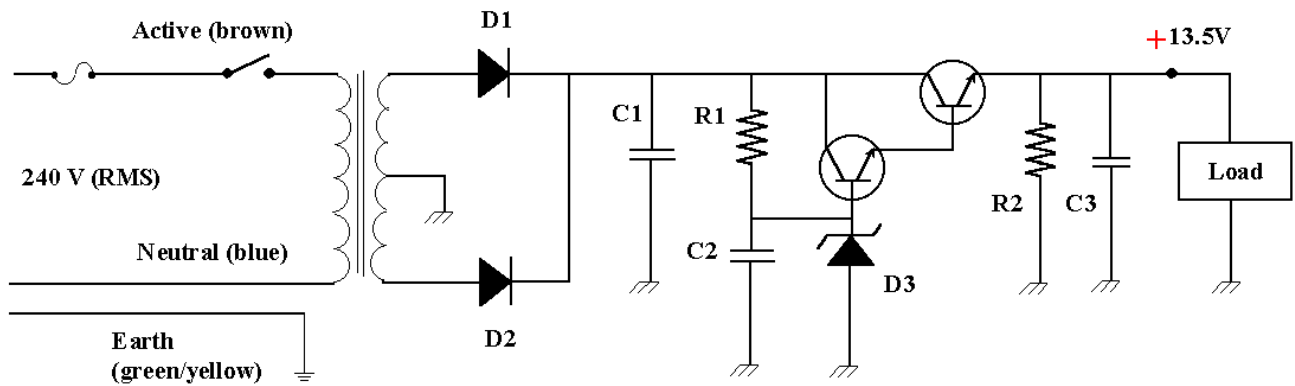
A Darlington pair can also be made with two pnp transistors.

In the next section we will discuss a circuit you may see using Darlington pair pass transistors.

A POWER SUPPLY WITH PASS TRANSISTORS

All of the regulator circuits we have looked at are **series regulators**. The pass transistors are in series with the load. Series regulators are more efficient. The shunt type regulator (which is in shunt or parallel with the load) is rarely used as they dissipate too much power.

Figure 3 is a full wave power supply. It uses two diodes, therefore it must have a centre tapped transformer. C1 and C2 are both filter capacitors, probably around 1000 microfarads and therefore electrolytics. R1 is the current limiting resistor for the zener regulator D3.



A full wave power supply using pass transistors

Figure 3.

The zener regulator has a Darlington pair as the pass transistors. The Darlington pair does a superior job compared to a single transistor, improving the regulation and the current handling abilities of this power supply significantly. Read the section on pass transistors and the Darlington pair again if you need to.

R2 supplies a light load on the power supply should the 'real' load not be connected. Its value would be about 1k Ω . C3 is an RF bypass capacitor about 0.01 μ F.

If you can answer the following questions you will be right with this and many other circuits:

Is the zener diode reverse biased? Yes, the positive rail (line) is at the top of the circuit. See if you can deduce this from the power supply diodes and how they would conduct and charge C1.

What is the purpose of R1? Current limiting.

Are the transistors connected the right way - in other words, does current flow against the arrow into the emitter? Yes, the bottom of the load is negative. Current flows up through the load and into the emitter of the top right transistor.

If you need to, revise the operation of the full wave power supply using two diodes only.

INTERFERENCE - CROSS-MODULATION

Cross modulation is the **transfer of intelligence from an unwanted strong signal to a wanted weak one**. Say a receiver is tuned to a signal on 3.9 MHz and a strong unwanted signal is at 3.8 MHz. The modulation of the second signal may be imposed on the wanted signal, even though the second signal is well outside the IF pass band of the receiver.

RECEIVERS - DESENSITISATION AND BLOCKING

When a receiver is tuned to a weak signal with another strong signal close to the tuned frequency, an apparent decrease in receiver gain or sensitivity may occur. This loss of gain is called desensitisation or blocking. It commonly occurs when the unwanted signal voltage is sufficient to overcome the operating bias of an amplifier or mixer stage.

Desensitisation and blocking are different degrees of the same thing. Blocking is when desensitisation is so bad the receiver is 'blocked' or totally deaf, so it receives nothing.

SUPERHETERODYNE RECEIVER - IMAGE INTERFERENCE

There are always two signal frequencies which will combine with a given frequency to produce the same difference frequency. This means there are two signals that can enter a superheterodyne receiver and be converted to the IF.

The figure 4 block diagram shows a superhet receiver tuned to a desired frequency of 1200 KHz. This 1200 KHz desired signal is converted to the IF by mixing with the local oscillator of 1655 KHz as shown.

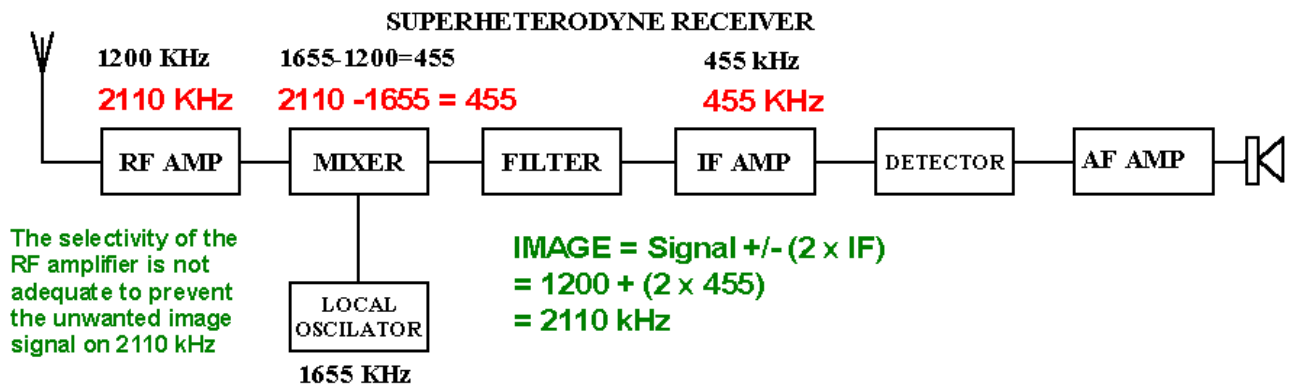


Figure 4.

Now an unwanted signal on 2110 KHz can also mix with the local oscillator and be converted to the same IF. If a signal were present on 2110 KHz then this receiver would most probably pick it up. The unwanted signal is called the **image frequency** and the type of interference known as **image interference**.

It is normal for most superhet receivers to have the local oscillator on a frequency above the signal frequency as shown. **The image frequency for any superhet is twice the IF away from the wanted signal.**

Improving the selectivity of the RF amplifier will help with image interference, however a better way is to move to a higher IF. An IF of 10.7 MHz will place the image 21.4 MHz away from the desired signal. The RF amplifier is capable of rejecting a signal so far away. The problem is that we rely on the IF for selectivity and gain, and at an IF of 10.7 MHz it's a lot harder to obtain the high selectivity and gain. **Therefore an ideal receiver will have two IF's - a high first IF for good image rejection, and a further low IF for improved gain and selectivity.**

CONVERSION EFFICIENCY (MIXER)

Conversion efficiency is simply a fancy name for how efficient a frequency converter is, or more precisely, the efficiency of the mixer in the frequency converter. A frequency changer or converter is just a mixer and its associated oscillator. Efficiency percentage is a function of: $(\text{power out}) / (\text{power in}) \times 100$. All you need to remember is that conversion efficiency is how efficient a frequency converter does its job with the signals fed to it.

PUSH-PULL AMPLIFIER CANCELS EVEN HARMONICS

A push-pull amplifier has the advantage of largely cancelling any **even** (2^{nd} , 4^{th} , 6^{th} etc.) harmonics. This is a great advantage for both AF and RF amplifiers. For the same reason, a push-pull amplifier cannot be used as a frequency doubler or quadrupler.

CRYSTAL - SERIES RESONANCE

We have learnt that a crystal actually has two resonant frequencies: a parallel resonant frequency where it behaves like a very high 'Q' parallel LC circuit; a series resonant frequency where it behaves like a very high 'Q' series LC circuit.

These two frequencies may be as close as 10 KHz. The lowest resonant frequency of the crystal is the series resonant frequency.

The point...

An **overtone oscillator** is one that operates on a multiple of the series resonant frequency of a crystal. This may also be put another way. An overtone oscillator is one that operates on a multiple of the lowest frequency to which the crystal is cut. These are of course the same thing, as the lowest frequency to which the crystal is cut is the series resonant frequency.

TRANSISTORS - UNIJUNCTION TRANSISTOR

I did not cover the unijunction in the readings. The only question I have ever seen asked is to identify the symbol (figure 5).

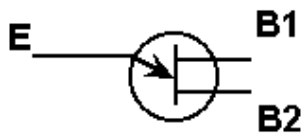


Figure 5.

The unijunction transistor (UJT) is a three terminal device. It has two bases B1 and B2 and an emitter. The unijunction does not appear to have 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.

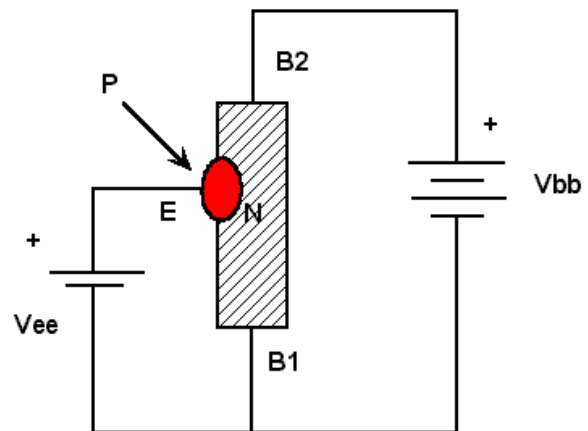


Figure 6.

Refer to figure 6. 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 P type region is fabricated to one side of the bar as shown in figure 6. 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 7. 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 B_2 and B_1 , and ηV_{bb} is the voltage drop in the silicon bar between the point opposite the emitter and B_1 .

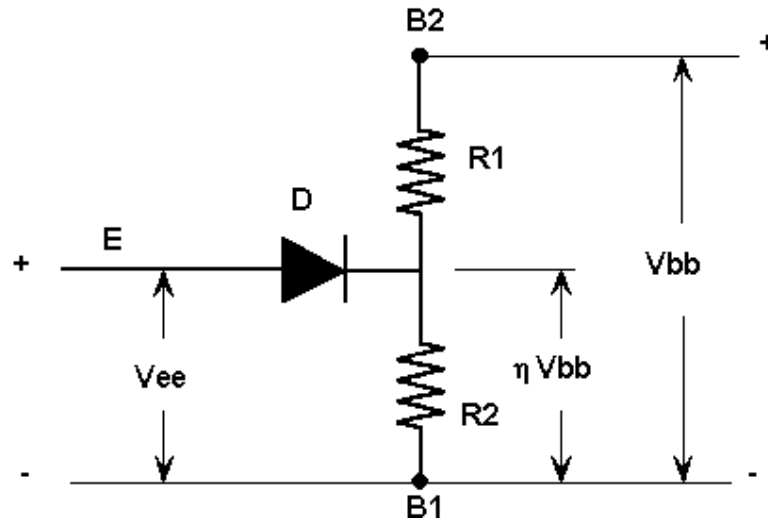


Figure 7.

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

Forward biasing of the emitter with respect to B_1 fills the N region between the emitter and B_1 with charge carriers (holes). In other words, the resistance between the emitter and B_1 decreases dramatically. This resistance is represented by R_1 . So to put it another way, when the emitter is forward biased the resistance R_1 decreases dramatically. As a result, the voltage across R_1 must decrease dramatically. Therefore, because emitter current flows, the current through R_1 increases.

(ADVANCED) CONCEPT OF NEGATIVE RESISTANCE

Note that the drop in voltage across R_1 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 Ohm's law. Since a rise in current is normally 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 a **generator of a sawtooth voltage**.

ANTENNAS - GAMMA MATCH

A dipole which forms the driven element of a Yagi antenna has an impedance too low for 50 ohm coaxial line. The exact impedance of the driven element will be determined primarily by the number of elements on the Yagi. Let's just assume it's 15 ohms.

Now if the centre of the driven element is 15 ohms, and the ends of a dipole are always a high impedance - perhaps as high as 2000 ohms, then if we move our feedline connection a little outwards from the centre of the dipole we will find an increase in impedance - in fact we will find 50 ohms. That's exactly what a gamma match does (figure 8).

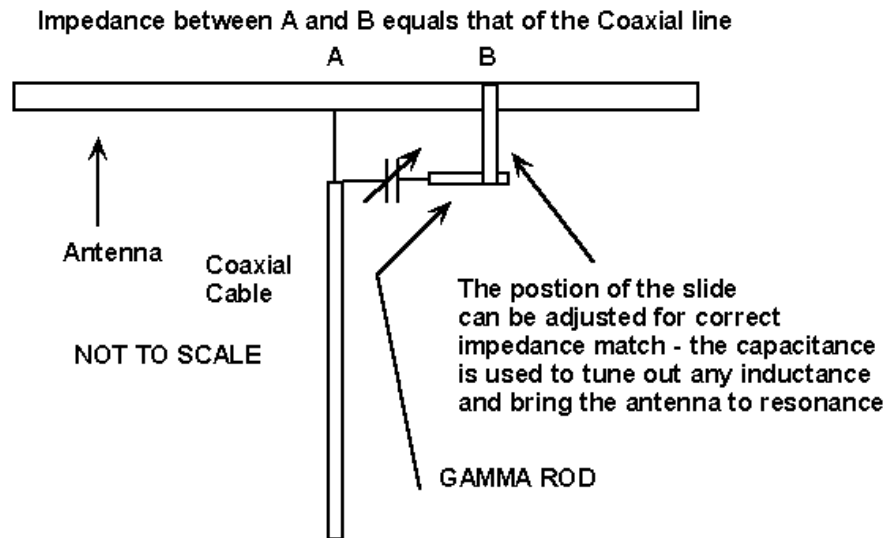


Figure 8.

Because the impedance at point B will have some reactance (the dipole is only resistive at the centre), the capacitor is used to tune out the inductance.

In many practical gamma matched antennas there is no capacitance per se. The gamma rod can be a tube and another insulated rod can go from the braid of the coaxial cable into the gamma rod. So this type of gamma rod has an inbuilt capacitor formed by the gamma rod inside the gamma tube. This capacitance can be adjusted by sliding the gamma rod within the gamma tube. The position where the gamma assembly clamps onto the antenna is adjusted for the correct impedance match. The gamma match has applications for antennas other than the Yagi. **The gamma match is simply a method of moving the feedline along an antenna electrically until the correct impedance is found for the feedline (and low VSWR).** The gamma match has a means of tuning out any reactance. For exam purposes - identify the drawing.

TRANSISTORS - PROTECTED GATE MOSFETS

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 contained within the symbol the zener diode symbol as well. You may have to identify this.

TRANSMITTERS - DEMANDS OF AN SSB TRANSMITTER ON THE POWER SUPPLY

A point that I overlooked when discussing SSB transmitters is the demands that such a transmitter places on the power supply. Think about AM (DSB) for a minute. The load on the power supply will vary with modulation, but not by very much because the large carrier is transmitted even without modulation. Think about FM. The load on the power supply will be constant. The demands an FM transmitter places on a power supply in terms of voltage regulation are very low - because the load drawn from a power supply by an FM transmitter is constant.

Now, take say a 100 watt SSB transmitter. With no modulation, no RF energy is transmitted and very little current is drawn from the power supply. At 100% modulation the full peak envelope power of 100 watts will be produced, and in a typical transmitter this will draw about 20 amperes from the power supply. For any modulation between 0 and 100% the current drawn from the power supply will vary between just above zero and about 20 amps.

I hope you can see how demanding an SSB transmitter is on the power supply. One instant almost no load, the next instant full load, and the next - well it's anyone's guess.

Through all of this huge variation in load the SSB transmitter's power supply is supposed to produce a constant steady voltage output. For this to happen, **power supplies in SSB transmitters have much more voltage regulation.**

BYPASS CAPACITOR FOR AUDIO

Suppose an audio system (stereo sound system) is getting RF interference. Suppose this interference is being picked up by the speaker leads. One way to eliminate the problem would be to connect a bypass capacitor in parallel with (across) the speaker leads. In the exam, which would you choose if two of the optional answers were 0.47 uF and 4.7 uF?

This is a bit of a sneaky and on the surface confusing one that the examiner may ask. The purpose of the RF bypass capacitor is to provide a low reactance path for RF but not for AF. Capacitive reactance will be lowest at higher frequencies. Let's say the RF is 30 MHz and the highest AF for our audio system is 30 KHz (could be as high as 45 KHz even though you can't hear those frequencies - nevertheless you should not eliminate them - I wonder why?).

Let's work out the X_C for both capacitors at 30 MHz and 30 KHz using $X_C = 1 / (2\pi fC)$.

For the 0.47 uF capacitor:

$$\text{At 30 KHz, } X_C = 1 / (2\pi \times 30 \times 10^3 \times 0.47 \times 10^{-6}) = 11.28 \text{ ohms.}$$

$$\text{At 30 MHz, } X_C = 1 / (2\pi \times 30 \times 10^6 \times 0.47 \times 10^{-6}) = 0.0112 \text{ ohms.}$$

So placing the 0.47 uF capacitor across the speaker leads is placing a bypass path of 11.28 ohms for audio and 0.0112 ohms for RF. (The impedance of an audio amplifier is typically 8 ohms).

Compare that to the 4.7 uF capacitor:

$$\text{At 30 KHz, } X_C = 1 / (2\pi \times 30 \times 10^3 \times 4.7 \times 10^{-6}) = 1.128 \text{ ohms.}$$

$$\text{At 30 MHz, } X_C = 1 / (2\pi \times 30 \times 10^6 \times 4.7 \times 10^{-6}) = 0.00112 \text{ ohms.}$$

So which capacitor do you use?

You would choose the 0.47 uF because the 4.7 uF has far too low a capacitive reactance for both audio and RF. I did not choose these values, they came from an old exam paper.

My opinion is that the X_C for both capacitors is too low. I would not like my good sound system bypassed at the speakers by 11 ohms of X_C . You would lose too much high frequency response in your AF amplifier. It would be like turning the treble right down, and sound flat. I would go for 0.1 uF to start with and see if it worked. Frankly, a far easier method is to wind the speaker leads 'slowly' onto a ferrite rod - stopping when the interference stops - then tape the rod with PVC tape - done!

However, if you are asked in an exam, choose the best option - you will only have to choose from two.

Note:

There are many bypass capacitors in a radio transceiver. They are fixed capacitors and need to be small, since there are many of them. They also need to be stable, for these reasons ceramic capacitors are used as bypass capacitors.

End of Reading 41

Last revision: November 2013

Copyright © 1998-present Ron Bertrand

E-mail: admin@res.net

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

Free for non-commercial use with written permission