

Online Radio & Electronics Course

Reading 27

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

OSCILLATORS

WHAT IS AN OSCILLATOR?

Oscillators are used in radio circuits to produce radio and audio frequency energy, generally with a sinewave output, though the waveform can be many shapes such as a square wave or saw tooth. The sinusoidal waveform may be AC or DC. Oscillators used in radio frequency circuits are always very low power devices, in contrast to AC generators in a power station. Nevertheless, the AC power generator and the electronic oscillator are related, in that they both produce sinusoidal electrical energy. Unlike the AC generator though, the electronic oscillator can produce output on frequencies measured in tens of megahertz. Special oscillators can produce output at microwave frequencies.

An oscillator producing a radio frequency output is actually a low power transmitter in its most basic form. In an actual radio transmitter and receiver, up to several or more oscillators may be employed.

We are going to look a number of different types of oscillators and their circuits in this reading. Don't be put off by the number of circuits as you don't have to remember them in detail. You do need to learn the identifying features of each oscillator circuit. In the exam you will be asked to name the type of oscillator shown in the circuit. The fundamental principles of oscillator operation will be explained for each of the types. You will find a repeating theme across all of the oscillator types.

REQUIREMENTS FOR OSCILLATION

If any circuit has the following properties in the required amount, then the circuit will oscillate whether it is supposed to or not: -

- a) Amplification.
- b) A frequency determining device.
- c) Positive feedback (regeneration).

In oscillators, the factors above are designed into the circuit intentionally. The requirements (a) and (c) also occur in many amplifiers. For this reason, care must be taken with amplifiers to prevent or control, in particular, the third requirement for oscillation, positive feedback. Any amplifier provided with sufficient positive feedback will begin to self oscillate. Amplifiers are *not* supposed to oscillate, they are meant to amplify, though many amplifiers can easily become undesirable oscillators.

An amplifier which gets unwanted positive feedback will become an oscillator, and potentially cause interference. Amplifiers that oscillate generate a signal, rather than amplify one. Such unwanted signal generation can cause interference.

A generic-oscillator as shown in figure 1 is any amplifier with positive feedback.

Example: When sound from the speaker of a public address system gets back into the microphone(s) of that system, oscillation will occur. In audio circles it is called feedback. The amplifier “squeals”. When this happens at a radio frequency you can’t “hear it”. The oscillation is well beyond human hearing, but the effect is the same.

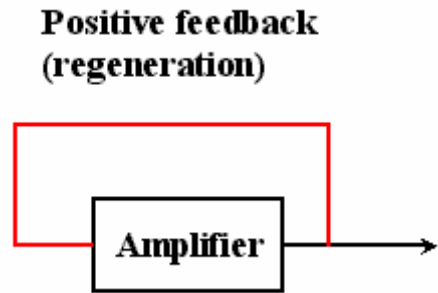


Figure 1.

FREQUENCY DETERMINING DEVICE

The frequency-determining device is usually a resonant LC circuit or a quartz crystal. Slices taken from quartz crystals make the most stable oscillators. See appended material on piezoelectric effect.

STABILITY

To ensure good stability, an LC oscillator should:

- a) Have a high C-to-L ratio.
- b) Have a well regulated power supply.
- c) Have good isolation between the oscillator and its load.
- b) Employ components which have low temperature coefficients.
- e) Not be exposed to large changes in temperature.
- f) Have all components mechanically rigid.

DRIFT

Drift is an unwanted *slow* change in the frequency output of an oscillator.

One of the main causes of drift in LC oscillators is unwanted capacitance changes in the circuit. These capacitance changes are mostly due to temperature effects. If the tuning capacitance is made *high* compared to the inductance in the **frequency-determining-circuit**, then such capacity changes will cause a smaller percentage change than if the tuning capacitance were *smaller*. Simply having a large capacitance compared to inductance produces a more stable capacitor, both in regards to mechanical rigidity and temperature effects. We say the stability is better with a **higher C-to-L ratio**.

BUFFER AMPLIFIER

A buffer amplifier improves the frequency stability of the oscillator by isolating it from the load. An oscillator is not able to deliver much, if any, power to other circuits. If too much power is taken from an oscillator, then it may be 'pulled' off frequency, or even damped so badly that it fails to oscillate. The buffer amplifier is placed immediately after the oscillator. The buffer amplifier has high input impedance and as such draws little or no power from the oscillator. The buffer has just enough gain to supply the following stage with usable power without loading the oscillator.

CHIRPING

In a telegraphy (Morse code) transmitter, the stage which is being keyed (by the Morse key) should never be too close to the oscillator as this can result in oscillator chirping. Chirping sounds like rapid short changes in frequency, very much like a canary chirping. What is happening is that sudden changes of load on an oscillator are occurring when the telegraphy key is closed, pulling the oscillator and hence the output of the transmitter off frequency. The chirp is the oscillator stabilising to the new frequency.

Interlude

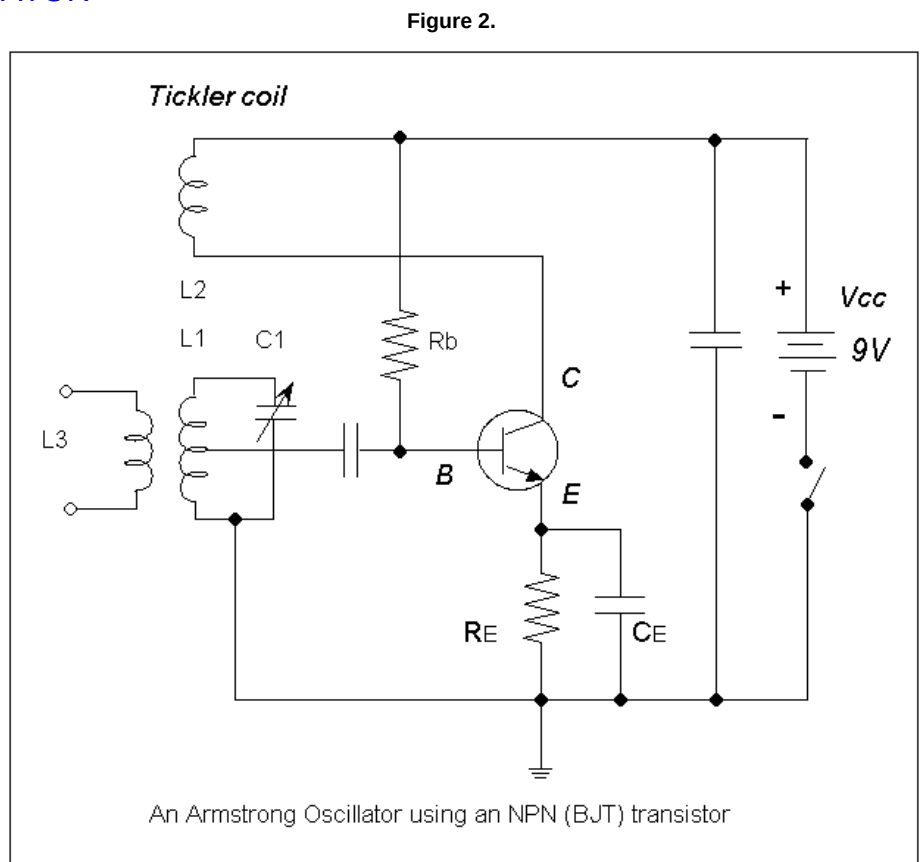
As we go through some different types of oscillator circuits, you will notice a common theme. I would like you to take notice of the:

- (a) Type of active device employed (however the type of oscillator is *not* determined by the active device used),
- (b) Polarity of the power supply and whether it is correct!
- (c) How feedback (regeneration) is achieved, and
- (d) What is the frequency-determining device.

If a portion of the circuit should be committed to memory, I will strongly emphasise this in the text. The shorthand alphanumerical labelling of components that I will be using is typical, though need not be remembered.

THE ARMSTRONG OSCILLATOR

In figure 2 an NPN BJT is used as the active device. L2 is called the **tickler coil** and is the distinguishing feature of an Armstrong oscillator. L2 provides regeneration to the input of the BJT. L2 does this by being inductively coupled to L1. Some of the signal in the output circuit is inductively coupled back to the input circuit by L2. The base circuit of the transistor contains a parallel tuned circuit consisting of L1 and C1. This circuit determines the frequency of operation. C1 is variable to change the frequency of oscillation. Provided the connections to the tickler are the right



way around, then feedback is positive (regenerative) and oscillation will be sustained. Connecting the tickler coil the wrong way would produce negative feedback (degeneration) and the circuit would not operate.

Rb provides for the correct amount of bias current. DC bias flows from earth (or negative) through R_E , into the emitter, out of the base, through Rb and then back to positive. The value of Rb and to a lesser extent R_E determines the amount of DC bias current.

Figure 3.

This is fairly straightforward. The amount of DC bias current is primarily determined by the value of the resistor R_b .

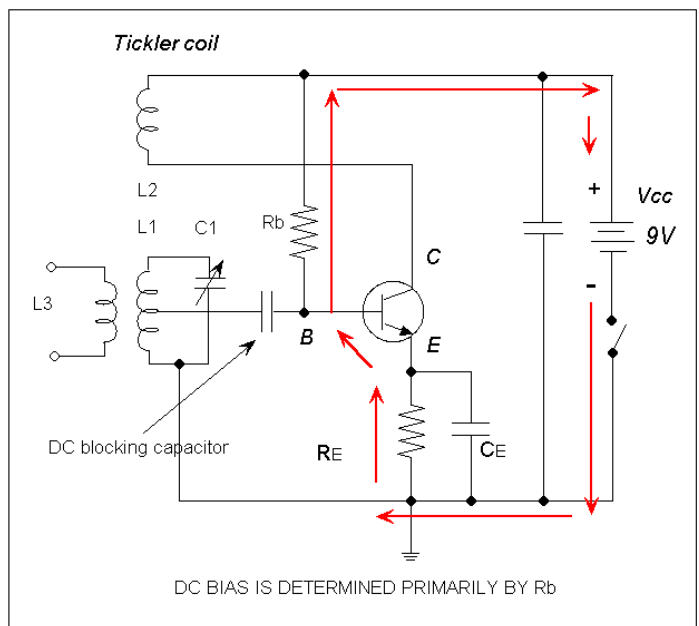


Figure 4.

A much larger current than the base current flows from the negative terminal of the battery – up through R_E , into the emitter out of the collector and back to the positive terminal of the battery.

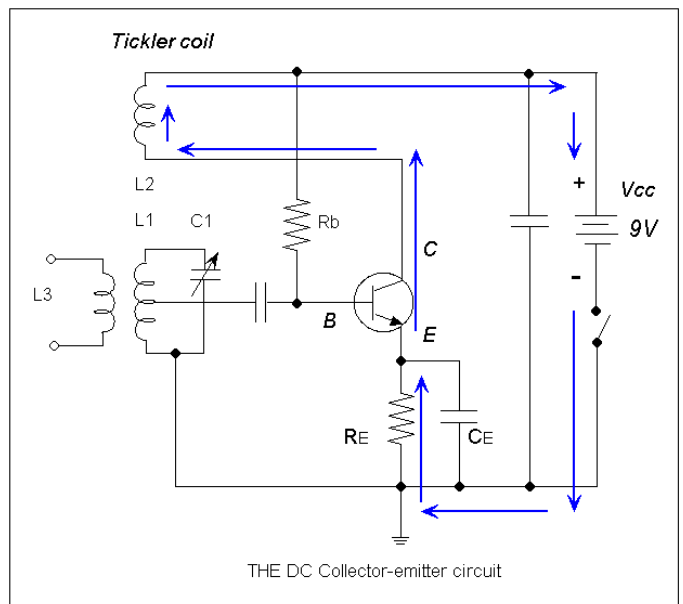
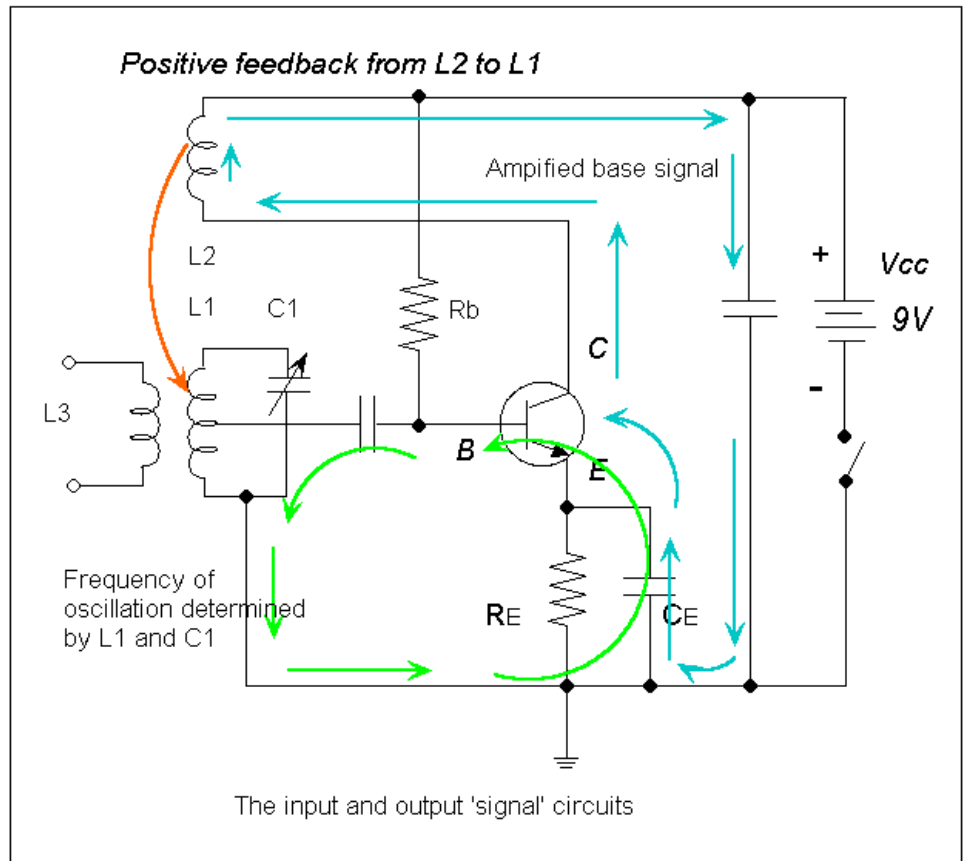


Figure 5.

The figure 5 circuit shows where the *signals* would flow in this oscillator. Assume that the oscillator is meant to produce a sinewave on 1 MHz. This will be a sinewave of varying DC not AC. Most active devices do not work on AC.

When the oscillator is turned on, L1 and C1 start producing an oscillation on 1 MHz. This oscillation would normally die down due to losses in the circuit of L1 and C1.

The oscillating voltage across L1 and C1 is superimposed on top of the DC bias current in the base circuit. So a 1 MHz signal current flows in the base circuit as shown in green. Notice how the signal flows through C_E and not R_E . (A little bit a signal current does really flow through R_E but not enough to be significant). The capacitive reactance of C_E at 1MHz would be $1/10^{\text{th}}$ the value of R_E .



Now this 1 MHz signal in the base circuit causes a 1 MHz signal in the collector circuit. The signal in the collectors circuit is much stronger and flows as shown in aqua. The capacitor across the battery bypasses the signal around the supply. We never want signal currents to flow through a battery or power supply. For one thing, the power supply is common to all stages. So if we allow signals from any stage into a power supply, they (the signals) can affect the operation of other stages via the power supply. You will nearly always see a power supply bypass capacitor, and often an RFC (radio frequency choke) in series with the power supply just to make it all that much harder for signal currents to get into the power supply.

Notice that the amplified signal flows in the tickler coil. The tickler coil (L2) is inductively couple to L1. If you like, think of the tickler coil as the primary of a transformer and L1 as the secondary. We have positive feedback from the tickler coil into L1 – so the oscillation is sustained.

L1 is also inductively coupled to L3, so we can take some of the signal current away from the circuit for use elsewhere. An oscillator would not be of much use if it did not provide us with an output – L3 is the output. We will discuss what we can do with oscillators in a future reading.

AN ELECTRON TUBE ARMSTRONG OSCILLATOR

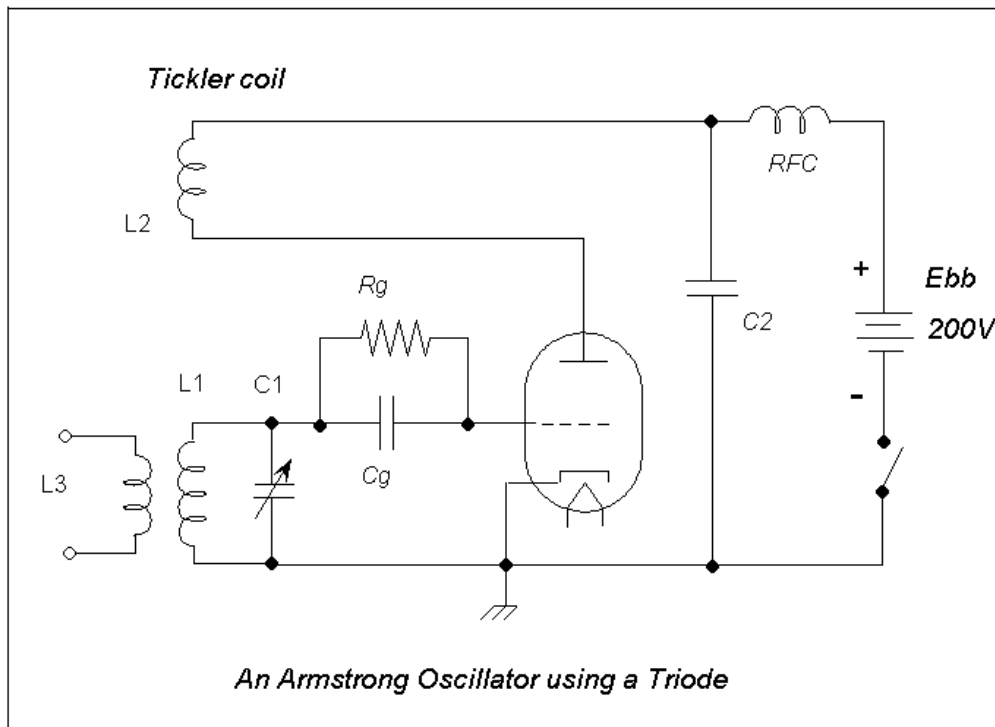


Figure 6.

As with the BJT circuit the **tickler coil** provides feedback. The tickler coil is the most identifying feature of the Armstrong oscillator.

L1 and C1 determine the frequency of oscillation. Output is taken across L3, which is inductively coupled to L1.

The oscillator is operating on a high frequency (radio frequencies). **The one stage that is common to all other stages in a radio system is the power supply, therefore the power supply is a potential path for each stage to interfere with the other stages.** C2 and the RFC (radio frequency choke), either alone or together, will be seen in many RF circuits. Because of the high frequency of the oscillator, C2 has a low reactance to RF energy created by the oscillator, and bypasses this energy around the power supply. The power supply is only shown as a battery for simplification. On the other hand, the RFC has a high reactance to high frequencies and blocks radio frequencies from entering the power supply. C2 and the RFC have no effect on the DC from the power supply getting to the oscillator. Looking at it another way, C2 and the RFC form a low-pass filter, allowing DC to pass from the supply to power the oscillator, but blocking RF from getting from the oscillator into the supply.

GRID LEAK BIAS

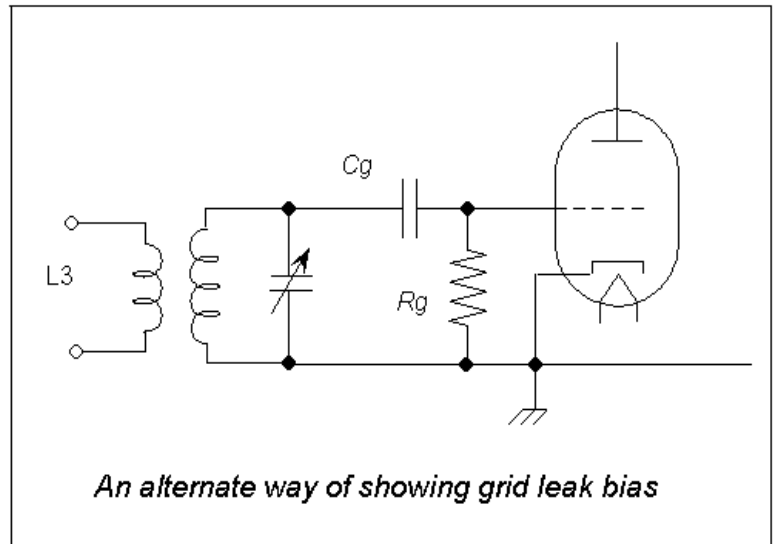
Bias for the electron tube is obtained by Cg and Rg. I have not discussed this bias arrangement before. In the case of an electron tube, it is called grid leak bias. This type of bias will also be found on a FET, in which case it is called **Gate Leak Bias**.

By now you should well and truly know that the purpose of bias on any active device is to place the correct DC operating potential onto the input element of the active device, to set the correct operating point and class of operation. However, in the case of BJT's it is bias current rather than voltage.

With grid or gate leak bias, you will always see the same configuration of R_g and C_g at the input of the device (the control-grid or gate). When the oscillator is first turned on, it will have no bias.

Oscillations in the parallel LC circuit ($L1$ and $C1$) will place an AC voltage on the grid. When the grid is positive it will draw current from the cathode (space charge), and C_g will be charged negative on the right and positive on the left. After a few cycles the negative voltage on the right will remain constant and provide the required negative bias.

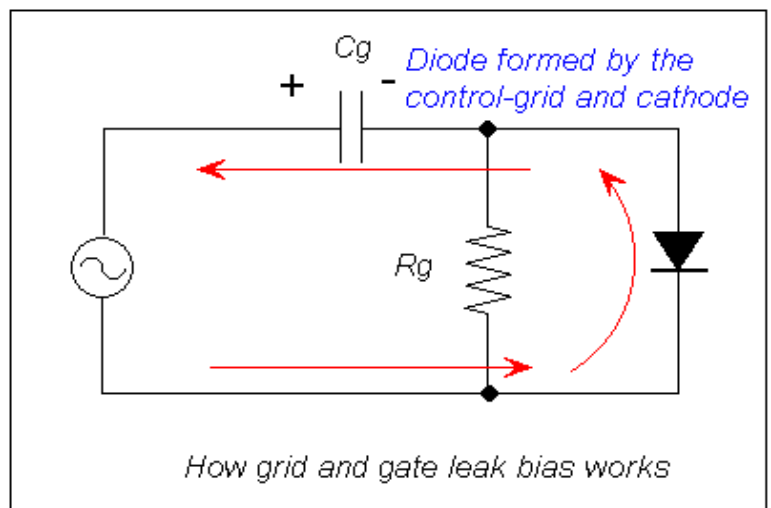
Figure 7.



The circuit of figure 8 shows how grid leak bias works. $L1$ and $C1$ are replaced with a source of AC. A diode rectifier is formed by the control-grid and the cathode. I have just shown a semiconductor diode here.

When the oscillator first starts up (when it is turned on), for a while positive half cycles will cause current to flow in the circuit as shown, and C_g will charge negative on the right hand side. It is this negative voltage that biases the control grid.

Figure 8.



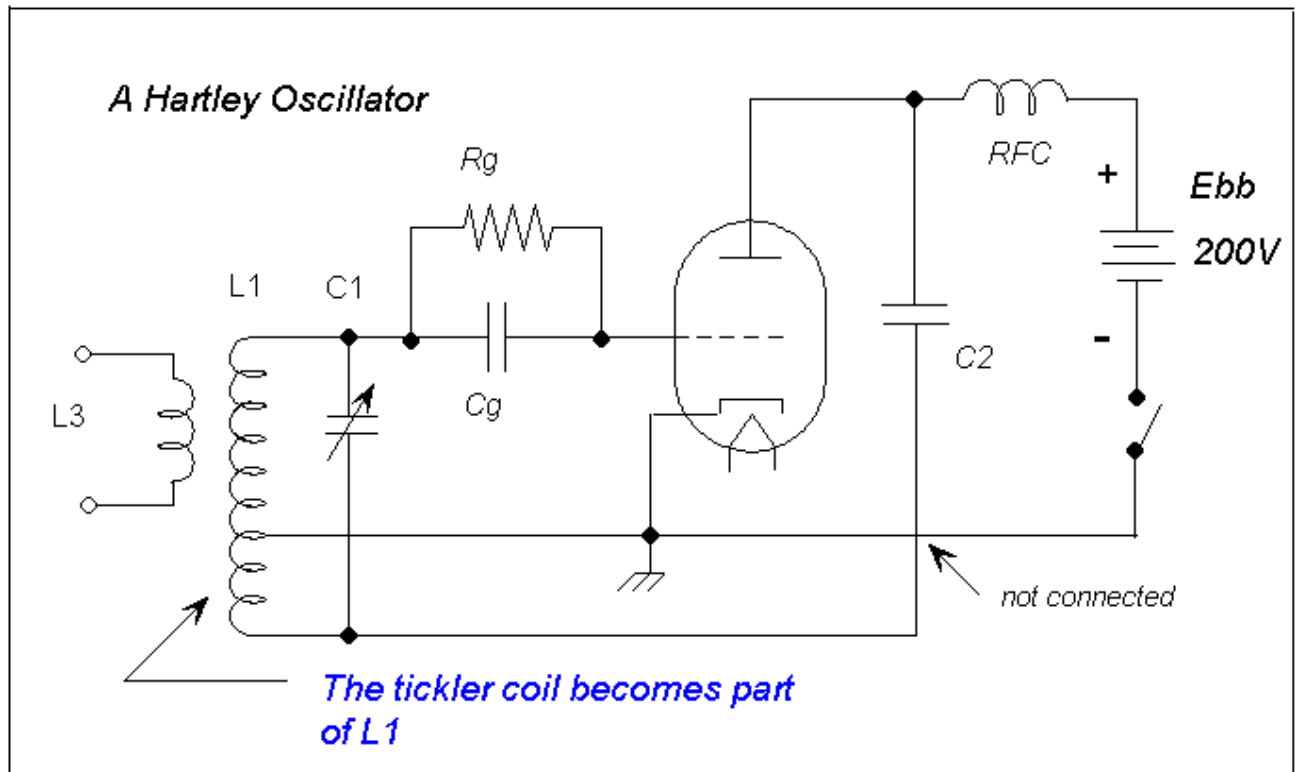
THE HARTLEY OSCILLATOR

The Hartley is a simple extension of the Armstrong. I suspect it came about because Hartley was lazy and worked out a way to make the Armstrong work without having to wind a separate tickler coil.

The tickler coil is now incorporated into part of the resonant tank inductor $L1$. The output in either circuit flows through a portion of $L1$ providing the necessary regeneration.

The Hartley oscillator is most easily identified by the tapped inductance. The tap position is adjusted to control the amount of feedback - it is not a centre tap as some diagrams suggest.

Figure 9 – Hartley Oscillator using Triode.



There is no difference between an Armstrong and a Hartley except the tickler coil is made part of L_1 . Instead of being mutually coupled, L_1 is an autotransformer, the primary of which is the tickler coil. All other operation of this circuit is the same.

A BJT HARTLEY OSCILLATOR

Again, just a Hartley oscillator (tapped inductance doing away with the tickler). L_1 and C_1 determine the frequency of oscillation. R_e and C_e are for emitter stabilisation and bypassing.

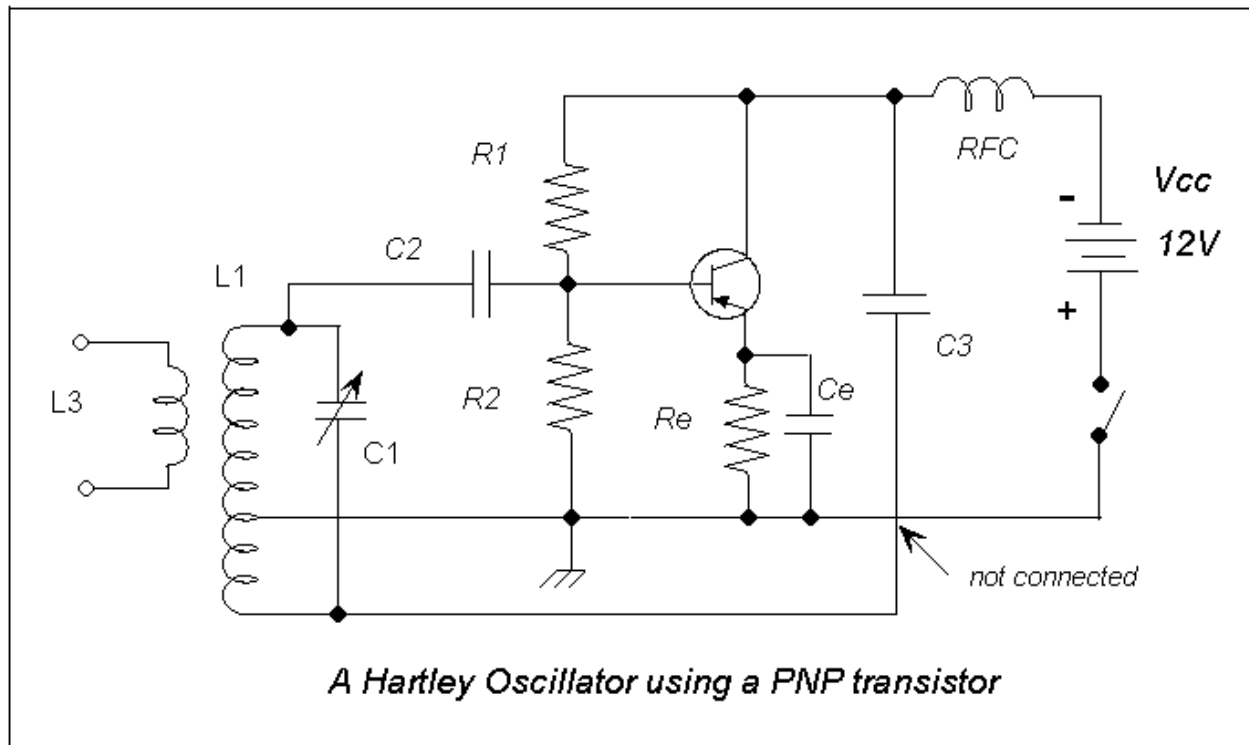
The input signal from the oscillatory circuit is taken from between the tapping and the top of L_1 . Feedback is injected back into the circuit because the output flows between the bottom of L_1 and the tapping.

About the only other difference with this circuit is that bias is now determined by R_1 and R_2 . This is called voltage divider bias.

C_2 is a DC blocking capacitor and C_3 is a power supply bypass.

Disconnecting the power supply from the 'signal' using RFC's and a bypass capacitor like C_3 , is sometimes referred to as power supply decoupling. So, you may see C_3 and the RFC described as 'power supply decoupling components'. This is just a fancy way of saying that the RFC's and bypass capacitors 'decouple' or disconnect the power supply from the point of view of the signal.

Figure 10 – Hartley Oscillator using Transistor.



THE COLPITTS OSCILLATOR

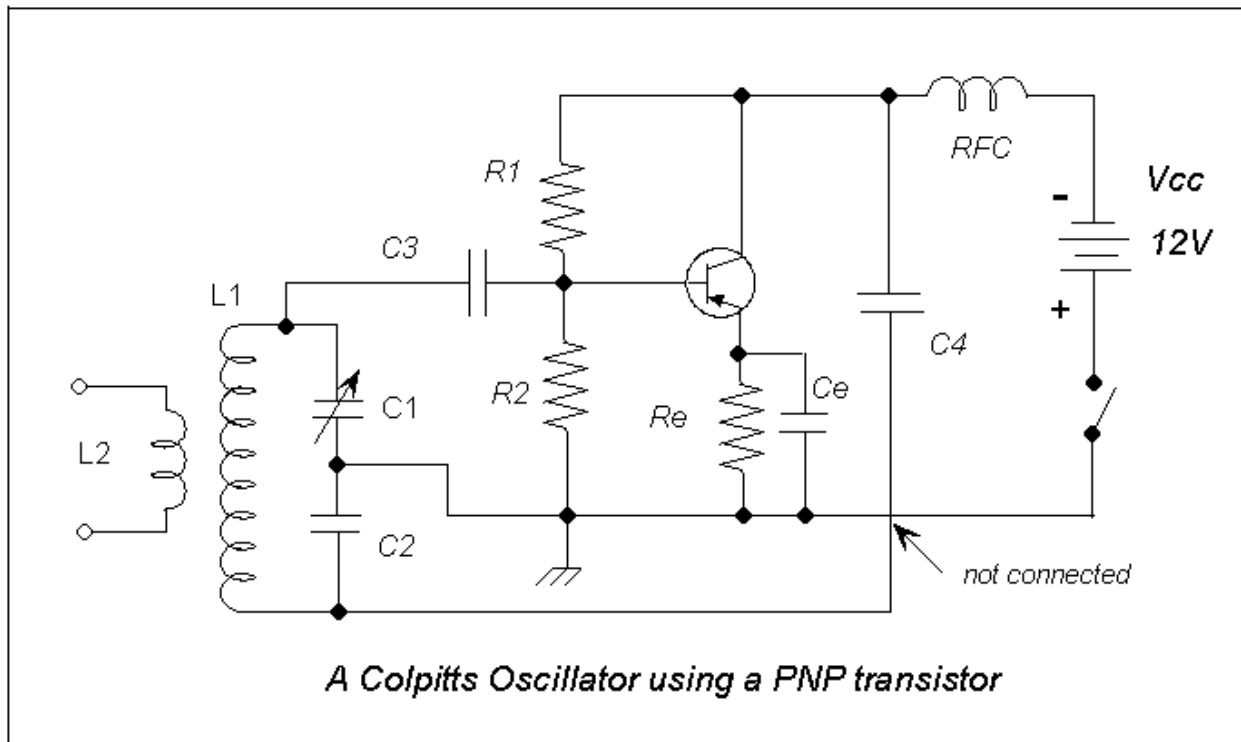
The configuration of the Colpitts oscillator resembles that of a Hartley in operation and appearance. The difference is that the tapping to the resonant circuit is now made with a **capacitive voltage-divider rather than with a tap inductance**. The output voltage is applied to the input via the voltage divider. Ratio of C1 and C2 controls the amount of feedback, a lot easier than fiddling with an inductance. Sometimes C1 and C2 are ganged to provide a fairly constant amount of feedback over a wide range of operating frequencies.

Ganged capacitors are two or more variable capacitors on the one shaft, with their movable plates connect to that shaft.

A good feature of the Colpitts oscillator is its comparatively good wave purity. This is due to the fact that C1 and C2 provide a low impedance path for harmonics, effectively shorting them to the emitter. The Colpitts is an exceptionally fine high-frequency oscillator and has been used as the VFO (variable frequency oscillator) in many amateur transceivers.

The Colpitts and the Hartley are the same – except the Colpitts uses a tapped capacitance rather than a tapped inductance to provide feedback.

Figure 11 – Colpitts Oscillator.



THE CLAPP OSCILLATOR

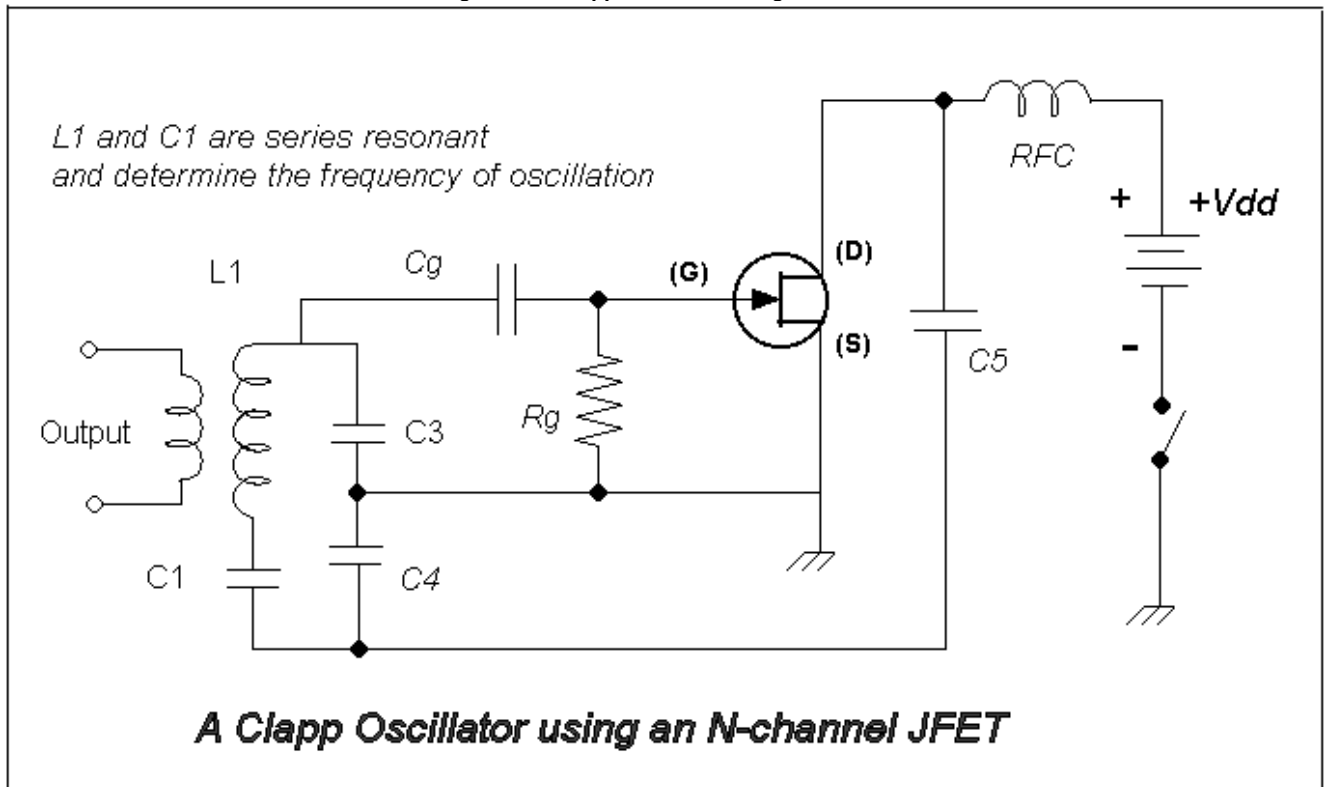
All of the oscillators we have discussed to now, have contained a *parallel* LC circuit at the input of the active device to determine the frequency of operation. The Clapp oscillator uses a *series* LC circuit.

In figure 12 below, $L1$ and $C1$ form a series resonant circuit to determine the frequency of operation.

The voltage-divider capacitors $C3$ and $C4$ perform the same function as in the Colpitts oscillator. The frequency of oscillation is slightly higher than the series-resonant frequency. Due to the fact that a series circuit has low impedance, the Clapp oscillator is less affected by variations in load conditions. The Clapp oscillator has excellent frequency stability and has frequently found applications in amateur transceivers.

Cg and Rg provide gate leak bias. It may look a little different from the previous circuits that used gate or grid-leak bias in that Rg is drawn between the gate and source rather than in parallel with Cg . Electrically it is the same thing. The purpose if Rg is to discharge Cg .

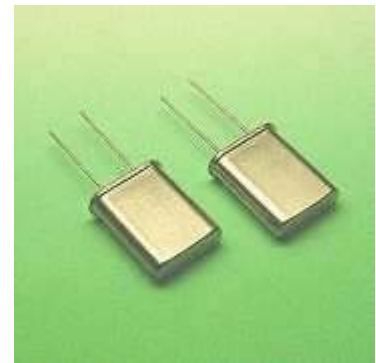
Figure 12 – Clapp Oscillator using JFET.



You have to be a little careful not to confuse a Clapp and a Colpitts. The Clapp has a *series* LC circuit as shown by L1 and C1 in figure 12. If C1 was removed in this circuit then it would be a Colpitts.

QUARTZ CRYSTALS

Quartz crystals or better put, thin slices or quartz cut from a larger crystal, exhibit the piezoelectric effect. They will oscillate just like a tuned circuit. The accuracy of the frequency of oscillation is extremely stable. Hence we have quartz watches, and many other timing devices.



FREQUENCY OF A QUARTZ CRYSTAL

The resonant frequency of a quartz crystal is **primarily** determined by its **physical dimensions**. However, cuts (the plane or angle of the slice through the main crystal) from the natural crystal will provide different frequency ranges and characteristics.

By proper selection of the type of cut, dimensions of the plate (the plates are the electrical contacts to the crystal) and mode of vibration, it is possible to obtain crystals with resonant frequencies from as low as 6 kHz, and as high as 75 MHz. For higher frequencies, the crystal plate becomes too thin and fragile, and is more susceptible to frequency changes with temperature.

MORE THAN ONE RESONANT FREQUENCY

A quartz crystal actually has two resonant frequencies, a **series-resonant** frequency and a **parallel-resonant** frequency.

A quartz crystal is capable of acting as a parallel resonant circuit or a series resonant circuit. The equivalent electrical circuit of a quartz crystal is shown in figure 13.

Figure 13.

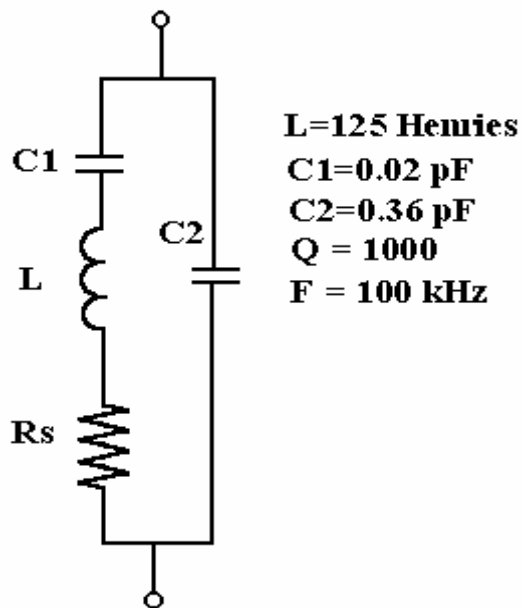
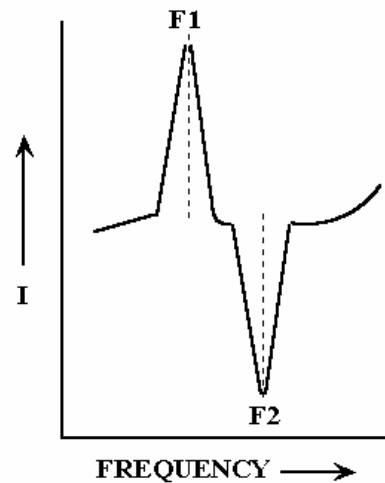


Figure 14.



The **equivalent circuit** is a combination of a series and parallel-tuned network. It is **not possible** to actually *construct* the equivalent electrical circuit of a crystal as any man-made inductor of the magnitude shown would have very large losses indeed. If a quartz crystal is placed in **series** with a signal path, then signals on the series resonant frequency will be passed easily through the low-impedance offered by the crystal. That is, it behaves like an LC series circuit.

In figure 14, current through the crystal is shown. The higher the current the lower the impedance. $F1$ is the series resonant frequency and the crystal has **low impedance** and consequently passes more current. At $F2$ the crystal is parallel resonant and acts as a **high impedance**. These properties become important later in order to understand how a crystal or a combination of crystals can be used to make a high quality (high selectivity) filter.

TEMPERATURE COEFFICIENT

The term "temperature coefficient" defines the way in which the crystal frequency will vary with temperature change. Crystals are usually rated in hertz-per-megahertz per temperature change in degrees celsius.

A crystal might have a positive, negative, or zero temperature coefficient.

If a crystal has a negative temperature coefficient (NTC) a temperature increase results in a frequency decrease. A positive temperature coefficient indicates that a temperature increase results in a frequency increase. A crystal with zero temperature coefficient will maintain a relatively constant frequency within the manufacturers stated frequency limits.

Temperature coefficients are usually expressed in hertz per megahertz per degree celsius (or centigrade), or more simply, in parts per million, with the degrees celsius understood.

When crystal stability is of paramount importance, the crystal is enclosed in a temperature controlled 'oven'. The crystal is then kept at a constant temperature via heating and a feedback mechanism to maintain the constant temperature. Crystal ovens were once used even in standard communications equipment. In modern times crystal stability has been much improved and ovens are not very common. However, they are still used where equipment is likely to be exposed to great temperature extremes (space, polar environments and the like).

OVERTONES

An overtone crystal is specially ground to obtain enhanced oscillation on **odd harmonics** of its fundamental frequency. **A crystal cut with a fundamental frequency of 10 MHz can be cut so as to enable oscillation on 30 and 50 MHz**, which are the third and fifth overtones respectively. The use of a crystal on overtone frequencies makes stable oscillator operation possible up into the VHF range. In many VHF/UHF amateur transceivers, the conversion of the incoming signal to the first intermediate frequency (IF) is accomplished by a local **overtone oscillator**.

Crystals that are to be used on overtone frequencies are always connected in series with a signal path because, for overtone operation, the crystal **MUST** operate on its series resonant frequency. For this reason, an overtone is more accurately defined as **an odd multiple of the series resonant frequency**. Most ordinary crystals can be used on the third or fifth overtone.

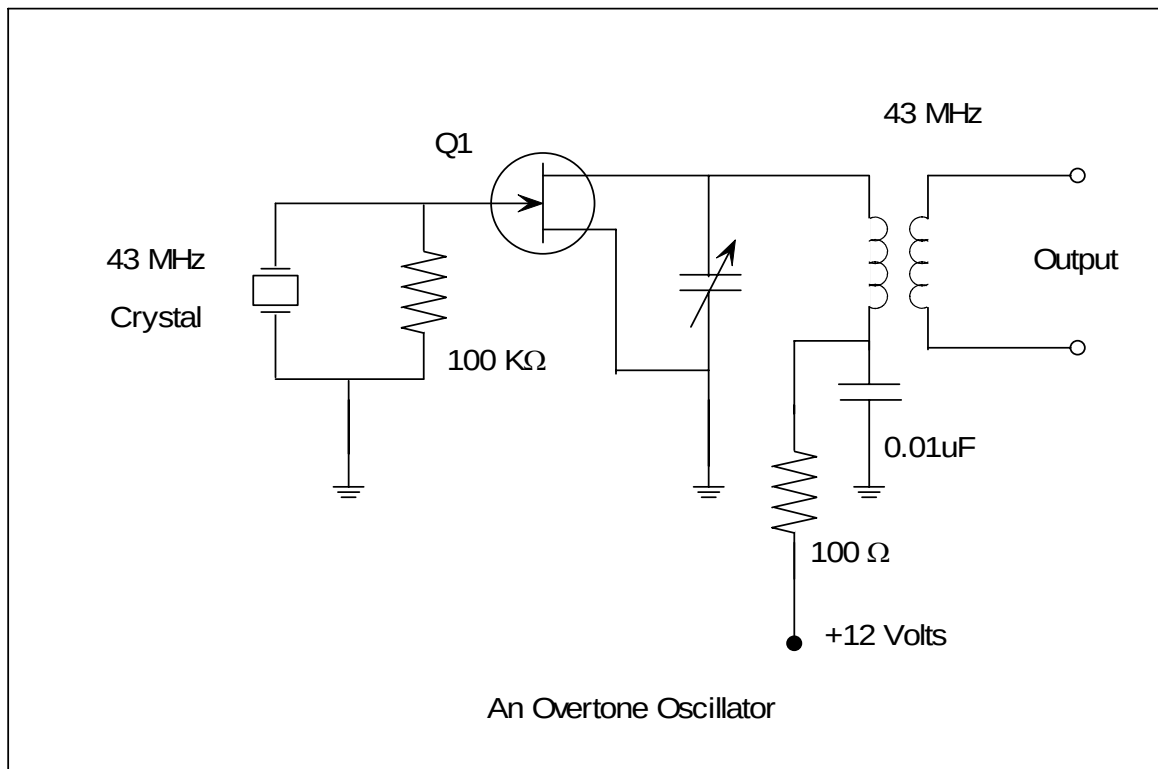
CIRCUIT OF AN OVERTONE OSCILLATOR

In the circuit of an overtone oscillator shown in figure 15 – the crystal frequency is shown as 43 MHz – this is very high for a crystal frequency. The fundamental frequency of the crystal would really be much lower for **stability**. The crystal is operating on an overtone frequency. The actual fundamental frequency of the crystal could be 14.333 MHz. However, the crystal has been cut in such a way that it will physically vibrate on an overtone. In this case 43 MHz, which is an odd multiple of the series resonant frequency of the crystal.

You can tell the circuit is not that of a harmonic oscillator, because the crystal is labelled with the same frequency as the output.

The variable capacitor at the output of the JFET would be used to tune the primary of the RF transformer to 43 MHz. The 100 ohm resistor and the 0.01 uF capacitor provide power supply decoupling. That is, they form a simple low pass filter to prevent RF from getting into the power supply.

Figure 15 – Overtone Oscillator.



THE HARMONIC CRYSTAL OSCILLATOR

A crystal oscillator with its output circuit tuned to any harmonic of the fundamental frequency of the crystal is called a **harmonic oscillator**. The harmonic oscillator takes advantage of the flywheel effect of the output stage to maintain oscillation in the same way as a frequency multiplier.

The harmonic oscillator operates on an entirely different principle to that of the overtone oscillator. In the harmonic oscillator, the crystal is physically vibrating on its fundamental frequency. In the overtone oscillator, no fundamental vibration or frequency is present anywhere in the circuit.

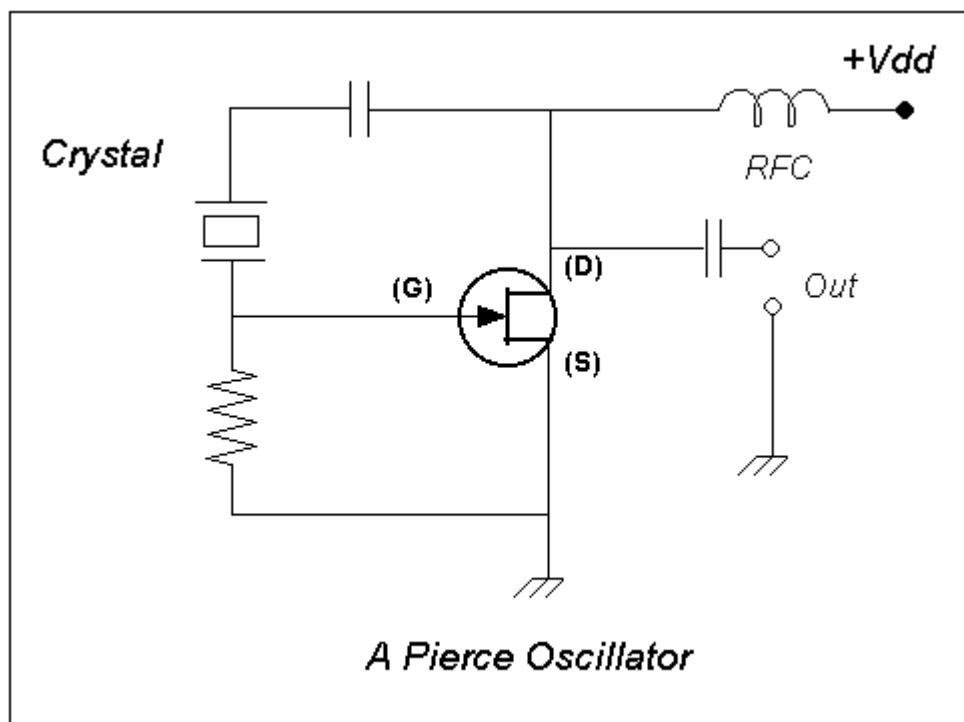
THE PIERCE CRYSTAL OSCILLATOR

The Pierce crystal oscillator **has no resonant tank circuit**. The crystal is connected directly between the output and input of the active device used. The crystal operates on its series-resonant frequency.

The Pierce oscillator is a crystal oscillator. It must have a crystal. An interesting thing about the Pierce is that the crystal provides the feedback path. The crystal shown below is connected between drain and gate. The crystal is series resonant meaning the crystal is operating on its series resonant frequency and has a low impedance. Any signal that can travel from the drain to the gate is positive feedback.

The resistor between gate and source provides the small bias voltage for the JFET.

Figure 16 – Pierce Crystal Oscillator



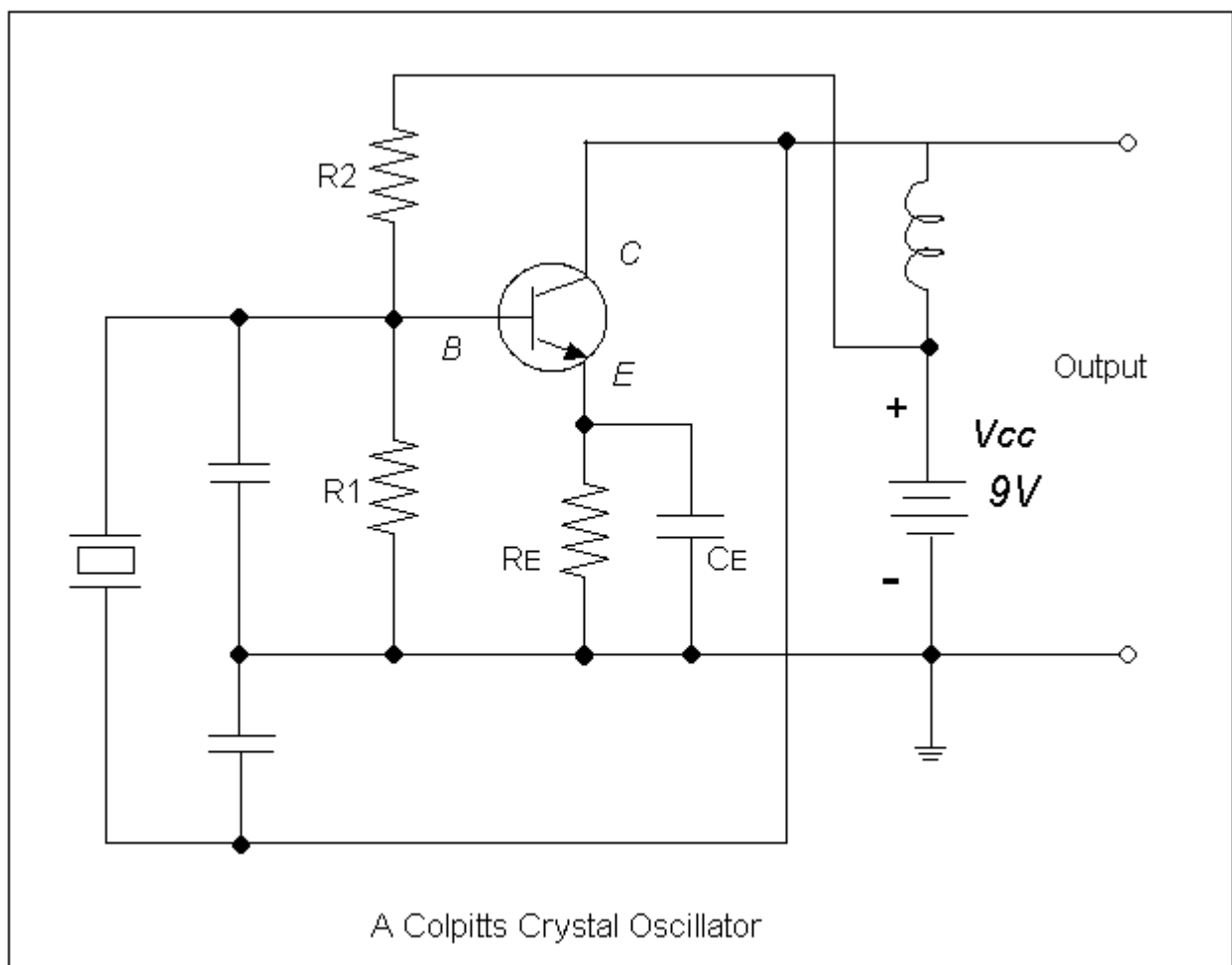
THE COLPITTS QUARTZ CRYSTAL OSCILLATOR

The Colpitts crystal oscillator is like any other Colpitts except a crystal is installed in place of the LC tank. As with a normal the Colpitts, feedback is provided by the capacitive voltage divider. Operation of the circuit is the **same** except that the crystal now determines the frequency of oscillation. The crystal operates on its parallel-resonant frequency. Since the Q of the circuit is high the feedback required is considerably less than with standard LC Colpitts.

A small 'trimmer' capacitor may be placed in series with the crystal to enable **small adjustments** to the crystal frequency (doing this is not a unique feature of the Colpitts). This capacitor would typically be 20 to 30 picofarads and one should not expect to vary the frequency by more than +/- a few kilohertz if predictable operation is to be maintained. Adjusting a crystal frequency in this way is called **pulling** the crystal. Try to change the frequency too much using this method and you may find that the crystal will jump to some unpredictable frequency, and operation becomes unreliable and unstable.

This excessive pulling is what is done in 27 MHz CB radios and referred by the users as a slider control. It is a bad engineering practice and illegal as it can cause the transmitter to operate on unauthorised channels (which is the intention). However, more important than the legalities is the interference to other radio services, which has the potential to endanger life.

Figure 17 – Colpitts Crystal Oscillator

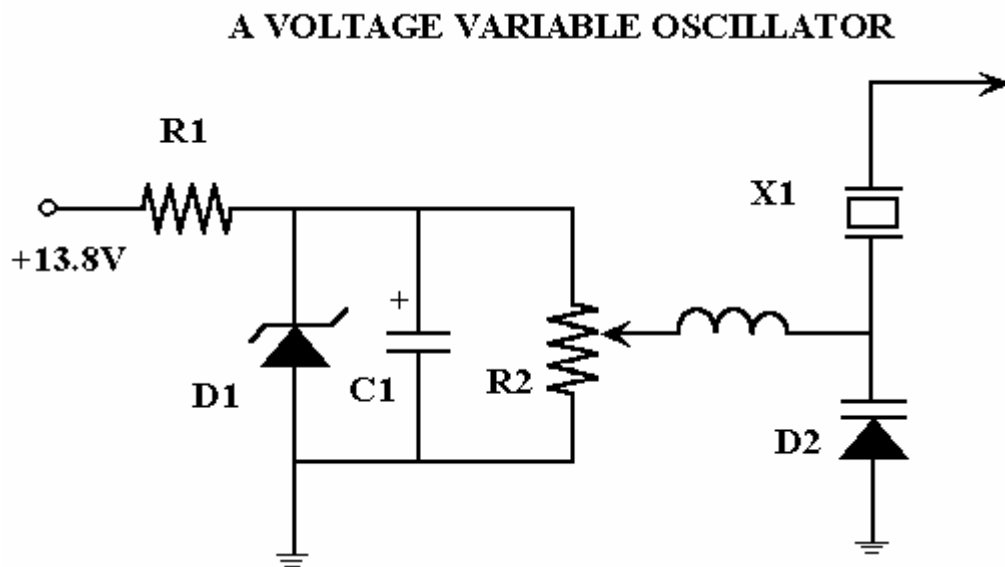


There is nothing new in the circuit of figure 17. The LC circuit has been replaced by the crystal.

A VOLTAGE CONTROLLED OSCILLATOR

Figure 18 shows one method of adjusting a crystal oscillator's frequency with a varactor diode.

Figure 18.



This technique is frequently employed in the clarify (single sideband is to be discussed) circuits of citizens band and other transceivers. The crystal would form part of a standard crystal oscillator. The reverse-biased varactor diode is connected in series with the crystal and the reverse bias voltage (and hence the junction capacitance) is adjusted by the potentiometer (variable resistor) R2. R1, D1, C1 and R2 provide an adjustable and regulated reverse voltage for D2.

You might ask why bother with such an arrangement? The alternative would be to have a variable capacitor in place of D2. In addition, this capacitance will, by necessity, have to be mounted on the front end of the radio. This creates all sorts of engineering difficulties. The wiring between the variable capacitor and the crystal all form part of the total capacitance in the oscillator circuit. However, with the varactor method just described, the variable capacitance is right at the crystal, and R2 can be mounted in any convenient position on the front panel, as the wiring between R2 and D2 has no appreciable effect on the operation of the oscillator. This is an excellent method to overcome the unwanted stray capacitance effects of wiring.

PHASED LOCKED LOOP

Modern radio equipment uses a method of frequency synthesis called Phased Locked Loop (PLL). Modern transceivers are required to operate over a very large range of frequencies and modern standards insist on frequency stability. Stability can be obtained through the use of crystal oscillators, though in a modern transceiver with perhaps thousands of channels, this would mean hundreds of crystals. This is where PLL comes in. A very clever technique to maintain the stability of a crystal using crystal oscillators, and yet provide a virtually unlimited number of operating frequencies.

PLL is covered in a separate reading on the website. This reading goes much deeper than required but is not a difficult reading. You are expected to know the fundamentals of PLL.

Please report any errors or problems with this reading (or any reading).

End of Reading 27.

Last revision: August 2006

Copyright © 1998-2006 Ron Bertrand

E-mail: manager@radioelectronicschool.net

<http://www.radioelectronicschool.net>

Free for non-commercial use with permission

APPENDIX – ADDITIONAL READING ON QUARTZ CRYSTALS

WHAT IS QUARTZ

The technical formula is SiO_2 and it is composed of two elements, silicon and oxygen. In its amorphous form, SiO_2 is the major constituent in many rocks and sand. The crystalline form of SiO_2 or quartz is relatively abundant in nature, but in the highly pure form required for the manufacture of quartz crystal units, the supply tends to be small.

The limited supply and the high cost of natural quartz have resulted in the development of a synthetic quartz manufacturing industry. Synthetic quartz crystals are produced in vertical autoclaves. The autoclave works on the principle of hydrothermal gradients with temperatures in excess of $400\text{ }^{\circ}\text{C}$ and pressures exceeding 1,000 atmospheres. Seed quartz crystals are placed in the upper chamber of the autoclave with natural quartz (lascas) being placed in the lower chamber. An alkaline solution is then introduced which when heated increases the pressure within the chamber. The autoclave heaters produce a lower temperature at the top chamber in comparison to the bottom. This temperature gradient produces convection of the alkaline solution which dissolves the natural quartz at the bottom of the chamber and deposits it on the seed crystals at the top. Alpha crystals produced by this method can have masses of several hundred grams and can be grown in a few weeks. If the temperature reaches $573\text{ }^{\circ}\text{C}$ a phase transition takes place which changes the quartz from an alpha to a beta (loss of piezoelectric property).

Quartz crystals are an indispensable component of modern electronic technology. They are used to generate frequencies to control and manage virtually all communication systems. They provide the isochronous element in most clocks, watches, computers and microprocessors. The quartz crystal is the product of the phenomenon of piezo-electricity discovered by the Curie brothers in France in 1880.

WHY IT WORKS

Piezoelectricity is a complex subject, involving the advanced concepts of both electricity and mechanics. The word piezo-electricity takes its name from the Greek piezein "to press", which literally means pressure electricity. Certain classes of piezo-electric materials will in general react to any mechanical stresses by producing an electrical charge. In a piezoelectric medium the strain or the displacement depends linearly on both the stress and the field. The converse effect also exists, whereby a mechanical strain is produced in the crystal by a polarising electric field. This is the basic effect that produces the vibration of a quartz crystal.

Quartz resonators consist of a piece of piezoelectric material precisely dimensioned and orientated with respect to the crystallographic axes. This wafer has one or more pairs of

conductive electrodes, formed by vacuum evaporation. When an electric field is applied between the electrodes the piezoelectric effect excites the wafer into mechanical vibration. Many different substances have been investigated as possible resonators, but for many years quartz has been the preferred medium for satisfying the needs for precise frequency generation. Compared to other resonators eg. LC circuits, mechanical resonators, ceramic resonators and single crystal materials, the quartz resonator has proved to be superior by having a unique combination of properties. The material properties of quartz crystal are both extremely stable and highly repeatable. The acoustic loss or internal friction of quartz is particularly low, which results in a quartz resonator having an extremely high Q-factor. The intrinsic Q of quartz is 10^7 at 1 MHz. Mounted resonators typically have Q factors ranging from tens of thousands to several hundred thousands, orders of magnitude better than the best LC circuits. The second key property is its frequency stability with respect to temperature variations.

End of document.