

Oscilloscope Traces in Pulsed Circuits

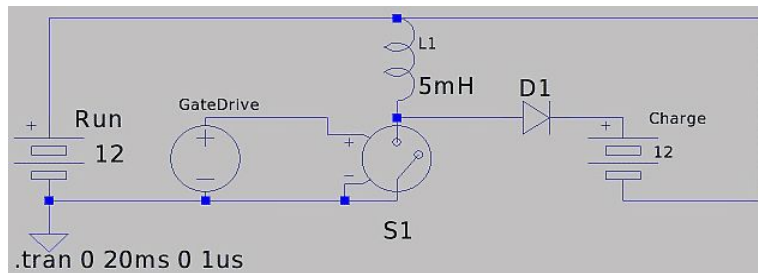
Basic Concepts

In this article I will give some analysis and descriptions of oscilloscope wave forms commonly seen in the pulsed circuitry used by those pursuing the anomalous charge characteristics of lead acid batteries that have been in electrical folk lore for many decades. My background is in electrical engineering, and my early career began as a designer of measurement circuits at Tektronix and at J. F. Fluke. In both cases I quickly learned what was needed to measure and design high speed circuits that could resolve pulses at the nanosecond level. For example, in designing practical instrumentation, input circuits are built on ground planes, and each wire is considered to be a transmission line with the need for proper termination to avoid reflections. In contrast to the pedestrian undergraduate electronic labs that strove to demonstrate first principles by using lumped, near ideal components and slow speeds, as one begins to turn the horizontal time per division knob on the oscilloscope up to the nanosecond range, one finds that the wave forms one encounters are filled with bumps and wiggles that were unexpected and unintended. The idea is that the lumped, idealized circuits used to teach theory are in reality made of wires, stray inductance, stray capacitance, resistance, and much else.

Every wire has self inductance, and has mutual inductance with every other wire in proximity. Every conductor also is capacitively coupled with every conductor in its vicinity. These factors will produce distributed effects that become more apparent as more transient energy is active in the circuitry. Add to this the fact that the oscilloscope and its probes can become part of the circuit (with a small capacitive load and a clip lead which is a small inductor), thus one should have at least a bit of healthy skepticism as to whether or not what one sees on the screen is actually reflective of what was hopefully being isolated and measured. It is therefore important to know what standard textbook theory has to say in such situations, so that when one encounters something that **deviates from expectation**, it can be delved into more properly such that any anomaly can be teased out with more clarity. Or to quote Yogi Berra:

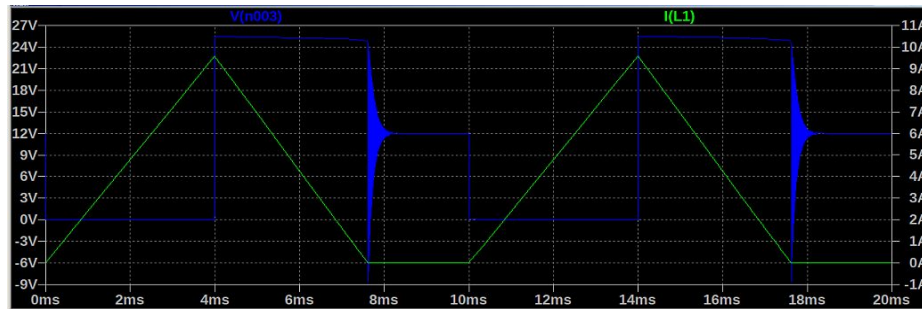
The difference between practice and theory is that, in theory, there is no difference between practice and theory, but in practice there is.

To begin, we start with the most basic switched inductor circuit that is used in countless DC to DC converters and switching power supplies, referred to as a boost converter, which is also the basis for the basic Bedini SG. Here I am using a Spice model with almost perfect switches, no interconnect effects, and ideal components in order to show theoretical behavior.



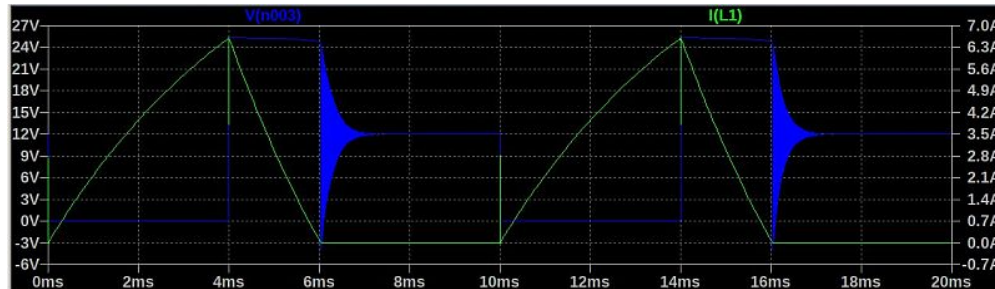
Basic Boost Converter

The switch S1 turns on and current linearly ramps up into L1 from the battery at a rate of $di/dt = V/L$. (As a rule of thumb, one can remember that 1 Volt across 1 mH will increase by 1 Amp in 1 mSec, applying proportional values as befits the situation.) The resulting wave forms look like this:



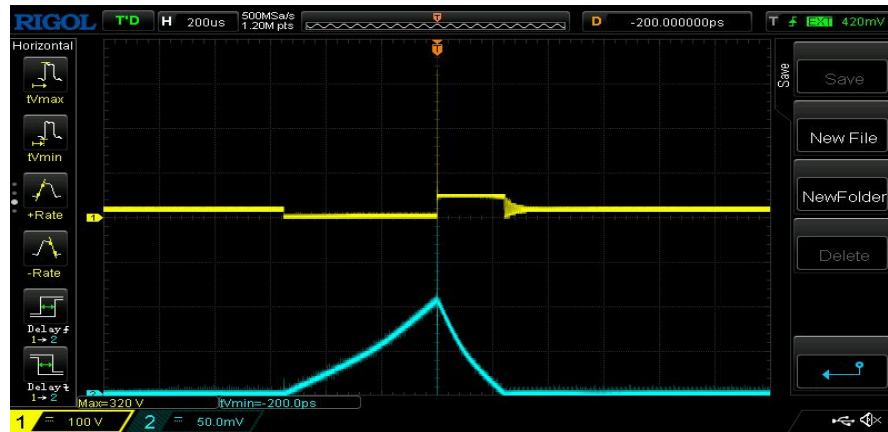
The blue trace is the voltage across the switch, and the green is the current through the inductor. The blue trace starts with the switch closure at 0V until 4 ms, at which point the switch opens. The current in the inductor has ramped up to almost 10A, and must continue to flow somewhere, so it switches to the diode D1 and into the charge battery loop. Note that the voltage across the switch now shows about 24.5V, which takes into account the diode voltage drop. Also, as the discharging current through the diode decreases, the diode voltage drop also decreases resulting in a small slope. As the current continues to decrease, the self inductance of L1 can no longer generate enough voltage to turn on D1, so the current through that branch goes to zero. But there is still some flux in the coil which must continue to flow somewhere, and it does so through the off resistance of S1 (which I set at 500K arbitrarily), and the capacitance of D1, to create a ringing damped oscillation. This is commonly seen in switching power supply circuits, and snubbers are often used to reduce such effects. Here I just wanted to highlight the idea that when an inductor is switched off, any remaining flux must go somewhere and be considered. In this case the actual energy contained in that damped oscillation is quite small, and most researchers looking for effects elsewhere can just let it be.

One sees a wide variety of inductors that researchers use to explore pulse circuitry. One factor to take into consideration is coil resistance. Here is how the above waves look with a more real inductor of 1 ohm series resistance:



Note that now the linear current ramps are rounded, due to the now apparent L/R time constant, and that there are small spikes on the switch voltage due to the coil's equivalent parallel capacitance, which I arbitrarily set to 500pF. If one is seeing such rounded wave forms, it is a sure sign of inefficient operation. This can be seen by the lack of symmetry in the charge and discharge portion of the current wave, with the area under the curve of the charge portion larger than the discharge portion, due to dissipation. If one sees this sort of thing, it is a sign that one would be better off with higher input voltage, and lower current, so that a more efficient, linear condition is used.

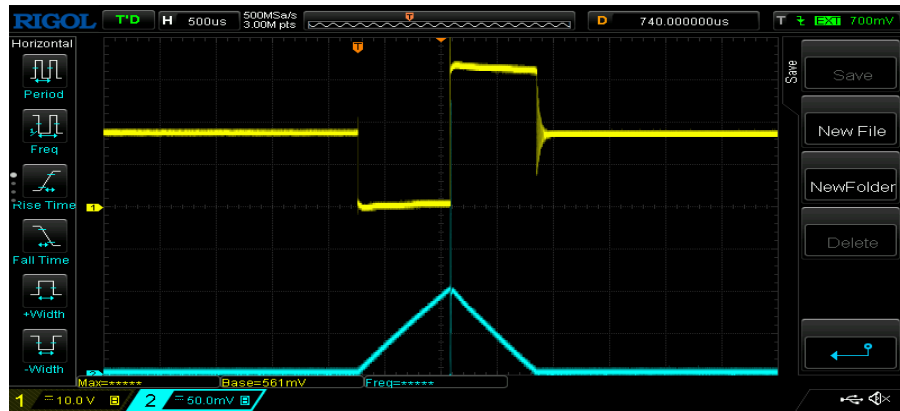
A somewhat different situation is seen in the oscilloscope trace below, showing an opposite curvature in the current wave form. This shows an inductor with a ferrite core. As the core sustains higher flux due to increasing amperage in the windings, it starts to lose its permeability, and the coil starts to reduce in total inductance. This causes the di/dt to increase, and hence the curved wave form. This is usually avoided in normal power supply design as it again results in reduced efficiency and potential heating in the switches due to current peaks.



Reality Strikes. Behold the Spikes

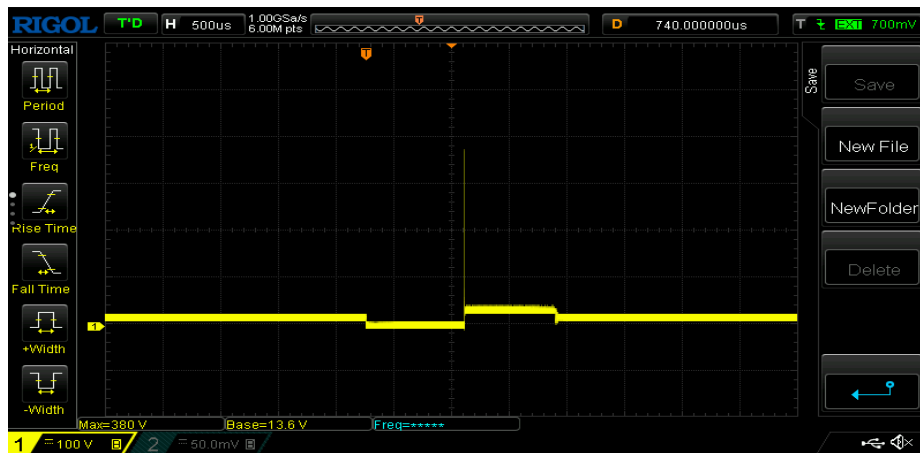


Now let us see what happens with this same circuit, but with real components and the usual sprawling (and in my case a bit messy) layout. Here I am using a fast SiC FET, a fast recovery diode, and a 12V starter battery on the charge end. The inductor is made of 8 parallel coils, giving a total of 1.2mH at .3 ohms. For testing, I am running the circuit from a variable DC supply rather than another battery, as it allows for testing a wide range of settings. The supply input to the circuit is very well bypassed at the switch itself so that it presents a very low impedance at all frequencies. If the anomalous battery charging is occurring mainly in the charge battery, that is where we will focus attention. Below is the trace for this circuit running at 14V input, and an on time of about 700 uS. The yellow is the voltage across the switch, and the blue is the current through the inductance:

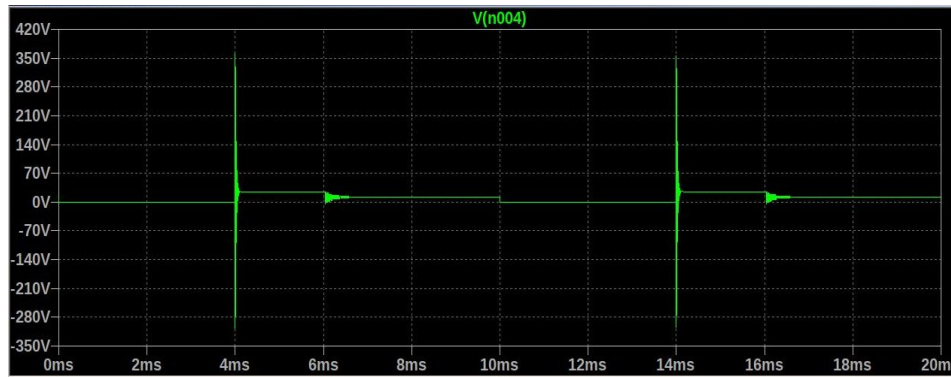


The scope shows substantially the same wave forms as in the model. In this case the peak of the current wave is 10A, and the wave shows good symmetry with low losses. The voltage wave shows the same slope we saw before, as the diode forward drop reduces with lower current. There is also a slight voltage rise when the switch is on, due to $I^2 \cdot R$ loss in the R_{ds} of the SiC FET. If one sees a substantial rise in that on state voltage with high current flow, that is an indication that a lower R_{ds} device is called for. The small bit of ringing on switch turn off is of no concern here (although in a much higher frequency (>30kHz) power supply, that condition becomes more important).

If we change the vertical trace on the voltage wave form below, we see that at the switch turn off there is a significant voltage transient of almost 400V peak. This is the MOSFET destroying impulse that requires attention lest one ends up with frustrating piles of dead devices. I use a device rated at 1200V, so there is no concern in this case. The main question I want to explore is, what are the factors that create this peak voltage? Is it avoidable or desirable? Or as we would say in the engineering world, “Is it a fault or a feature?”



We can return to the Spice model for a little insight. If we put a small inductor in the loop containing the charge battery, we can see that similar behavior is shown. I choose the .3uH value as it results in about the same peak voltage as above, and is consistent with the self inductance of typical wiring. This value is very small compared to the main coil, and is generally considered in conventional power supply design technique under the heading of **stray inductance**, i.e. the unspecified inductance in the wiring and connections. The Spice model below also shows additional ringing due to the idealized components, which I don't need to pursue as I have the real circuit to look at:

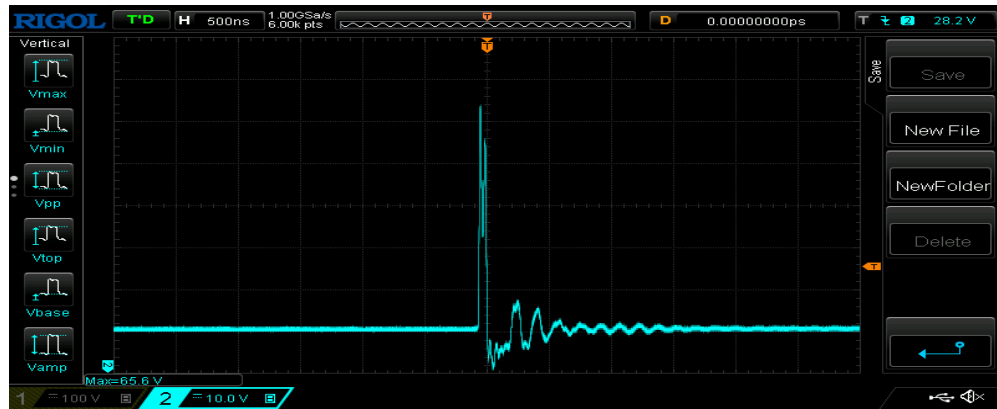


It is necessary to look into this in more detail, taking care to use the most direct and shortest leads possible to connect the scope to the circuit. First let us look in on the peak impulse itself with much higher sweep rate, and also overlay the SiC FET voltage trace with that of the diode output:



It can be seen that the transient is over in about 100ns. The blue wave is the diode output, and the yellow is the SiC FET drain as before. This shows that the diode is not significantly delayed in turning on, and has no major voltage drop, so it is not seen to be a contributing factor to the peak voltage of the transient. Note that the most important characteristic of the diode in these circuits is its **turn on** time, not the turn off time. The latter is what is being referred to in “fast recovery” diodes. Most diodes in the off state show mainly a series capacitance, and so will turn on relatively fast compared to their turn off time.

If the scope probe is now clipped across the battery terminals directly, a substantial reduction is seen. Note that in order to do this measurement, one needs to float the ground on the scope, and to just use the one probe by itself, otherwise short circuiting through the scope ground leads can occur. With the same conditions that gave the ~400V peak at the SiC FET, the trace now appears thus:



Here the peak voltage directly across the battery is now just 65 volts. This means that the two wires that connect the battery to the diode, and to the positive of the power supply, briefly have a total potential of over 300V across just ~25cm of length each. A quick look, using one of many online calculators, shows that the self inductance of just a simple wire gives about 250nH for a 25cm length with 2.5mm diameter (#10AWG). This result may be surprising for someone new to high speed switching circuits. In my case, I have also spent much time designing and using antennas, and at high frequency an antenna will routinely have hundreds of volts of RF potential between points on a radiator just 20 cm apart, so the pulsed situation above is not to be seen as anomalous or even uncommon.

The previous Spice model would indicate that this much peak voltage would be expected with wires of this length. We can confidently say that conventional stray inductance is definitely present and needs to be taken into consideration with these circuits. We can also say that the voltage spike is not occurring strictly “across the battery”, but rather that the battery is shunting a small portion of the total voltage in relation to its own internal equivalent inductance. We still need to come to a better understanding of what factors the transient depends on, and how to increase or decrease its impact. Typical battery pulser circuitry has focused on large, open core coils with long leads winding through relays and connectors. It is not surprising to hear tales of spikes of ~1000V.

A few measurements are immediately indicated. I changed my above setup to test the effects of various lead lengths and diameters between the pulser and the battery. Here is a list of wire size, peak voltage at the SiC FET, and peak voltage across the battery, for the same pulse width and supply voltage:

Wire Size	Wire lengths	Vpeak at MOSFET	Vpeak at battery
#10 AWG solid	8cm and 10cm	420V	300V
#6 AWG Solid	25cm and 25cm	480V	110V
#10 AWG Stranded	25cm and 25cm	480V	100V
#14 AWG Solid	25cm and 25cm	480V	95V
#4 AWG Stranded	60cm and 60cm	620V	88V

This shows that peak voltage at the SiC FET is mainly determined by lead length, and to a lesser degree by diameter. It is interesting that if one wants to have the largest spike “across the battery” it is necessary to use extremely short leads, and in fact would really necessitate mounting the circuit as much as possible right at the battery terminals, which tends to be a bit inconvenient. Others that try this test will certainly find somewhat different results, according to the details of their circuit, but the effects will be there.

It is necessary to see what the effect of the battery itself has on the circuit, which I expected to show some sort of nonlinearity due to the fast switching that would contribute to the peak voltage. To test this I removed the battery and made a noninductive (as much as possible) resistor of 1 ohm to be used instead of the 1.2mH inductor. In other words, my pulser was now just switching current on and off through a single resistor, with lead lengths as shown in the table below. My circuit otherwise had as short connections as possible, and used a multiply connected small copper ground plane under it to minimize any ringing or reflection of its own. The pulse width was adjusted to

give 10A of current as before. I connected the scope directly to the SiC FET leads and to the ground plane. With this simplest possible test, just turning on and off a resistor, I would expect to see minimal spikes. Here is what I saw:

Wire type	Lengths (2 leads)	V _{peak} at MOSFET
1 ohm soldered to MOSFET directly	10mm	84V
#20 AWG stranded twisted	15cm	200V
#14 AWG solid open loop	22cm	300V
Standard clip leads, twisted together	30cm	334V
Standard clip leads, open loop	30cm	400V
#6 AWG Stranded, close parallel	60cm	430V
#6 AWG Stranded, open loop	60cm	500V

Again, the effect is mainly with length, and also whether the leads are close to each other (like a noninductive transmission line) or whether they formed a loop of significant diameter. This table and the previous one give some clear indications of how circuitry layout can be designed to minimize the V_{peak} levels. Short lengths, close or twisted together, of moderate diameter or even of flat strips, will give the best reduction of stray inductance.

But it is surprising that even with 1 ohm noninductive, which was made from a number of smaller carbon resistors paralleled together, and soldered directly across the SiC FET leads, there is still a 84V spike generated. This shows that, given the lengths I went to minimize ground reflections and such, pesky stray inductance can still enter into the circuit. My resistor was perhaps not completely noninductive (hard to measure at this level), and there are still effects from the SiC FET leads.

Some summary comments

At this point, I can make some points that I hope will be useful to experimenters, and help them to interpret what they see in their builds.

- * The high voltage transient that is generated by turning off the switching device is higher and narrower the faster the turn off speed is. It's peak value is directly proportional to the amount of current that is switched from the inductor to the diode.

- * That switched current is abruptly shunted into the charge loop by the diode. The amplitude of the resulting spike can be mainly attributed to the distributed inductance of the lengths of the wiring particular to a given setup, and largely obeys the formula of $V=L*di/dt$.

- * The voltage peak can be reduced by anything that reduces that inductance, mainly shorter lengths and fatter conductors. While perhaps impractical, the use of a ground plane to create the lowest possible inductance in common connections can help.

- * Typically, only a small proportion of the total peak amplitude of the transient voltage will appear across the charge battery, in relation to its fraction of the overall charge loop inductance. Anything that can be done to reduce the length of the interconnects will increase the amount of the peak amplitude that appears across the battery terminals.

- * The received wisdom has been to assume the spike voltage is desirable, and that more is better. There are some conditions to this statement, however. If one were to compare two circuits, one with long interconnects routed through battery swapping relays, the other with very short and direct wiring, the latter would show a lower peak voltage but also show higher voltage peak across the battery.

- * Changing a given MOSFET, perhaps with a less than excellent gate drive, to an SiC FET with fairly good gate drive will show a higher peak transient voltage, but also a narrower pulse. In my case the peak went from around 300V to 450V, and from around 200nS to something like 100nS. Other circuitry will show different results, but it gives an idea of the range of variation which can be seen.

* Another piece of received wisdom with these circuits uses the concept of negative energy, which is said to be the result of discharging an inductor directly into a battery. Bearden has theorized about this, and has very often brought up the E. T. Whittaker papers concerning decomposition of electrostatic potentials. The notion of the importance of the spike potential has been inferred from this. While I am not particularly wedded to any theoretical notions at this point, if this theory has validity then I would expect that the actual voltage **as measured across the battery** is what does the work, and not the peak voltage at the switching device. This would then be an impetus to create a layout that, as mentioned above, increases the battery spike by reducing the total spike at the FET.

* The other energy form, that of positive energy, is said to be obtained by discharging the inductor into a capacitor, accumulating the charge, and then discharging the capacitor into the battery. As a side note here, author Moray King has, since his early work *Tapping the Zero Point Energy*, promoted an alternative theoretic idea, that of vacuum polarization. This is occurring when a large, sudden **current** pulse is applied to a battery, which causes disequilibrium between the vacuum and the ions and free electrons in the electrolyte. The two theories, one current based and the other voltage based, are far from proven but are worthy of consideration.

In any case, if one is going to make use of the capacitor discharge method, there arises another option to further reduce the voltage spike at the FET. One can split the charged capacitor into two parallel parts. One portion can be small, a few microfarads at most, of good quality and low ESR, which can be located right next to the diode output using very short leads. The rest of the capacitor can be paralleled to that small part, and be located where convenient. The inductance of the resulting leads is effectively bypassed by the small, low ESR capacitor. This has worked well for my tests.

The above considerations and tests have been mainly focused on solid state or slow rotation devices, which show mainly inductive effects. In part 2 of this document, I will do some similar treatments for higher speed rotating devices, which begin to depart from simple inductive effects due to the motion of salient pole rotors using permanent magnets. It becomes more complex if one is to derive both battery pulse charging and maximized shaft torque at the same time.