



Introduction

Switching regulators have been with us for many years. They were considered tricky to design – and still are. In 1976 Silicon General (later Linfinity, then Microsemi), introduced the first monolithic (IC-based) switching controller, the SG1524 “Pulse Width Integrated Circuit”. A little later, this chip was improved and became the “SG3524” historic industry workhorse. And very soon thereafter, it was available from multiple chip vendors. Keep in mind that switching stages based on discrete designs were already gaining ground, particularly in military applications. In fact, some resourceful engineers had even made “switch-mode power supplies” by adding related circuitry around one of the highest-selling chips in history: the “555 timer” (sometimes called the “IC Time Machine”), introduced in 1971 by Signetics (later Philips, then NXP). The SG1524 was however the *first* IC in which *all* the required control functionality was present on a single chip/die. With the rapidly escalating concurrent interest in switching power supplies at the time, it is no surprise that as early as 1977 the very first book on the subject, written by the late Abraham Pressman, appeared on the scene. Together, these events spurred interest in an area well beyond most people’s expectations, and ushered in the world of switching power conversion as we know it today.

The SG1524/3524 drove a pair of (bipolar) switching transistors with a “*duty cycle*” (*ratio of switch ON-time to the total time period*) which was proportional to the “*control voltage*”. By using switching transistors to switch the input voltage source ON and OFF into an LC low pass filter, a relatively efficient voltage regulator was produced. At the heart of this regulator was the “*PWM (pulse width modulator) comparator*”. The output pulse-train to drive the transistors with was a result of applying a (relatively) smooth control voltage on one of this comparator’s input terminals, along with a “sawtooth” or a “PWM ramp” *generated from the clock*, on its other input. See *Figure 1*. This technique is known as “voltage-mode programming”, or “voltage-mode control” (“VMC”) – since the duty cycle is proportional to the control voltage. The control voltage is in effect the difference between the actual output voltage and the “reference” value (the value we want to fix the output at, i.e. the “setpoint”). We will discuss this figure in more detail shortly.

Another well-known technique today, which has also been around since the 80s, senses the peak current in the power switch or inductor, and turns the switch OFF at a programmed level of current. This technique is called current-mode control (“CMC”). Keep in mind it was not “brand-new” at the time. In fact it had been discovered years ago, but few had realized its significance till Unitrode Corp. (now Texas Instruments) came along. It received a huge boost in popularity in the form of the world’s first current-mode control (CMC) chip, the Flyback controller UC1842 from Unitrode. In CMC, there is in effect a (fast-acting) “inner” current loop along with the “outer” (slower) “voltage loop” which carries out the output regulation. See the note on ramp generation within *Figure 1*, indicating how this particular aspect is different between VMC and CMC. Prima



Microsemi TN-206 Voltage-Mode, Current-Mode (and Hysteretic Control)

by Sanjaya Maniktala

facie, CMC *seems to be better*. “Pulse-by-pulse” became synonymous with CMC. It was once even thought to be the silver bullet, or magic wand, to fix *everything* that *voltage-mode control* was not. The UC1842 was later improved to UC3842, and shortly thereafter, following the success story of the SG3524, the UC3842 was soon available from innumerable chip vendors. But a few years into this success story, *expectations got somewhat blunted*.

The disadvantages of CMC surfaced slowly. That growing realization was succinctly summed up in a well-known Design Note – DN-62 – from Unitrode, which said: “*there is no single topology which is optimum for all applications. Moreover, voltage-mode control – If updated with modern circuit and process developments – has much to offer designers of today's high-performance supplies and is a viable contender for the power supply designer's attention.*” It also says: “*it is reasonable to expect some confusion to be generated with the introduction of the UCC3570 – a new voltage-mode controller introduced almost 10 years after we told the world that current-mode was such a superior approach.*”

To put things in perspective, the above-mentioned design note was written by “the father of the PWM controller IC industry,” Bob Mammano, who developed the first voltage-mode control IC, the SG1524. Later, as Staff Technologist in the Power IC division of Unitrode (a division that he had jointly created with two others from Silicon General), Mammano led the development of the first current-mode control IC, the UC1842.

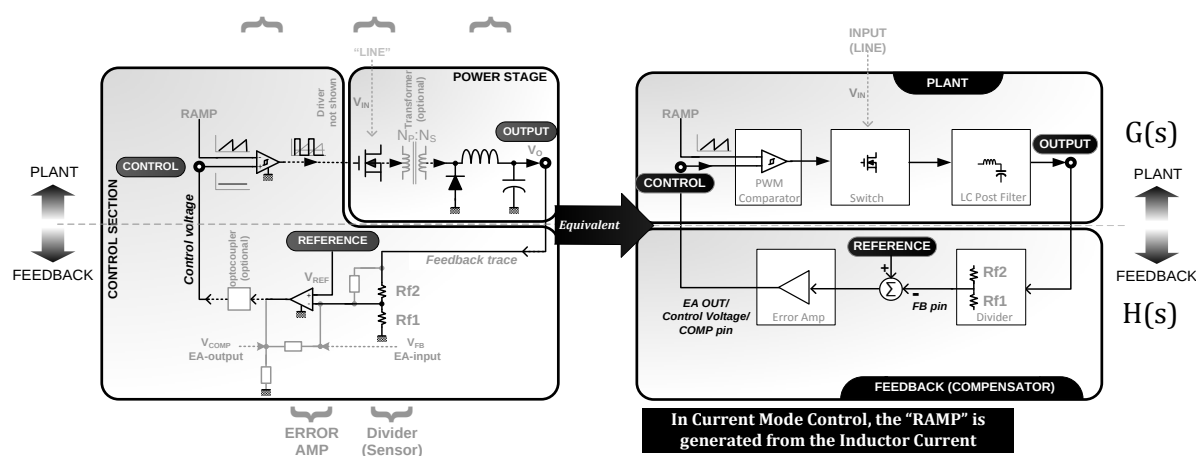


Figure 1: Voltage and Current Mode Regulators with their shared Functional Blocks



Building Blocks of Switching Regulators and Stability

In *Figure 1*, we see the building blocks of a typical switching regulator. There is a “power stage”, consisting of switch/diode, inductor/transformer, and input/output caps. The input “ V_{IN} ” comes into this block and gets converted into the output, V_O . Around this block is the “control section” block, consisting essentially of a voltage divider, an error amplifier, and a PWM comparator. In classic voltage mode control, the voltage ramp to the PWM comparator is fixed, and is artificially generated from the clock. In current mode control, this ramp is the sensed inductor/switch current mapped into a proportional voltage ramp that is applied to the PWM comparator. Keep in mind that in both VMC and CMC there is a clock, and its basic function is to determine the moment the *switch turns ON* in every cycle. The moment at which the *switch turns OFF* within each cycle is determined by the “feedback loop”. By using a clock, we ensure a constant repetition rate, or *constant switching frequency*, something that is considered desirable for switching regulators, particularly for complying with EMI limits.

Note: “Hysteretic controllers”, discussed later, typically dispense with the error amplifier and the clock (and though they retain something quite similar to the PWM comparator, they implement it in a very different manner). Therefore, trying to keep a constant switching frequency in that case becomes a major design challenge.

On the right side of *Figure 1*, we show how the switching regulator can be mentally partitioned and visualized when discussing *loop stability*. Notice the terminology in use. We see that the PWM comparator is considered part of the “Plant”, along with the power stage, and the rest falls into the “Feedback” section also called the “Compensator”. Their respective transfer functions are denoted as $G(s)$ and $H(s)$ respectively.

In the mathematical treatment of loop stability, we define a “transfer function”, which is basically the output of a given block divided by its input. The output and input do not have to be voltages, or currents, or even similar parameters. For example, the output of the PWM comparator is the “duty cycle” whereas its input is the “control voltage” (output of the error amplifier). The magnitude of this transfer function is called its “gain” and its argument is its “phase”. Sometimes, the transfer function itself is just called the (complex) “gain”.

Objectives and Challenges of Loop Design

The entire purpose of loop stability consists of two tasks. *First* is to discover what $G(s)$ (the plant transfer function) is, i.e. its Gain-Phase plot with all its inherent “poles” and “zeros”. The *second* is to design the compensator (feedback section) accordingly, such that its poles and zeros are “correctly located” with respect to those from the plant. What does “correct location” mean? The criterion for that is based on what we want the “open-loop gain” to look like. In *Figure 1*, the *overall gain* going



once fully around in a loop, passing through the plant and feedback blocks, is called the open-loop gain. Clearly, it is the product of two gain blocks in succession, or the product " $G(s) \times H(s)$ ". Once we know what $G(s)$ is, we can design $H(s)$ such that $G(s) \times H(s)$ (the combined gain) is *very close to a straight line of slope "-1"* (10× fall in gain every 10× increase in frequency). That is the basic simplified criterion or design target for stabilizing the loop of any switching converter, whether it is CMC or VMC. The difference is that in CMC the plant transfer function $G(s)$ is very different from the plant transfer function based on VMC. But the *final result*, the shape of $G(s) \times H(s)$, is intended to be the same (-1 slope). What unfortunately happens is that as *line and load variations occur*, $G(s)$ can change quite a lot in one control method compared to the other. So we may set the loop correctly at some "sweet spot" only to find it changing a lot, usually undesirably, as line and load change. And that is indeed where the real differences between CMC and VMC become more apparent. This is discussed next, and highlights the challenges to compensator design using CMC or VMC.

Note: A slope of "-1" is a straight line on a "log Gain" versus "log Frequency" plot, with a slope such that the gain decreases by a factor of 10 (or "20 decibels") every decade (10×) increase in the frequency. Alternatively put, that is a decrease of 6 decibels (a factor of 2) every octave (2×) increase in frequency.

Note: The product " $G \times H$ " is called the "open-loop gain" even though the loop itself is closed. However, for that reason it is sometimes erroneously called the "closed loop gain", which is actually a different term in loop theory.

Frequency Domain Analysis

For the feedback loop's most generalized analysis, and the overall response to an arbitrary stimulus or disturbance, a transfer function is usually written out in the complex-frequency plane, or "s-plane". This is done because it actually simplifies the analysis significantly, compared to trying to do the same thing as a function of time. See *Figure 2*. The former is called frequency domain analysis, the latter is time domain analysis. But eventually, the entire concept of the frequency domain is actually just a *mathematical construct*, not even necessarily intuitive, considering we now even have negative frequencies to integrate over. So eventually, we do need to map the responses from this "frequency domain" back into the "time domain", and that is what is eventually *observable* to us.

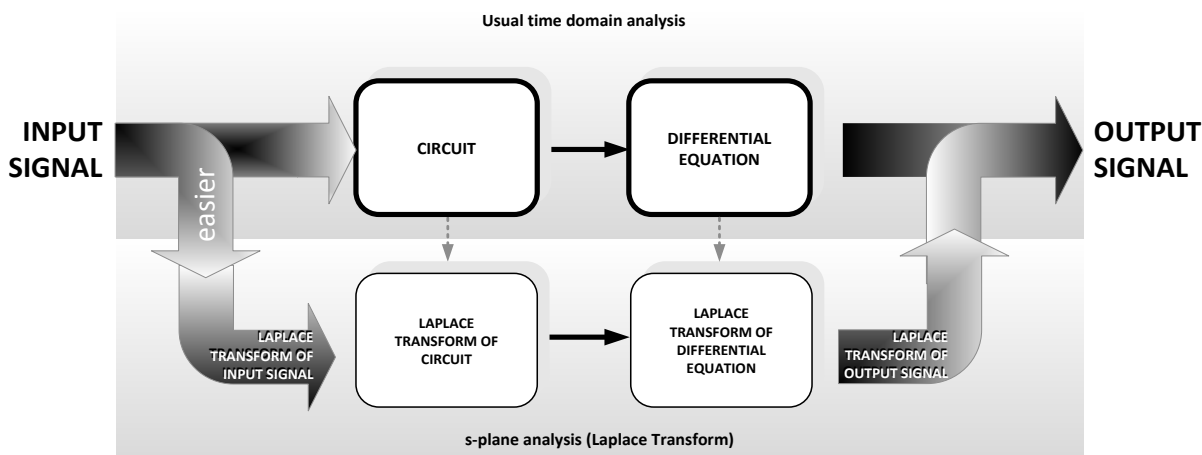


Figure 2: Working in the Frequency Domain versus in the Time Domain

Note: We define " $s = \sigma + j\omega$ ". Here, " σ " is a number that mathematically speaking just helps functions converge when we carry out integration, but is also ultimately responsible for imparting to any real-world response – a real-world exponential decay with respect to time, based on the function $e^{\sigma t}$. " j " is the usual imaginary number $\sqrt{-1}$ and $\omega = 2\pi f$, where f is the frequency. Together these are responsible for the oscillatory part of the response, based on the function $e^{-j\omega t}$. Keep in mind that all this is just a more generalized extension of what we learned in high school: any repetitive ("periodic") waveform of an almost arbitrary shape can be decomposed into a sum of several sine (and cosine) waveforms of frequencies. That is what Fourier series analysis is. In Fourier series, though we do get an infinite series of terms, the series is a simple summation consisting of terms composed of discrete frequencies (the harmonics). When we deal with more arbitrary wave shapes, including those that are not necessarily periodic, we need a continuum of frequencies to decompose any waveform, and then understandably the *summation* of Fourier series terms becomes an *integration* over frequency. That is, how the Fourier series evolved into the "Fourier transform," and from there on to the Laplace transform (mathematics in the s-plane). In general, decomposing an applied stimulus (waveform) into its frequency components, and understanding how the system responds to each frequency component, is called "frequency domain analysis." The "s-plane" is in effect a complex frequency plane, and Laplace transform is the way to conduct the most generalized form of frequency domain analysis. Eventually, though, we should not get carried away, for the frequency domain is just a *mathematical construct* --- not even intuitive anymore, considering we now integrate over *negative* frequencies (what is a negative frequency?). So eventually, we do need to map the calculated responses from the "frequency domain" back into the "time domain", and that is what is eventually *observable* to us, and is literally "real".

Note: If we are talking about *steady repetitive* waveforms, we can set $s = j\omega$, and we will get back the well-known Fourier series decomposition.



Plant Transfer Functions

This is often called the “control to output” transfer function, since it is in effect the output voltage divided by the control voltage. It is a product of three *successive* (independent) gain blocks: the PWM comparator, the switching power section, and the output LC filter. Note that only in the Buck, do we have an actual LC post filter. In the case of a Boost or a Buck-Boost, there is something (the switch) connected between the L and the C (output cap). So strictly speaking, we cannot consider it as a separate (independent) LC gain block and just multiply all the successive gain blocks. However, using the “canonical model”, it can be shown that an *effective* LC post-filter (a separate gain block) can be visualized *for non-Buck topologies* too, provided we use the effective inductance $\underline{L} = L/(1-D)^2$. The important point to note is that this effective inductance varies with input voltage (for a given output) since D decreases as input raises, causing a higher effective inductance and thereby tending to make the loop responses more sluggish. In general, large inductors and capacitors need several more cycles to reach their new steady-state energy-levels, after a line or a load transient. This sluggishness, based on larger effective L, is of greater impact with VMC, because as we will see very shortly, when using CMC, the inductance is not even part of the plant transfer function to start with. With VMC it is.

We will go through several figures next, to sum up the Plant transfer functions for VMC and CMC.

Figure 3: This is a Buck with classic VMC. The DC gain (gain at low frequencies) changes as a function of V_{IN} in classical VMC, because V_{RAMP} – the amplitude of the sawtooth applied to the PWM comparator – is traditionally fixed. This causes a change in loop response characteristics with respect to input. Furthermore, the *line rejection* is not good under suddenly changing conditions. The reason is that a sudden change in line is not “felt” by the PWM comparator *directly*, and so it continues with the same duty cycle for a while. But any change in input voltage ultimately requires a change in the steady-state duty cycle (as per the steady DC transfer function equation of the converter: $D=V_O/V_{IN}$). So not changing the duty cycle quickly enough leads to an output overshoot or undershoot. The system has to “wait” for the output error to be sensed by the error amplifier, and that information to be communicated to the PWM comparator as a change in the applied “control voltage”. That eventually corrects the duty cycle and the output too, but not before some swinging back and forth (ringing) around the settling value. However, if we could just change the ramp voltage directly and *instantaneously* with respect to the input voltage, we would not need to wait for the information to return via the control voltage terminal. Then the line rejection would be almost instantaneous, and furthermore, the DC gain in *Figure 3* would not change with input voltage.

To implement the above, we would need to do the following:



$$V_{\text{RAMP}} \propto V_{\text{IN}}, \text{ so } \frac{V_{\text{IN}}}{V_{\text{RAMP}}} = \text{constant}$$

This is called “Line Feedforward”. We will see that CMC has similar properties very *naturally*, which was one of the reasons for its perceived superiority for a long time. Pure VMC, on its own, is certainly impaired, especially in this respect. But *VMC with Line Feedforward (when implemented in a Buck) actually offers superior line rejection to the one coming “naturally” from CMC.*

Figure 4: This is a Boost with VMC. Note that Line Feedforward is not practical here considering the complexity of the terms. We would need to somehow set V_{RAMP} proportional to $V_{\text{IN}} \times (1-D)^2$, based on the DC gain of $G(s)$. Also note the appearance of the Right Half Plane (RHP) zero. It also appears in the Buck-Boost. It is present for any duty cycle and for either VMC or CMC. In a “well-behaved” (left half plane) zero, the gain *rises* (or changes) by the amount “+1” at the location of the zero, and its phase *increases correspondingly*. With an RHP zero, the phase *falls* even though the gain *rises*, making this particular zero very difficult to compensate or deal with.

The existence of the RHP zero in the Boost and Buck-Boost can be traced back to the fact that these are the only topologies where an actual LC post-filter *doesn't* exist on the output. So even though we created an “effective” LC post-filter by using the canonical modeling technique, in reality there is a switch/diode connected between the actual L and C of the topology, and that is what is ultimately responsible for creating the RHP zero. The RHP zero is often explained intuitively as follows (in the earliest Unitrode App Notes): if we suddenly increase the load – the output dips slightly. This causes the converter to increase its duty cycle in an effort to restore the output. Unfortunately, for both the Boost and the Buck-Boost, energy is delivered to the load only during the switch-OFF time. So, an increase in the duty cycle decreases the OFF time, and unfortunately there is now a smaller interval available for the stored inductor energy to get transferred to the output. Therefore, the output voltage dips even further for a few cycles, instead of increasing as we were hoping. This is the RHP zero in action. Eventually, the current in the inductor does manage to ramp up over several successive switching cycles to the new level, consistent with the increased energy demand, and so this strange counter-productive situation gets corrected. Of course provided full instability has not already occurred! As mentioned, the RHP zero can occur at any duty cycle. Note that its location moves to a lower frequency as D approaches 1 (i. e., at lower input voltages). It also moves to a lower frequency if L is *increased*. That is one reason why bigger inductances are *not* preferred in Boost and Buck-Boost topologies.

The usual method to deal with the RHP zero is literally “pushing it out” to higher frequencies where it can't significantly affect the overall loop. Equivalently, we need to reduce the bandwidth of the open-loop gain plot to a frequency low enough that it just doesn't “see” this zero. In other words,



the crossover frequency must be set much lower than the location of the RHP zero. In effect, *the bandwidth and loop response suffers on account of the RHP zero.*

Figure 5: This is a Buck-Boost with VMC. As for the Boost, Line Feedforward is not a practical goal here. The RHP zero is also present here, though its location is a little different as compared to a Boost.

Figure 6: This is a Buck with CMC. We see the differences compared to VMC. The DC gain is *not* a function of input voltage (at least to a first approximation). The reason is that the PWM ramp is derived from the current ramp, and we know that the current ramp swing (ΔI) is a function of the voltage across the inductor during the ON-time, i.e. it depends on $V_{IN}-V_O \approx V_{IN}$. So in effect, there is *pseudo-line-feedforward* – not as perfect as we can implement by choice in VMC. Nevertheless, CMC does offer good line rejection, which was historically one of the key reasons for its wide popularity. However, we can see that the DC gain is a function of “R” – the load resistance. This causes a change in the loop characteristics with changes in load. But there is an interesting property shown in the diagram because the *location* of the pole also varies with load. Therefore the *bandwidth* remains unchanged with load (and line). That is actually the same as in VMC with Line Feedforward implemented.

We do see that CMC has a “single pole” at the resonant frequency of the load resistor and the output cap. In contrast, we saw in *Figure 3* that VMC has a double (two single) poles at the resonant frequency of the inductor and output cap. That by itself is not really an issue, because what it eventually means is that we need two zeros from the compensator to cancel out the double pole in VMC, but only one zero to cancel out the single pole of CMC. So the compensator can be *simpler* for CMC than for VMC. Typically that means we can use a Type 2 compensator (often based around a simple transconductance error amp) for CMC, whereas we usually need a more complicated Type 3 compensator for VMC. Other than that, what is the difference? The difference is that despite “cancellation” of the double pole arising from VMC, there can be a huge residual *phase shift* (rather, a huge back and forth phase swing) in the region around the cancellation frequency. This can lead to conditional stability issues, especially under non-linear (large) line/load disturbances. So CMC has somewhat more predictable and acceptable responses in general.

Figure 7: This is a Boost with CMC. It also has the troublesome RHP zero.

Figure 8: This is a Buck-Boost with CMC. It also has the troublesome RHP zero.

This almost sums up our overview of the key differences between CMC and VMC. There is one last issue as discussed next. Note that all along we have restricted ourselves to continuous conduction mode (CCM), mainly for simplicity sake. In any case, discontinuous conduction mode (DCM) is encountered only for much lighter loads, and further, in many modern synchronous topologies, we



Microsemi TN-206 Voltage-Mode, Current-Mode (and Hysteretic Control)

by Sanjaya Maniktala

may continue to stay by choice in CCM down to zero load (that is called "FCCM" or forced continuous conduction mode).

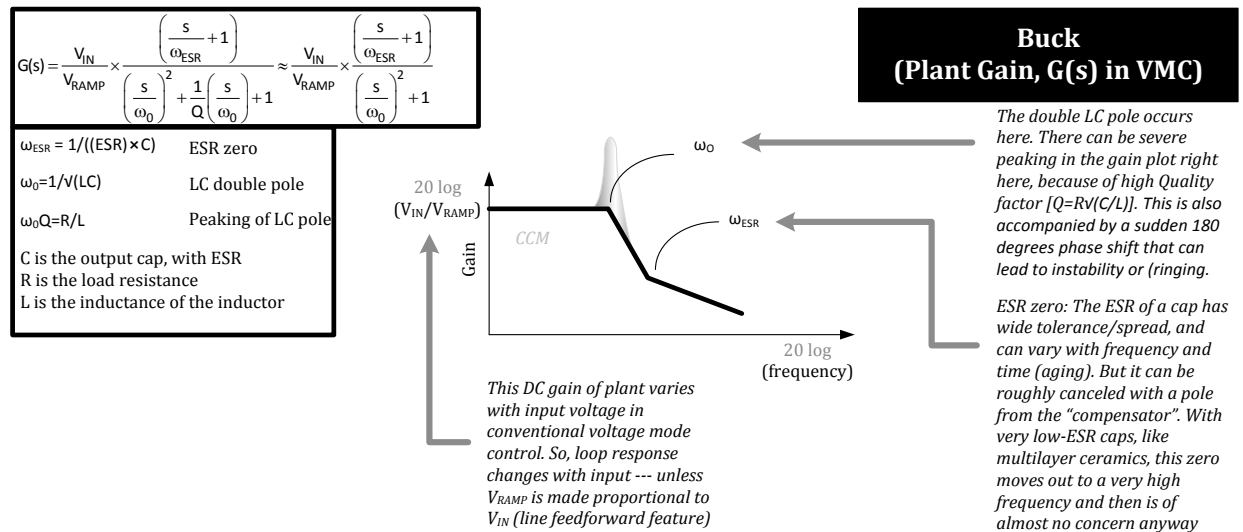


Figure 3: Plant Transfer Function for a Buck, Using VMC

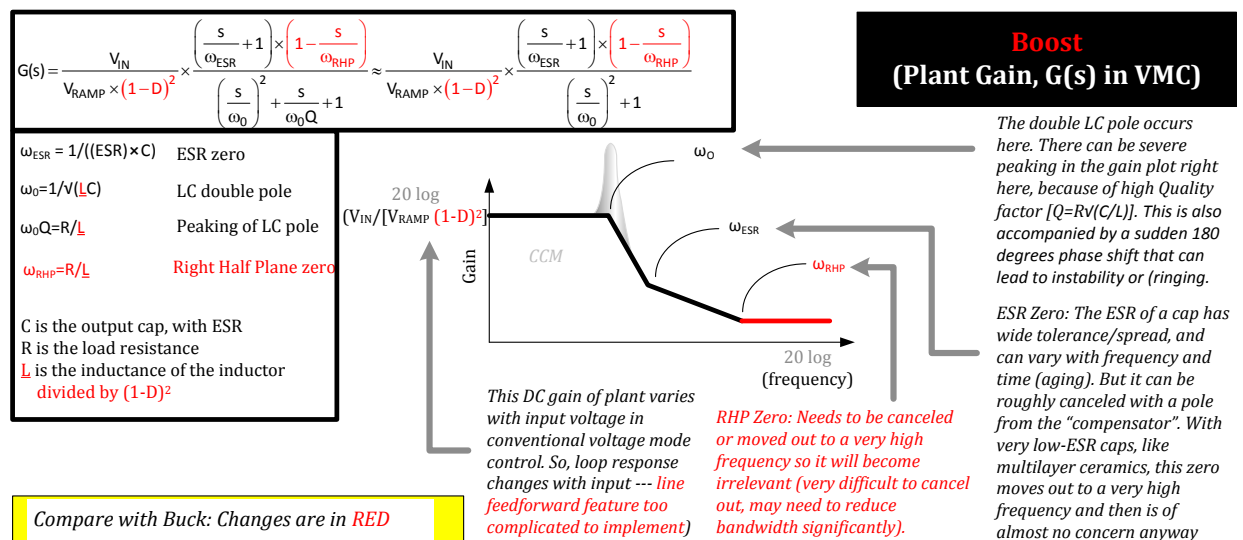


Figure 4: Plant Transfer Function for a Boost, Using VMC

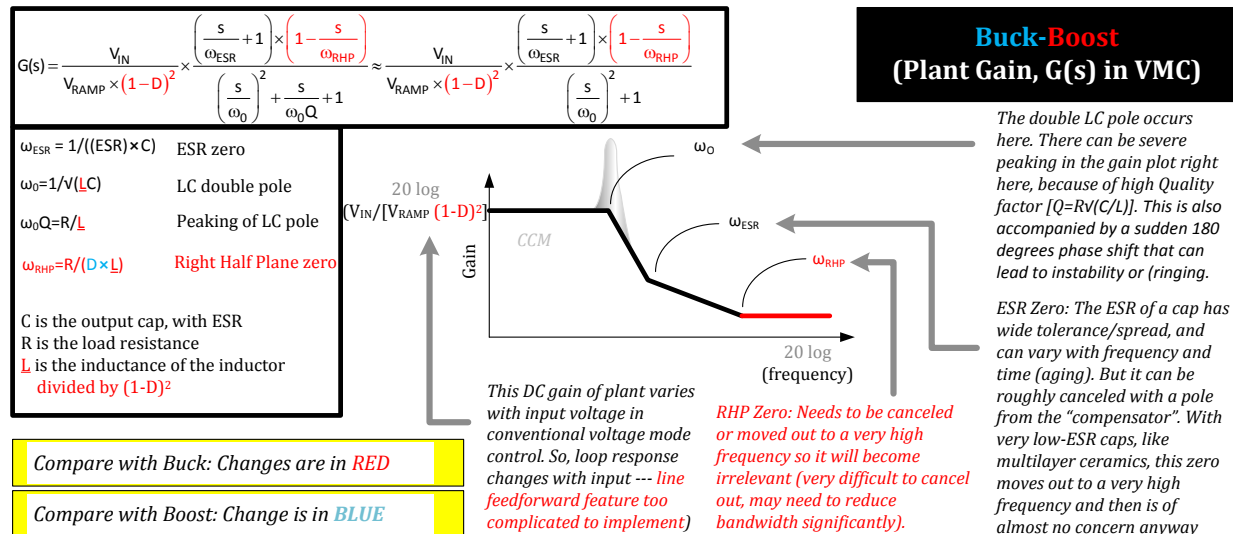


Figure 5: Plant Transfer Function for a Buck-Boost, Using VMC

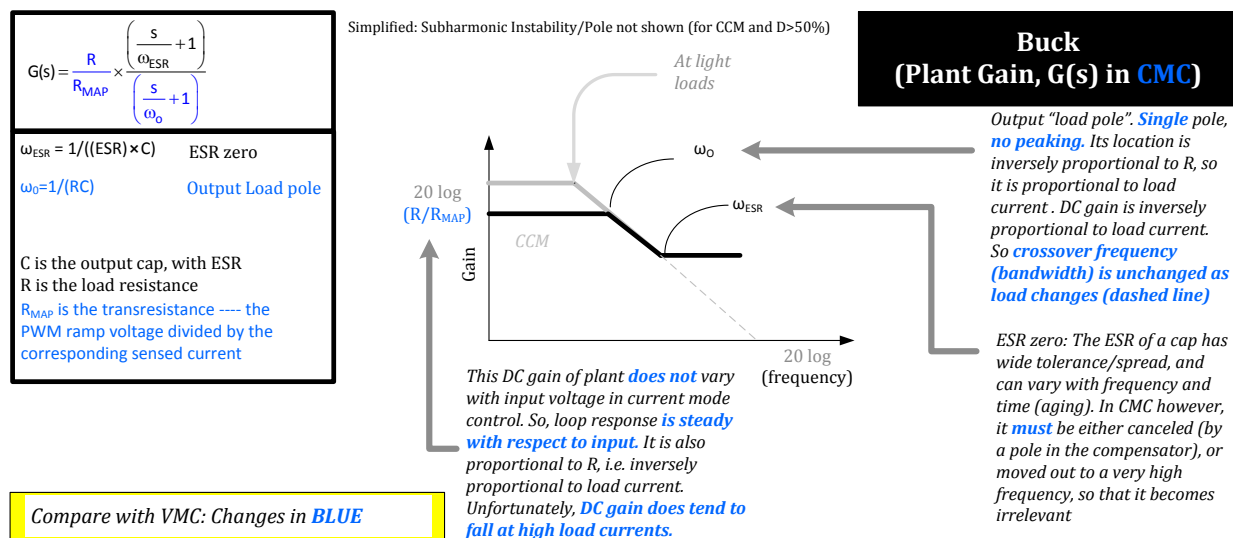


Figure 6: Plant Transfer Function for a Buck, Using CMC



Microsemi TN-206 Voltage-Mode, Current-Mode (and Hysteretic Control)

by Sanjaya Maniktala

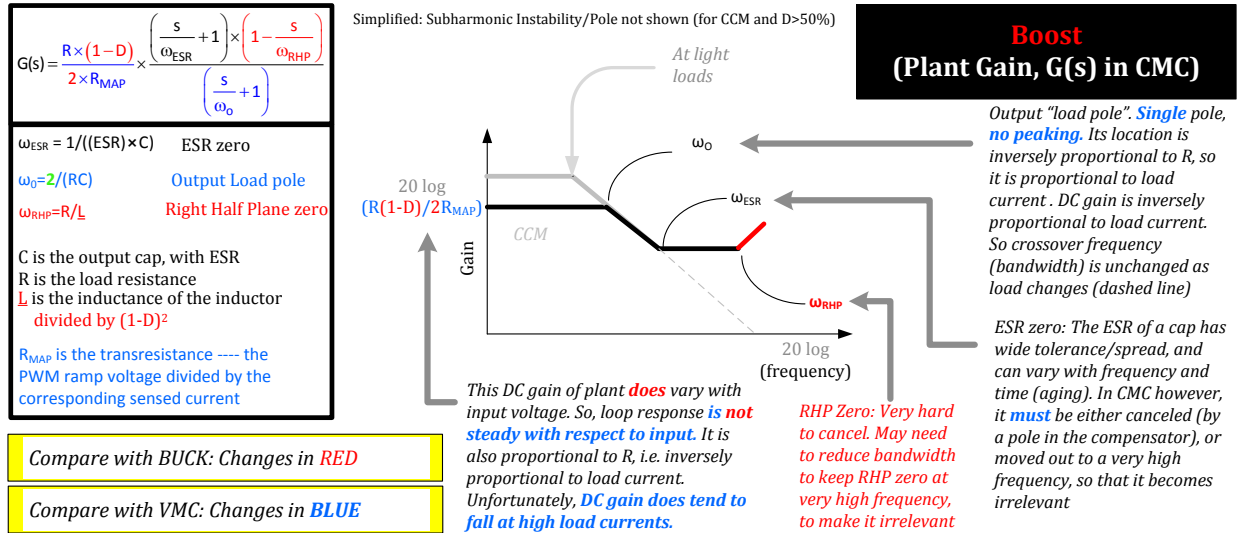


Figure 7: Plant Transfer Function for a Boost, Using CMC

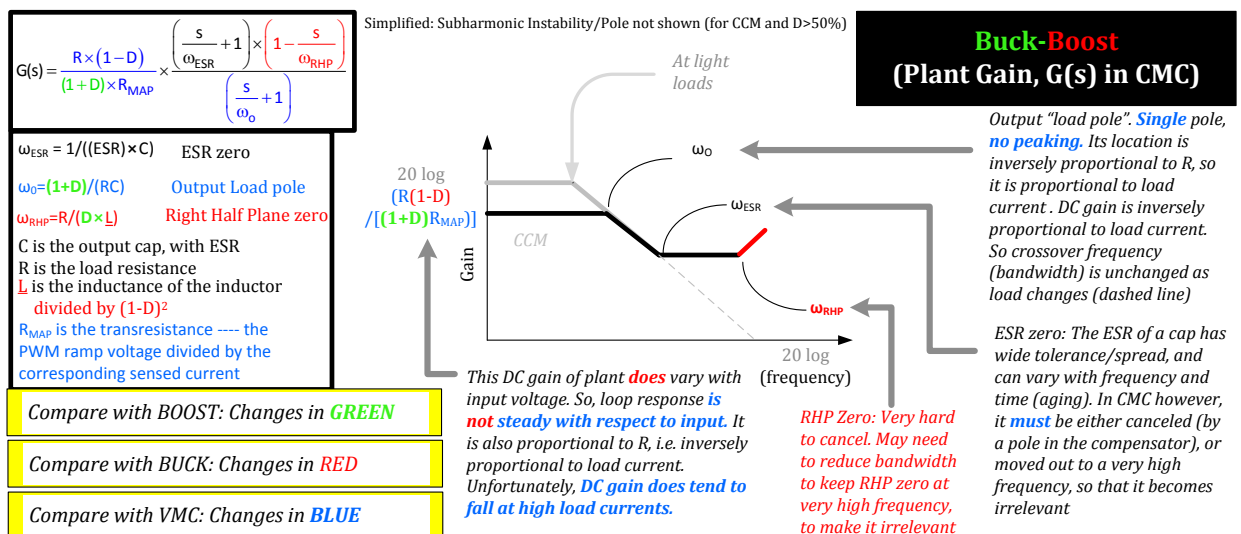


Figure 8: Plant Transfer Function for a Buck-Boost, Using CMC



Subharmonic Instability in CMC

The model we used for the CMC figures was very simplified. We ignored something that is actually very relevant to understanding the *practical limits* of CMC. The phenomena of *subharmonic instability* in CMC is well-known and represents one of the most major drawbacks of CMC, one that can easily force us into setting a much lower bandwidth for CMC than we had thought initially (based on incomplete models or by reading very early, and rather optimistic, Unitrode App Notes).

A converter already in the midst of this particular (subharmonic) instability will usually show no outwardly symptoms of anything amiss at all, especially in *steady state*. But under sudden line or load disturbances, we will notice highly impaired and sluggish loop response. The Bode plot will not look comprehensible. If we connect a scope to the switching node, we will see a pattern of one wide pulse (almost max duty cycle) followed by a very thin pulse (almost min duty cycle). Now, the repetition frequency of this pattern is not the switching frequency, but *half* the switching frequency (two pulses in every repeating pattern). Therefore, subharmonic instability is often called “half-switching frequency (or $f_{sw}/2$) instability”.

Note: All cases of one wide pulse followed by one narrow pulse are not necessarily related to subharmonic instability of the type under discussion here. The same pattern can be caused by a leading edge noise spike, causing early termination of one pulse which is then automatically followed up by the loop as one (or two) wider “make-up” pulses.

Note: Repeating patterns of three (or more) pulses are most likely noise artifacts or traditional instability, not subharmonic instability

Well before a system actually enters this irrecoverable subharmonic instability state, we can ask what are the *symptoms*, or signs, of *impending* subharmonic instability? Because if we recognize that, perhaps we can avoid subharmonic instability *before* it happens.

To answer that, first of all we must remember that subharmonic instability is entered only in continuous conduction mode (CCM) at duty cycles greater than 0.5, and of course only when using current-mode control (CMC).

Being forever “practical”, let's suppose we take the Bode plot of any current mode controlled converter – one that has *not* yet entered this wide-narrow-wide-narrow state. We will discover a *mysterious peaking* in the gain plot at exactly *half the switching frequency* (very similar to the gain peaking for VMC at the LC double pole, see *Figure 3*). This is the “source” or origin of future subharmonic instability. Subharmonic instability is therefore nowadays modeled as a *complex pole at half the switching frequency*. Quite like the LC double pole of VMC, it too has a certain Q (Quality Factor), or “peaking” that we need to limit.



Note that based on sampling theory, we never try to set the crossover frequency (loop bandwidth) higher than half the switching frequency. So in effect, this subharmonic pole will always occur at a frequency *greater* than the crossover frequency. But it could in fact be *uncomfortably close* to the crossover frequency, especially if we try to push the crossover frequency (bandwidth) higher and higher, going from, say, 1/6th to 1/3rd the switching frequency. This half-switching frequency pole is ominous because of the fact that if it peaks too much *it can end up* causing the gain plot to intersect the 0dB axis *once again*, just past its actual (first) crossover point. Even though this is an unintended crossover, any phase reinforcement at any crossover point can provoke full instability. A system will then quickly transgress and settle down into an *irrecoverable alternate pulsing* (fsw/2) pattern. But if we are not so aggressive in our bandwidth goals, we are much better poised to see only minimum effects coming from this fsw/2 pole. Keep in mind however, that the effect of this pole on the *phase angle* may start at a *much lower frequency*.

The subharmonic instability pole has a certain “Q” that we can calculate., It has been shown by actual experiments that a Q of *less than 2* usually allows stable conditions. A Q of 1 is preferred by conservative designers, and though that choice does quell subharmonic instability even more firmly than Q=2, it does lead to a *bigger inductor*, and often to a rather non-optimum inductor current ripple ratio ($\Delta I/I$) --- of less than $\pm 20\%$. Alternatively, we need to apply greater “slope compensation”. But too much slope compensation is akin to making the system more and more like VMC, and pretty soon, especially at light loads, the double LC pole of voltage mode control will reappear potentially causing instability of its own (since we didn’t plan for it).

To keep the Q of this subharmonic pole to less than 2, we need to set the inductance of the converter to higher than a minimum value. The equations for that are

$$L_{\mu H} \geq \frac{D - 0.34}{\text{Slope Comp}_{A/\mu s}} \times V_{IN} \quad (\text{Buck})$$

$$L_{\mu H} \geq \frac{D - 0.34}{\text{Slope Comp}_{A/\mu s}} \times V_O \quad (\text{Boost})$$

$$L_{\mu H} \geq \frac{D - 0.34}{\text{Slope Comp}_{A/\mu s}} \times (V_{IN} + V_O) \quad (\text{Buck-Boost})$$

So basically, with CMC, we may need to do one or more of the following:

- Add slope compensation (but not so much that we just get unplanned-for VMC)
- Increase the Inductance (per the limits presented above)
- Reduce loop bandwidth (almost down to the conservative expectations for VMC)



Conclusions on CMC versus VMC

We have realized that when we take a closer look, CMC and VMC are just alternative ways of achieving loop stability. The type of compensation scheme we need may be simpler for CMC (Type 2) than for VMC (Type 3). But then, CMC also needs slope compensation, and so on. Pros and Cons as usual.

CMC also suffers from high PCB sensitivity, since we usually depend on a small sensed current to generate the voltage ramp applied to the PWM comparator. Since a lot of noise is generated when the switch turns ON, we usually need to introduce a “blanking time” to avoid triggering the comparator for at least 50-200ns after the switch turns ON. This unfortunately leads to indirectly establishing a minimum ON-time pulse width of also 50-200ns, and the corresponding minimum duty cycle can play havoc with high-voltage to very low-voltage down-conversion ratios, especially with high switching frequencies. In comparison, VMC is inherently more robust and noise-resistant.

A modern preference seems to be in the direction of “VMC with Line Feedforward”, just as Bob Mammano implicitly predicted in DN-62. But another recently emerging major thrust is actually heading towards hysteretic control, as we will shortly see.

In closing, we present an overview of compensator design strategies in *Figure 9*.



Microsemi TN-206 Voltage-Mode, Current-Mode (and Hysteretic Control)

by Sanjaya Maniktala

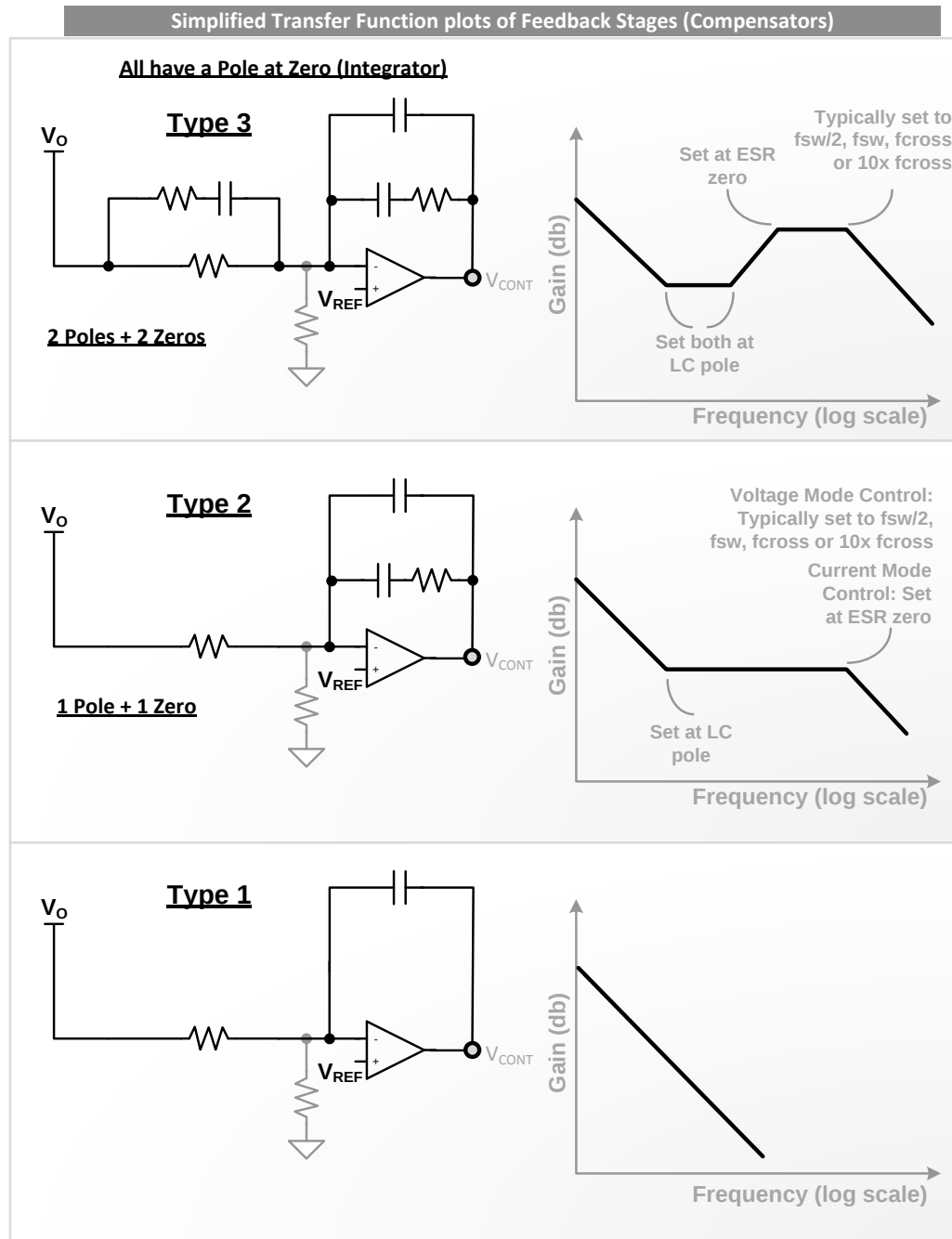


Figure 9: Overview of Compensator Design for VMC and CMC



Hysteretic Control: Energy on Demand

Looking at *Figure 1* again, we see that the basic way a PWM comparator works is by creating ON-OFF pulses from the intersection of two voltage profiles at its input terminals: one *steady voltage* level (the control voltage), and another sawtooth voltage profile (ramping up and down). We can well ask why not applying the reference voltage directly as the “smooth” voltage level (instead of the error voltage) on the other terminal, and use a sawtooth based on the inductor current (as in CMC)? This is shown in *Figure 10* and in *Figure 11*. Note that as shown in *Figure 12*, this is not really a “PWM comparator” anymore, at least not in the sense we were used to so far. It is now a “hysteretic comparator” that terminates the ON-pulse if the ramp voltage exceeds the reference voltage by a certain amount Δv (“ $\Delta HYS/2$ ”), and turns the switch back ON when the ramp falls below a certain threshold slightly lower than the reference voltage ($-\Delta v$). It is therefore often called a “bang-bang” regulator. If there is a sudden line or load transient, it can react by either turning OFF completely for several pulses in succession, or by turning ON fully. Therefore its transient response is excellent – “Energy on Demand” in effect. Or “Energie vom Faß” as the Germans would perhaps like to call it.

Early forms of bang-bang regulators have existed for decades, based on SCRs (silicon controlled rectifiers) or bipolar transistors, but *without inductors*. This ancient technique has been literally re-invented in modern switching power conversion, and offers tremendous advantages as can be confirmed on the bench. The bandwidth of the loop response is close to the switching frequency itself. There is no feedback or compensator to design, nor poles and zeros to manipulate. But the mathematical models of hysteretic controllers are still very complicated and just evolving. This however has not stopped designers from trying to eke out the full commercial advantage of hysteretic control. One of the biggest advantages is that because there is no clock and no error amplifier, nor any compensation circuitry, the quiescent current (I_Q , zero load, but still switching), is very low (typically less than $100\mu A$). This makes the hysteretic converter very suitable for modern battery-powered applications in particular. In the world of cell phones and tablets, hysteretic has started leading the way.

Hysteretic control does have its limitations. Because there is no formal clock, it is hard to assure constant frequency. There can also be a lot of erratic pulsing, usually accompanied by unacceptable audio noise (squealing), and also unpredictable EMI (and audio). The way to avoid erratic pulsing is to ensure the ramp waveform applied to the hysteretic comparator is an exact replica of the actual inductor current. This way we get a chicken and egg situation where *the duty cycle created by the comparator is exactly what the system naturally demands under the given line and load conditions*. Then there are no missed pulses accompanied by audible low-frequency harmonics. The way to adjust the frequency to an acceptable constant level is by *symmetrical* variation of the comparator



thresholds, as indicated in *Figure 12*. If we do not do this “symmetrically”, there will be a resultant DC offset, causing a drift or off-centering in the output voltage.

Another way to try obtaining hysteretic control with almost constant frequency is to use a “COT” (constant ON-time controller). We know that (for a Buck):

$$D = \frac{T_{ON}}{T} = T_{ON} \times f = \frac{V_O}{V_{IN}}$$

Therefore

$$f = \frac{V_O}{(V_{IN} \times T_{ON})}$$

In other words, if we fix the ON-time of the hysteretic converter, but also make that ON-time inversely proportional to the input, we will get a constant “f”. That is the underlying principle of the COT (Buck) converter. Note that the COT converter has a pre-ordained ON-time, so there is no upper comparator threshold required anymore. But for the same reason, we can get some DC offset here.

In the case of a very-high-voltage to low-voltage conversion (like 48V to 3.3V), we know that in traditional converters we get rather limited by the minimum (practical) pulse width of the converter, especially in the case of CMC (due to the blanking time issue as discussed earlier). But in the COT Buck, by fixing the minimum ON-time, we in effect not only lower the frequency, but also achieve smooth down-conversion without any unexpected overshoots during line and load transients as in traditional control methods.

In a similar manner it can be shown that a constant OFF-time will give a constant frequency when applied to a *Boost*. We'll remind ourselves once more of the intuitive reason for the RHP zero: under a sudden load demand the output dips momentarily and therefore the duty cycle increases. But in the process, the OFF-time decreases. Since in a Boost (and Buck-Boost) energy is delivered to the output only during the OFF-time, a smaller OFF-time leaves less time for the new energy requirement to be met, which temporarily causes the output to dip even further before things get back to normal. So we intuitively realize that fixing a certain minimum OFF-time will help in this case. And that is in fact true: *the RHP zero is not present when operating the Boost in constant off-time mode*. And we can get constant-frequency operation too, by setting $T_{OFF} \propto V_{IN}$.

For a Buck-Boost, the relationship for achieving constant frequency is too complicated to implement easily, without sacrificing the key advantages of hysteretic controllers: simplicity and low I_Q . So we will go past this roadblock.



Autotuning

Nowadays, system houses want the flexibility, not only of voltage margining via I²C control, but also the ability to change switching frequencies, say, over a typical $\pm 50\%$ range. The idea is to be able to avoid certain frequencies after a pre-release EMI/audio scan. In general, beat frequencies also need to be avoided in cases of multiple free-running regulators switching at a frequency fairly close to each other. It may be too late to return to the drawing board and start changing components only to find that line/load transient response has been affected. So the concept of autotuning is gaining ground. Design houses are spending a lot of time starting out with CMC or VMC and learning to re-position the poles and zeros automatically in case the frequency is changed. Hysteretic control, in particular COT, offers the advantage of having no traditional loop compensation components. Being “energy on demand”, it has inherent autotuning capabilities. Microsemi’s future hysteretic controllers are thus being designed to provide effective support for autotuning features. In contrast, digital methods tend to create a large I_Q and silicon area requirement.

Microsemi Proprietary Hysteretic Control

Microsemi has a proprietary hysteretic engine that creates an *artificial ramp* that mimics the inductor current, to help overcome PCB layout sensitivity issues. It also changes the hysteretic thresholds *symmetrically* to achieve *constant frequency with no DC offset*. In *Figure 13* we have presented the Microsemi hysteretic parts currently available, or in the process. Please contact your specific region’s Microsemi sales representative for more information.

This concludes a brief summary of the pros and cons of hysteretic control. See Table 1 for a summary of pros and cons of the various control methods.



Microsemi TN-206 Voltage-Mode, Current-Mode (and Hysteretic Control)

by Sanjaya Maniktala

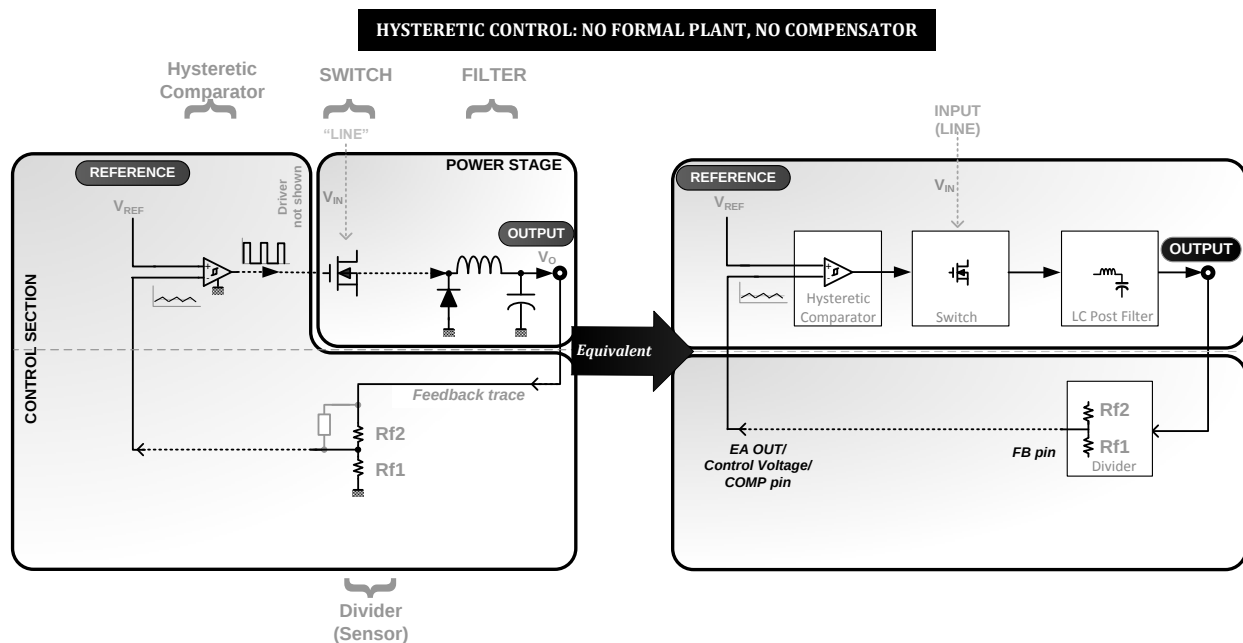


Figure 10: Functional Blocks of the Hysteretic Converter

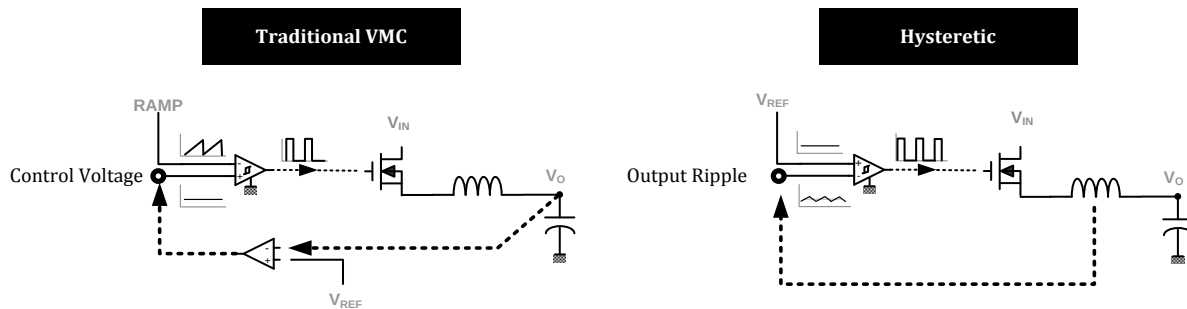


Figure 11: Basic Changes to Achieve Hysteretic Control



Microsemi TN-206 Voltage-Mode, Current-Mode (and Hysteretic Control)

by Sanjaya Maniktala

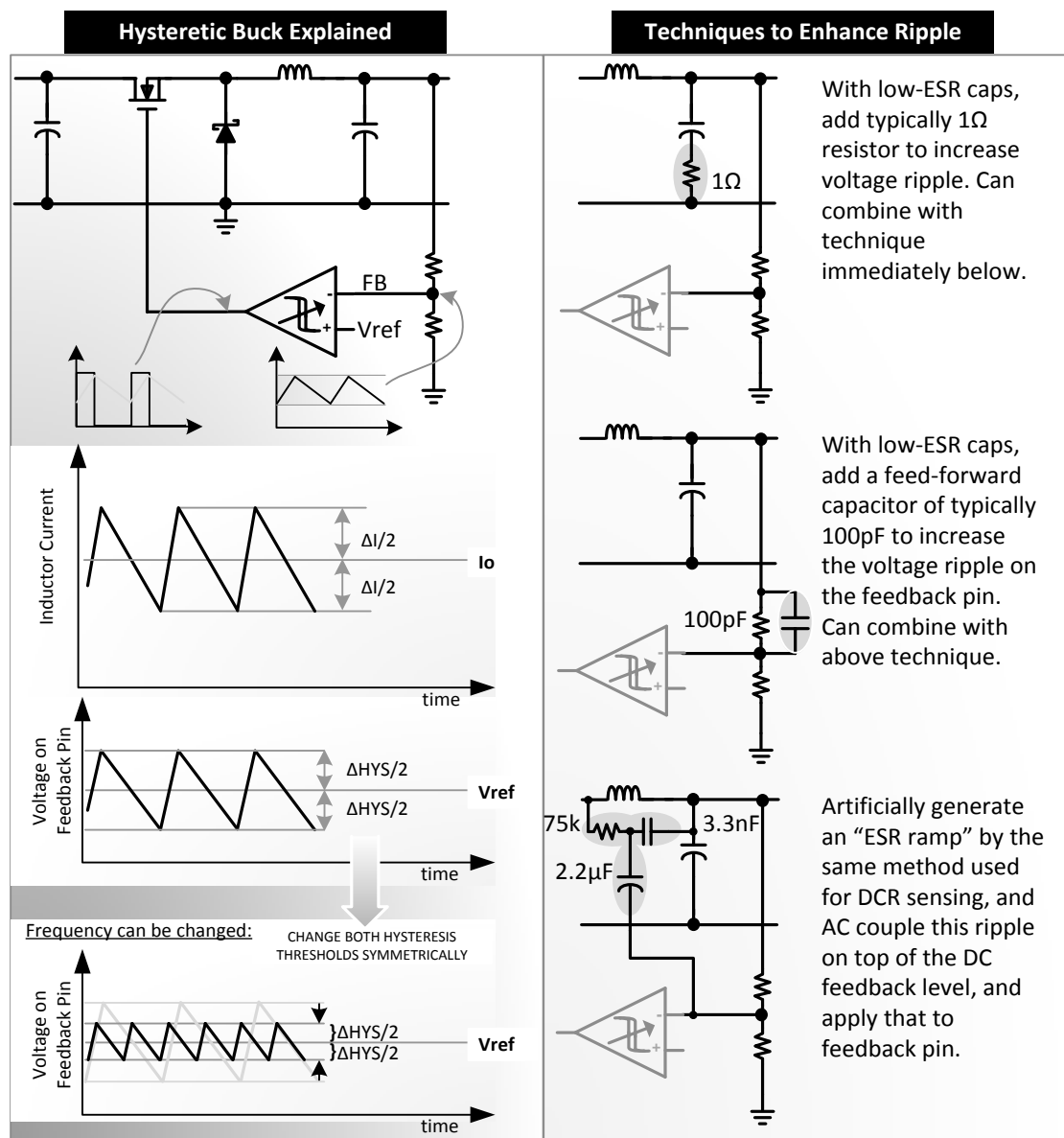


Figure 12: Hysteretic Control Explained in a Simplified manner



Microsemi[®] TN-206 Voltage-Mode, Current-Mode (and Hysteretic Control)

by Sanjaya Maniktala

Part	I _{out}	V _{in}	F _{sw} (Mhz)	Serial Interface	PSM	Package
LX7165	5A	3 to 5.5V	2	I2C	Yes	WLCSP 1.8x2
LX7169	3A	3 to 5.5V	3	No	No	DFN 3x3.5
LX7176	3A	3 to 5.5V	2	No	Yes	DFN 2x2
LX7167	2A	3 to 5.5V	3	No	Yes	DFN 2x2

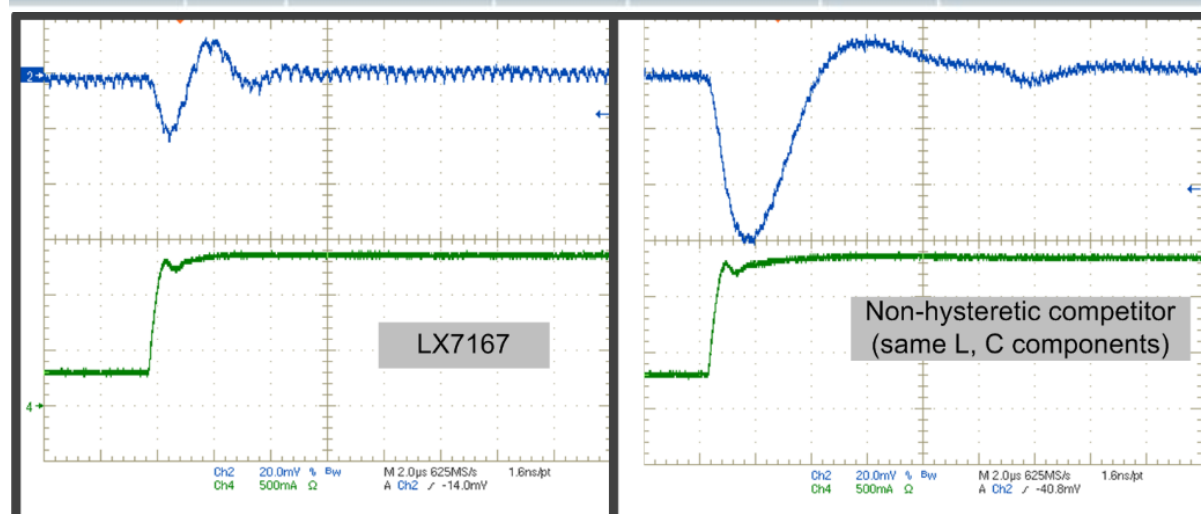


Figure 13: Hysteretic Switcher offerings from Microsemi



Summary of Pros and Cons of the Different Control Techniques

	CMC	VMC	Hysteretic
<i>Rejection of Line Disturbances (Dynamic Line response)</i>	Good (Inherent)	Very Good (with Line Feedforward)	Excellent (Inherent)
<i>Rejection of Load Disturbances (Dynamic Load response)</i>	Good (Constant Bandwidth)	Good	Excellent (Inherent)
<i>Constant Frequency</i>	Excellent	Excellent	Poor, need to vary hysteresis band, OR use Constant ON-time (Buck), OR use Constant OFF-time (Boost)
<i>Predictable EMI</i>	Excellent	Excellent	OK (with above COT techniques)
<i>Audible Noise Suppression</i>	Excellent	Excellent	OK (with above COT techniques)
<i>Extreme Down Conversion (Buck)</i>	Poor	Good	Excellent, with COT techniques
<i>Insensitivity to PCB Layout</i>	Poor	Excellent	Good, with Artificial Ramp Generation, otherwise poor
<i>Excellent Stability of Loop Responses (Tolerances and Long-term Drifts)</i>	Excellent	Good	Fair
<i>Simplicity of Compensation</i>	Good	Poor	Excellent
<i>I_Q (Quiescent Current)</i>	Good	Poor	Excellent
<i>Loop Stability with use of Output Ceramic Caps</i>	Excellent (with Type 3 compensation)	Very Good	Good, with Artificial Ramp Generation, otherwise poor
<i>Autotuning</i>	Complex	Very Complex	Inherent

Table 1: Summary (Voltage Mode versus Current Mode versus Hysteretic Control)



Microsemi TN-206 Voltage-Mode, Current-Mode (and Hysteretic Control)

by Sanjaya Maniktala

Revision History

Revision Date	Level /	Para. Affected	Description
0.3 / 08-July-2012		-	
0.4/18-July-2012		Added Figure 13	Description of Autotuning and Microsemi Hysteretic ICs
0.5/28-July-2012		Autotuning/Table 1	Table 1 has an added row

© 2012 Microsemi Corp.

All rights reserved.

For support contact: **sales_AMSG@microsemi.com**

Visit our web site at: **www.microsemi.com**

Catalogue Number: **DCDC_TN_206**