

FREQUENTLY ASKED QUESTIONS

February 16, 2017

Content Questions

If we have poles, shouldn't they manifest themselves on Bode plots as large valued $\hat{H}$?

Sometimes you do get large $|\hat{H}|$ associated with poles, but both poles and zeroes matter for the frequency response.

The poles manifest themselves as very large (in fact infinite) $|\hat{H}(\hat{s})|$ on the complex $\hat{s}$ plane. However, the Bode plots we've been drawing show the value of $|\hat{H}(\omega)|$ only as a function of ω , the (positive) imaginary part of the complex frequency, i.e., the value as you go up the imaginary ω axis on the $\hat{s}$ plane. This regime is relevant for sinusoidal frequencies (going off this imaginary axis corresponds to transients, which matters for some applications, but we won't be treating it much).

The $|\hat{H}(\omega)|$ vs. ω Bode plot corresponds to the product of distances to zeroes from a point on the positive ω axis, divided by the product of the distances to poles from this point. So there might be a big denominator at some ω , but this could be canceled by an also-large numerator.

Examples: the single-pole low-pass filter does have its largest value at $\omega = 0$, which is the closest you can get to the pole (which lives on the negative real axis). The single-pole high-pass filter has one pole and one zero, and these both get larger (and their ratio approaches 1) as you slide up to very high ω on the imaginary axis.

Why do we know that ω_1 and ω_2 are corner frequencies (for the RLC filter examples)? Why do we switch to ω_1 and ω_2 instead of ωRC and such?

When you write the denominator of all the *RLC* examples as $(1 - \frac{\omega^2}{\omega_1\omega_2} + j\frac{\omega}{\omega_1})$, you can see that ω_1 and ω_2 divide different regimes by eyeballing the expression for different values of ω .

We assumed $\omega_2 \gg \omega_1$. Then the $\frac{\omega^2}{\omega_1\omega_2}$ term will be guaranteed to be very large, and $|\hat{H}| \propto \omega^{-2}$, for $\omega \gg \omega_2$. (Since ω_1 is smaller than ω_2 , then if $\omega \gg \omega_2$, then $\omega^2 \gg \omega_1\omega_2$). When does that $\propto \omega^2$ behavior in the

denominator take over? When $\frac{\omega^2}{\omega_1\omega_2} \sim \frac{\omega}{\omega_1}$, so at $\omega = \omega_2$. So that's one corner frequency.

Now what if ω is very small? Then the ω^2 term can be ignored. Then it looks like a single-pole filter. The behavior transition in the denominator from 1 to $\propto \omega$ is when $1 = \frac{\omega}{\omega_1}$, so when $\omega = \omega_1$.

What if $\omega_2 = \omega_1$? Then both breakpoints are in the same place.

But what if $\omega_1 \gg \omega_2$? Then it gets a bit less obvious, and only extreme behaviors are clear.

As for the second question: the formulation in terms of ω_1 and ω_2 is just a convenient one for understanding the behavior in different frequency regimes, as these quantities correspond to corner frequencies.

What is the physical (conceptual) importance of corner frequencies?

They are the transition frequencies, with values depending on the components of the network, between different regimes of frequency response. I'm not sure there's anything deeper than that, in general, although for specific networks you can think of them physically. For example, for the low-pass filter, the corner frequency is $\omega_c = 1/(RC)$, which is 1 over the time constant of the network; for frequencies larger than that, capacitor never fully charges up so you never develop the full voltage on the output (so the sinusoidal signal is attenuated); for frequencies smaller than that, you charge and discharge the capacitor up and down all the way in each cycle and output follows input (see also some of the answers in FAQs 8 and 9).

How are buffer and op-amp circuits used in real life?

Op-amps (and buffers made with op-amps) are actually pretty ubiquitous in real-life circuits. We will be seeing more about them and applications of them towards the end of the course. An op-amp is a kind of basic amplifier (and a buffer is a unity-gain amplifier). There are many reasons you might want to increase the voltage amplitude in a circuit. A canonical example is a sound-amplifying circuit: when you turn up the volume on a music-playing device, you are increasing the gain of an amplifier. An example from my own field of physics research is an amplifier that turns tiny signals from particles into more robust signals which are easy to digitize. In practice, amplifier

frequency response matters quite a bit, as you often want to suppress noise or shape a signal in a particular way.

Buffers are commonly used to avoid loss of amplitude when connecting one device to another (as we saw in the example of sequential filters).

Real life op-amps have a lot of imperfections, and actual circuits will usually use op-amps together with other circuit elements in clever ways to achieve specific performances.

Why does the amplifier act like a low-pass filter?

In real life, any device has capacitance and resistance. Typical real-life amplifier devices have an equivalent circuit looking approximately like that of a low-pass filter (e.g. capacitance at the output, and equivalent resistance). Not every amplifier need have this response though; specific circuits can be devised that have different frequency responses.

Do amplifiers have any inherent impedance $\hat{Z}$?

Yes, they do. We will often idealize them as having infinite input impedance (i.e., they draw no current) and zero output impedance (i.e., they can supply infinite current). However real amplifiers have non-infinite input and non-zero output impedance. We'll be seeing more of this later.

How exactly does the feedback stabilize the amplifier?

One way to see it is from the math: from $\hat{H}(\hat{s}) = \frac{\hat{A}}{1 + \hat{A}(\hat{s})\hat{F}(\hat{s})}$, for $|\hat{A}\hat{F}| \gg 1$, $|\hat{H}| \sim \frac{1}{|\hat{F}|}$, which is independent of $|\hat{A}|$, so quite insensitive to any variation of the value of $|\hat{A}|$.

But here's a qualitative way of thinking about it: an op-amp gives an output voltage proportional to the *difference* between its inputs, by a large amplification factor. What the feedback network does is to send back a fraction of the output to the input. The amplifier then sees a smaller difference between the inputs... so it adjusts the output to be smaller. A fraction of this smaller output then gets sent back to the inputs again, and once again the output adjusts... this process keeps happening until the fractional voltage fed back from the output gives a difference between the inputs that no longer results in a change in the output. If the fed-back voltage at the input, is, say, $F = 0.1$ of the output, then the actual voltage at the output is $V_{\text{out}} = 10 V$

for an input difference of $V_{\text{in}} = 0.1 \text{ V}$, so the gain is $G = 1/F = 10$. For a real op-amp this all happens very fast and it all comes to equilibrium very quickly.

How does the idea of “feedback” get reflected in $\hat{F}$? Is there analogy of the function of $\hat{F}$ to the thermal control example?

Well, the most general concept of feedback is as follows: you take information from the output of a system and feed it back into the input to adjust the output. If it’s “negative feedback”, an output value is used in the input to adjust the output negatively (reduce an increase, increase a reduction). If it’s “positive feedback”, the output value is used to adjust the output positively (make something positive more positive, or something negative more negative).

In the thermostat example, you might measure the temperature at the output; if it’s higher than the target, you send a signal to the input of your thermal control unit to reduce the temperature; if the temperature is lower than the target, then you send a signal to the input to increase the temperature. The “feedback box” implementing $\hat{F}$ is then a somewhat complicated thing (maybe a thermistor, some logic to send a control voltage to the input or something), rather than just the simple voltage divider that we had in our example. But it’s still performing an $\hat{F}$ function of converting an output into some control value for the input.

Is there a ceiling on gain with feedback, even with freedom to choose $\hat{F}$?

In an idealized system you can get any gain you want by choosing $\hat{F}$ for $\text{gain} \sim \frac{1}{F}$, so long as $|\hat{A}\hat{F}| \gg 1$. You can’t get a gain larger than $|\hat{A}|$ by this method (but typically $|\hat{A}|$ is quite large, like 10^6).

Are amplifiers with smaller gain more stable than amplifiers with larger gain? If they are, do you still need negative feedback?

Hmm, I’m not sure that amplifiers with smaller gain are *necessarily* more stable, but I think it’s often true that large-gain amplifiers will be somewhat unstable... a large gain means that you get a big change at the output for a small change at the input, which means that small effects at the input

could make things change a lot. (The specific properties of amplifiers we'll see depend on solid state physics; we'll cover this later in a bit more detail.)

In our practical applications though, negative feedback for amplifiers is ubiquitous. We'll pretty much always be using devices that either have negative feedback built in to a package (so it's invisible to you, the user), or else if you are making an amplifier from "scratch" (i.e., from op-amps) you will nearly always use some kind of feedback circuit.

Can you explain the voltage divider in the feedback network?

In the amplifier network shown in the handout, a simple voltage divider with resistors fills the feedback box shown in the generic negative-feedback diagram. If you imagine a power supply $\hat{V}_1$ connected to the + and - inputs of the amplifier, $\hat{V}_4$ in the generic diagram corresponds to $\hat{V}_B$ in the amplifier picture ($\hat{V}_B$ is at the point between R_f and R_G). $\hat{V}_4 = \hat{F}\hat{V}_3$ then corresponds to $\hat{V}_B = R_G/(R_f + R_G)\hat{V}_{\text{out}}$, so $F = R_G/(R_f + R_G)$. You can then choose resistors to make $F = 0.1$ (or some other value < 1); the reciprocal of it will be your closed-loop gain.

Are there any real world applications of amplifiers with negative feedback?

Yes, practically every real-world amplifier uses some kind of negative feedback (although this might be invisible to the user, if the amplifier comes in some package).