

Prediction, Control, and the Regulating Mind

A Proposed Model of Regulation, from Outside the Field

Dr. Greg Moody

August 17, 2026

1. What I'm claiming, and what I'm not

Before I was a licensed psychotherapist and ran businesses and did martial arts and blah blah blah... I was an aerospace engineer. Guidance, navigation, control. Feedback loops with an estimate on one end and a correction on the other, and a short list of ways they fail.

I'm not a research psychologist. I don't run a lab, and I haven't spent thirty years inside this literature. A document from me proposing a model of psychology reads like somebody wandering in from another building to announce he's figured the place out. Fair enough. That isn't what this is.

The claim is smaller than the title makes it sound.

Two bodies of work already exist. One says the mind is a control system defending targets. The other says the mind is a prediction machine guessing at the world. Each one has a hole exactly where the other has something solid, and as far as I can find nobody has bolted them together to see what falls out of the join.

When you do, you get a handful of parameters that come apart cleanly. Those parameters predict different things about who gets better with which kind of help. That's the argument.

Why would somebody outside the field notice a join? Not because he's smarter. Because from inside either one, the other is somebody else's vocabulary. I spent years in one and then years in the other, so the fit was the first thing I saw instead of the last. Most cross-field ideas are wrong. This one might be too, and I can't tell that from where I'm standing.

Which is why Section 9 is in here. It lists, up front, the findings that would show this is wrong. If you read one part skeptically, read that one. I'd rather somebody knocked this down than nodded at it, and the people who could knock it down are the ones who've spent a career in one half of it.

A note on scope. I worked this out on a hard case first, autism and repetitive behavior, and there's a companion pair of documents covering that in full, including a long look at how thin the evidence there actually is. None of it is repeated here.

2. The mind as a control system

Cybernetics says the mind is a regulator. It has targets, it senses how far it is from them, and it acts to close the gap.

The vocabulary is small and worth having.

Set point. The target value the system defends. **Sensor.** What reports the current state. **Error signal.** The gap between the two. **Controller gain.** How hard the system acts on that gap. **Damping.** Whether corrections settle or overshoot. **Delay.** How stale the information is by the time it drives a correction.

A word on **cybernetic**, or **cybernetics**, since it has been dragged off in the wrong direction. It has nothing to do with cyborgs, and it predates them. It comes from the Greek *kybernetes*, the steersman on a ship, and Norbert Wiener took it in 1947 for the study of how any system steers itself by feedback. A man at a tiller, a thermostat, a cruise control, a body holding its temperature. The system senses where it is, compares that to where it means to be, and corrects. Wiener chose the word because the first serious paper on the subject was Maxwell's 1868 study of *governors*, the spinning weights that hold a steam engine at speed.

The clinically interesting move, and the one this whole document runs on, is that **an emotion is the error signal in a behavioral control system**. Hunger is what being below a set point feels like from the inside. Fear is what it feels like to have too much estimated danger against your tolerance for it. And **error** here doesn't mean bad. It's a signal that something is off, that's all, and it comes out as a feeling.

The clearest recent statement of this is a self-published series called *The Mind in the Wheel* by a group writing as Slime Mold Time Mold. Their argument, across a prologue and twelve parts, is that psychology never had a paradigm because it never proposed actual entities and rules, only abstractions. Their candidate entity is the negative feedback loop, which they call a **governor**: a sensor, a set point, a comparator, and an output that acts on the world. Hunger, thirst, fear and shame are all error signals from different governors, all competing for one body through an arbitration layer they call the selector. From there they rebuild personality as the parameter settings on your governors, and recast depression not as one disease but as several mechanically distinct faults that happen to share a surface. It's self-published, it has essentially no new data, and it's the most complete recent statement of a control-theoretic psychology anyone has written.

Then what is happiness?

It's the first question anybody asks, and it's a good one. Hunger, fear and shame all behave like error signals. Joy doesn't. Nothing in you is working to drive joy to zero.

The Slime Mold answer is that happiness isn't an emotion at all.

"Emotions are easy to identify because they are errors in a control system. Like any error in a control system, successful behavior drives the error to zero. This means that happiness is not an emotion."

What happiness is, on their account, is what it feels like when a governor's error gets corrected. Their line: "Happiness is what happens when a thirsty person drinks, when a tired person rests, when a frightened person reaches safety." Correct a large error, or correct one quickly, and you get more of it than from a slow incremental fix. That matches ordinary experience better than it has any right to. The best meal of your life was eaten hungry.

It also predicts something odd that I think is right: there's no such thing as unhappiness. Happiness can't go negative. What gets called unhappiness is an uncorrected error somewhere, which is a different thing wearing the same word.

And they give it a job. Happiness marks which behaviors worked so the system repeats them, and helps calibrate how much to explore versus stay with what already works.

Now the hole, because there is one.

That account handles joy, relief, satisfaction and pleasure by making them all the same thing. Their words: "there's only one way to feel good," and "all of our words for positive emotion, joy, excitement, pride, are really referring to the same thing, just in different contexts."

That's asserted rather than argued, and it leaves amusement out entirely. What error gets corrected when something is funny? What set point does a joke defend? Nothing obvious, and the series never raises it. Amusement, awe and aesthetic pleasure don't turn up anywhere in fourteen posts.

They're straighter about curiosity, which they flag themselves as an enigma: "acting on your fear should make you less afraid, acting on your thirst should make you less thirsty, but acting on your curiosity often seems to make you more curious." A signal that grows when you act on it runs backwards from every other loop in the model.

So my position is narrower than theirs. Error signals account for the emotions that push you toward a target. Happiness is a reasonable read of what arriving feels like. Amusement, awe and curiosity are unaccounted for, and I'd rather leave them that way than stretch the model to cover them, which is the exact flexibility problem Section 9 is about.

If the list of regulated variables reminds you of Maslow, the resemblance is real but the logic runs the other way. Maslow arranged needs into a hierarchy of priority. This arranges them as parallel control loops competing for one body, which is why a person can be hungry and lonely and cold at the same time without any of it queuing politely.

THE BASIC CONTROL LOOP

a set point, a sensor, an error, and something that acts

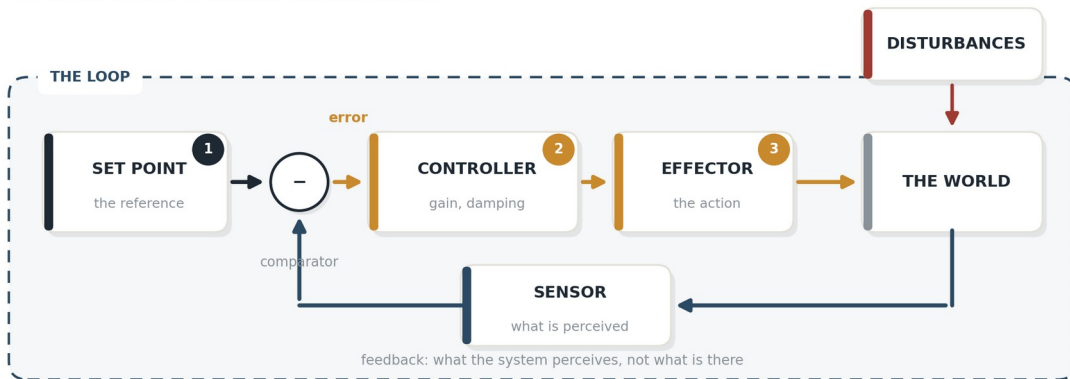


Figure 2.1. The basic control loop.

What this buys you is a failure taxonomy. A loop can break in a small number of specific ways: the sensor misreports, the set point is wrong, the gain is mis-set, the damping is insufficient, the effector fails, or in a system with competing loops the arbitration fails. Those are different problems with different interventions, and symptom-level description doesn't distinguish them.

The one that matters most here is gain and damping. A loop with high gain and insufficient damping doesn't just respond too strongly. It **oscillates**. It overshoots, corrects, overshoots the other way, and either settles slowly or never settles at all. Wiener identified cerebellar tremor as exactly this problem in 1948. Somebody with cerebellar damage reaches for a glass and the hand doesn't travel smoothly to it. It overshoots, comes back too far, overshoots again, and hunts around the target, getting worse the closer it gets. Wiener's point was that this isn't weakness or paralysis. The muscles work. What's broken is the correction: it arrives too hard and too late, which is the textbook signature of a feedback loop with its damping removed.

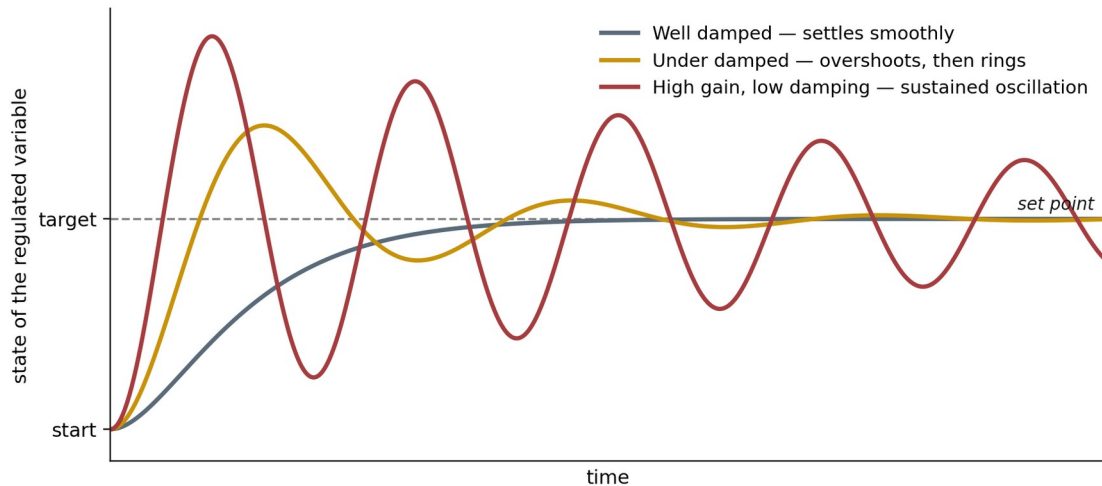


Figure 2.2. Gain and damping. Same system, three parameter settings.

Where it falls short. A thermostat reads a number and takes it at face value. No biological system gets that deal. Its sensors are noisy, slow, and measure proxies rather than the regulated variable. Plasma osmolality is not thirst. A racing heart is not danger. The control account has a sensor-shaped hole in it.

2a. The human thermostat

Before going further, work the thermostat all the way through. It's the cleanest way to see why a controller alone isn't enough, and it's the example I use when I'm explaining this out loud.

A thermostat on your wall already has predictions built into it. Not obvious ones. It waits a little while before it adjusts. This is called **hysteresis**, which means the point where it switches on isn't the same as the point where it switches off. Once it turns the air conditioning on it holds until the temperature passes the set point, rather than flipping the instant the number crosses. It waits a few degrees on the way back. None of that is comfort engineering. It's there so the unit isn't cycling on and off and on and off on sensor noise.

Now take the thermostat off the wall and put a person in its place.

Hand somebody the dial. Their only sensor is their own body. They feel warm, so they crank the AC up. Then they feel cold, so they crank it down. They do it the instant they feel it.

Here's what they don't know at the beginning. There's a delay between turning the dial and their body registering the change (it takes a bit to turn the AC on, then for the air to circulate and get the temperature reduced and then for their body to sense it).

So they crank it up when they're hot, and nothing happens, so they crank it further. Then the room goes cold, and they crank it way down. Then it's hot again. They keep cranking,

harder each time, because they aren't getting the response they expect on the timeline they expect it.

That person doesn't have a broken thermostat. They have a **broken model of the system**. They haven't learned that the environment has a lag in it.

Now let them learn. Once they know about the delay, the behavior changes completely: crank it up a little, wait, then decide whether to adjust again. Same dial, same body, same room. What changed is the prediction, and the prediction is what damped the controller.

That splits the job in two, and the split is the whole point of this document.

The estimator's job is to predict what's going to happen in the environment given the sensory input coming in. **The controller's job** is to decide how much to adjust.

Two different failures live in those two jobs, and they look similar from outside.

Go back to the dial. Suppose the sensor itself is bad. The thermostat reads 70, then 71, then 68, then 71, then 68, and it reacts to every reading. Now the AC cycles constantly, and nothing is wrong with the decision rule at all. The reading is the problem.

Or suppose the sensor is perfect and the built-in patience is missing. It's exactly 70.0. At 70.1 it turns on, at 69.9 it turns off. Perfect information, no prediction, and the same cycling.

A worked contrast, because it isn't all failure. Put your hand on a knife and cut yourself. The pain arrives large and the response is immediate, which is the controller doing its job correctly, because that's a case where fast and hard is the right answer. Meanwhile the other half of the system is doing something slower and more useful. It's taking that sensory input against what you expected and learning from it: *that hurt because the knife is sharp*. Now you carry a prior that the knife is sharp, and your caution around knives is set at a more accurate level than it was an hour ago. One event, two jobs. The controller handled the moment. The estimator changed the setting.

Which is also why the same parameter is right in one loop and wrong in another. For pain and danger, act fast. For whether the room is a little warm, you can afford to be slow. React to discomfort on the timescale you'd react to danger and you're the person cranking the dial.

3. The mind as a prediction machine

Predictive processing says the brain doesn't receive the world so much as guess at it, continuously, correcting the guess only where the evidence insists.

Predictions, with estimations of the expected environment, run down the hierarchy, and only the mismatch (the prediction error) runs back up. What you end up perceiving is a settlement between what you expected and what arrived, weighted by how much the

system trusts each one. That weighting is called **precision**, and it's the parameter that does all the work.

Aside... if you have read Plato's The Allegory of the Cave, this is kind of congruent with that... but enough intellectualizing here....



Figure 3.1. The allegory of the cave. The prisoners see shadows, not the objects casting them, and take the shadows for the world.

DOWN what the brain expects

UP only what it did not expect

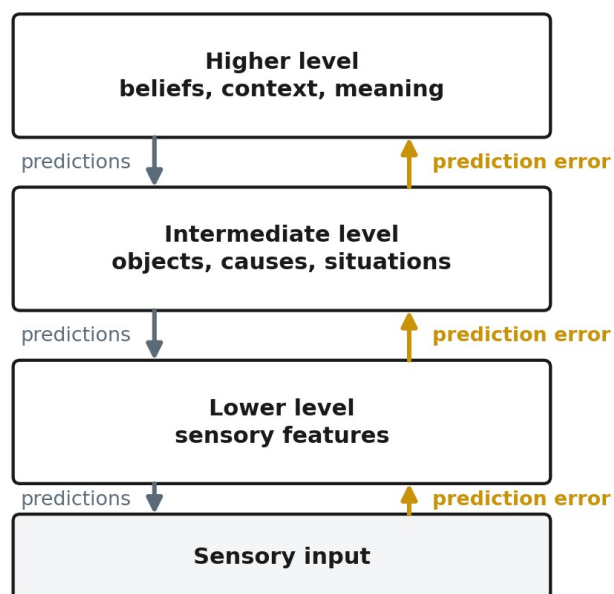


Figure 3.2. Predictions descend. Only prediction error ascends.

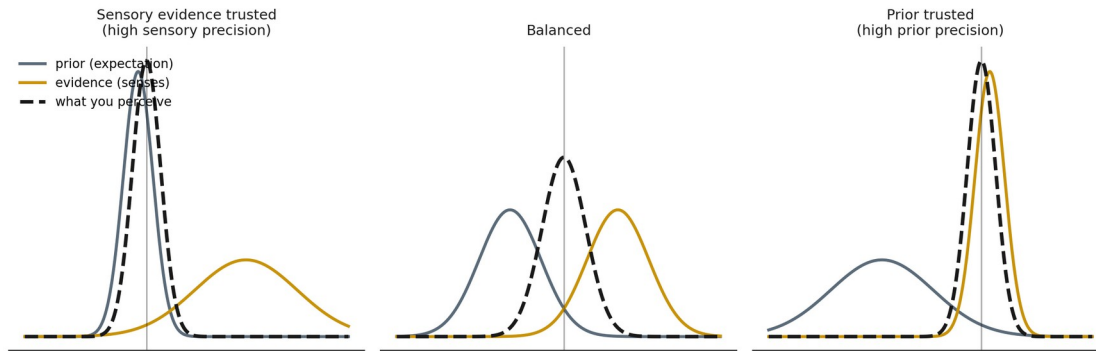


Figure 3.3. Precision weighting. Identical inputs, three different settings.

Where it falls short. It explains beautifully how an estimate gets built and says very little about why the organism cares, or how hard it acts once it knows. And its evidence base is thinner than its citation counts suggest. Section 8 gives the numbers.

4. Putting them together

Here's the move, and it's the only genuinely new thing in this document.

The estimate establishes the target. The estimator combines what you expected with what you sensed and produces an estimate of the current state. That estimate is also your updated prediction of what should be happening. And that prediction is what the controller compares against. There's no separate goal sitting off to the side. The thing you predict IS the set point you defend.

THE COMBINED MODEL

the estimate becomes the prediction, and the prediction is the set point

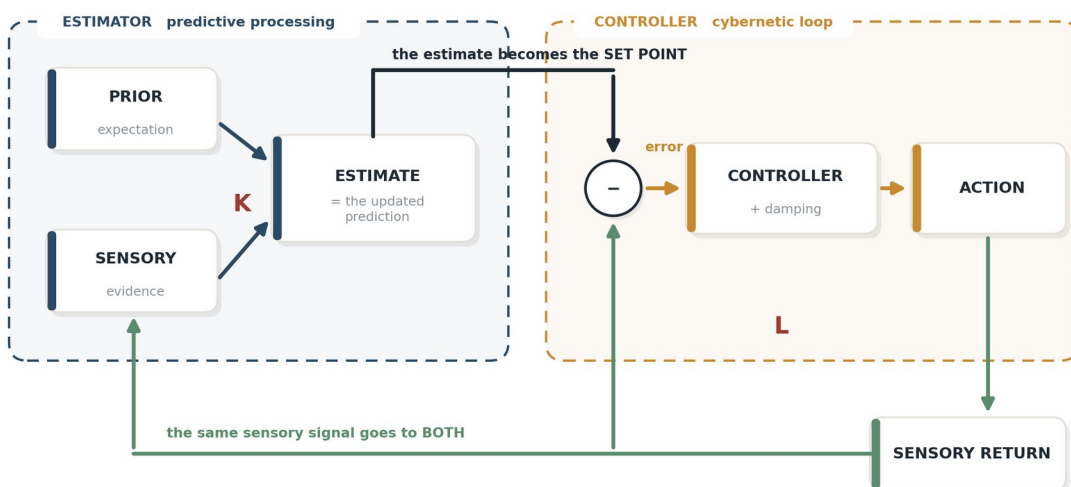


Figure 4.1. The combined model. The estimate of state becomes the updated prediction, which serves as the set point. The sensory return from action feeds both the controller and the estimator.

That gives two gains, and they are formally distinct.

Observer gain (K) is how much new sensory evidence is allowed to move your estimate, relative to what you already expected. Predictive processing calls this precision.

Controller gain (L) is how hard the system acts once an error exists. It has nothing to do with what you believe.

The separation isn't my invention. It's the separation principle in control theory, worked out by Kalman in 1960 and Joseph and Tou in 1961: for a linear system with Gaussian noise, the optimal estimator and the optimal controller can be designed independently. What's new here is the claim that **because the estimate becomes the set point, K determines where the target sits**, not just how confident you are about it. That's what makes the two dials clinically different rather than just mathematically separable.

Emotion, precisely

If an emotion is the error signal, and the estimate feeding that comparison is an *inference* rather than a measurement, then two things follow. Both of them happen constantly, and anyone who does this work has watched them happen.

The reason emotions respond to reframing at all is because inferences move when the evidence or the expectation moves. A thermostat doesn't feel differently about 62 degrees once you explain the situation to it.

And reframing has a ceiling.

Changing a high-level belief doesn't automatically change a low-level estimate. "I know it's irrational and I still feel it" isn't a failure of insight.

That sentence is what it feels like when one level of the system has updated and a lower level has not. The part of you that talks has changed its mind. The part of you that runs the estimate is still using the old weighting, and it's the one generating the feeling.

The learning theory

None of these parameters are fixed at birth...there may be an initial set point but then each is estimated from experience. That is, if the system is working properly.

Volatility teaches observer gain. In a stable world a surprise is noise and should be discounted. In an unpredictable one, it's a signal and should move you a long way. People demonstrably track this. Someone who grew up in a chaotic house learned to weight incoming evidence heavily, because in that house it WAS heavily weighted evidence.

Consequence teaches controller gain. Where small problems reliably became large ones, high gain isn't a malfunction. It's the correct setting for that environment. It was learned correctly.

Overshoot teaches damping. A system punished for waiting learned not to wait.

Volatility and consequence set those parameters together, not separately, and laying them on two axes is more useful than either one alone.

	Low consequence	High consequence
Stable environment	Dampened reaction, little prediction adjustment	Reliable prediction, strong reaction
Unstable environment	Fast prediction updating, reaction either way	Rapid prediction updating, strong reaction

An aside... attachment theory fans may consider this model in development of stable, anxious and avoidant styles.

The two problem quadrants aren't the ones people expect. Stable and high-consequence produces someone reliable and intense, which looks like competence until the environment changes. Unstable and high-consequence produces someone who updates fast and reacts hard, which is exhausting to run and correct for the world that built it.

Incongruent, not wrong. The setting was right where it was learned. The environment changed and the setting didn't. Maybe when we were young we learned the best possible reaction to the environment... then it persists even though the environment has changed.

It shows up two ways. An ineffective K means the estimate itself is off, so the target sits in the wrong place and everything downstream is aimed at it. An ineffective L means the estimate is fine and the correction is disproportionate, or adjusts in the wrong direction entirely. Same presentation, different repair.

Which gives the claim:

Most of what looks like dysregulation is a correctly learned parameter running in an environment that no longer matches it.

Aaron Beck, the psychiatrist who built cognitive therapy in the 1960s and gave us the idea that what you believe about a situation drives how you feel about it, would call the learned part a **schema**: a stored belief about yourself or the world, formed early and applied automatically ever after (Beck, 1967, 1979). He'd be right. What the model adds is where a schema comes from mechanically: it's the stored prior, built from the statistics of an environment the person actually lived in, still being emitted as a prediction in an environment that no longer matches it.

Expectation runs in front of evidence. We decide and then manufacture the logic. A Vulcan couldn't run a business because rationalization almost always

arrives before reasoning does, and the antidote is engineering discipline: estimate, test, and let the results outrank the argument you fell in love with.

Writers have a version of this, and it arrives with a lesson attached. It's usually quoted as Faulkner: *in writing, you must kill all your darlings*. Faulkner never said it. It's Arthur Quiller-Couch, lecturing at Cambridge on January 28, 1914: "Whenever you feel an impulse to perpetrate a piece of exceptionally fine writing, obey it, wholeheartedly, and delete it before sending your manuscript to press. *Murder your darlings*" (Quiller-Couch, 1916). The Faulkner version got popular through screenwriting guides in the 1990s and has been loose ever since. Samuel Johnson gave the same advice 150 years earlier and credited it to an unnamed old tutor of a college. So the advice about killing your darlings arrives with a darling of its own that nobody has managed to kill.

It carries a warning. A parameter learned over fifteen years from consistent evidence won't be revised by one good afternoon of argument. Which is why insight is a poor lever and repeated disconfirming experience is a good one.

5. Fitting this into CBT

Before going anywhere near application, this has to line up with cognitive behavioral therapy, because CBT is the model I like, it's the model many people reading this relate to, and a framework that can't be stated in its terms isn't worth much.

It lines up pretty well. Term for term, with one piece left over that turns out to matter.

The model I teach

Here's the diagram I've learned from my mentor Ryan Sheade, LCSW. I have put in front of students and clients for years.

COGNITIVE BEHAVIORAL THEORY

the chain as it is taught

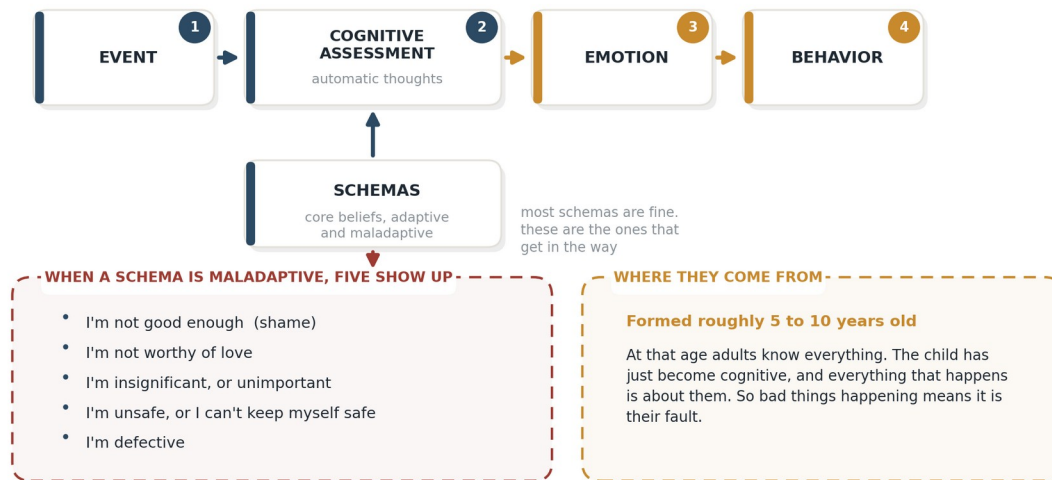


Figure 5.1. The cognitive behavioral model as I represent it.

An event happens. You run a cognitive assessment of it, which shows up as automatic thoughts. That assessment produces an emotion. The emotion produces behavior.

Underneath the assessment sit the schemas, the core beliefs, and those are what the assessment is running on. When they're maladaptive, five show up over and over:

I'm not good enough. That one arrives as shame. I'm not worthy of love. I'm insignificant, or unimportant. I'm unsafe, or I can't keep myself safe. I'm defective.

They form early, roughly five to ten years old. From the child's point of view, at that age adults know everything. The child has just become cognitive enough to build general rules out of specific events, and not yet cognitive enough to consider that a rule might be about somebody else. So everything that happens is about them. Bad things happening means something is wrong with them, and the rule gets written down and filed. To be clear, not all core beliefs are maladaptive, it's the maladaptive ones that often get in our way.

The same diagram, drawn as a loop

Now redraw it.

THE SAME MODEL, DRAWN AS A LOOP

the sensory return is taken to two places

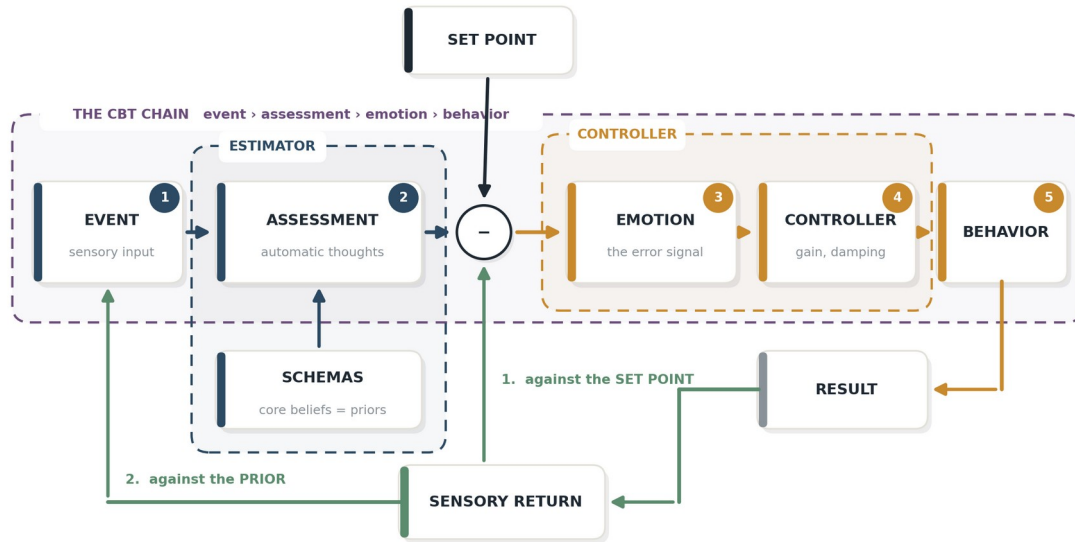


Figure 5.2. The same model as a control loop.

Every term carries over with nothing left standing.

The event is sensory input. Something arrives.

The cognitive assessment is the estimator. It takes the input, weighs it against a prior, and produces an estimate of the situation. Not a description of the situation. An estimate. **The schemas are the priors.** That isn't a metaphor and it isn't a rename, it's the same job: a stored expectation that shapes what ambiguous evidence gets read as.

The emotion is the error signal. It's the gap between the estimate and the set point the system is defending. (The positive emotions, and why happiness may not be an emotion under this account, are in Section 2.)

The behavior is the controller's output. What gets done about the gap, and how hard.

Then the piece the standard diagram doesn't have. The behavior produces a **result**, and the result is the next event. The chain isn't a chain. It's a loop, and it runs continuously.

Two things follow immediately, and the second one is why the redraw was worth doing.

What the loop makes visible

The step from assessment to emotion stops being a mystery. In the standard teaching that arrow is the part that gets waved at, and students accept it because the sequence obviously works even when the mechanism is vague. Under the control reading the emotion is a comparison, and its magnitude is the size of a discrepancy. Which is why the same thought produces a different feeling in two people, and a different feeling in one person on two different days. The thought didn't change. What it was being measured against did.

That also gives the four of the five core beliefs a job, and separates out the fifth. *Not good enough, not worthy of love, insignificant, defective* are all priors. They bias what the estimator concludes from thin evidence. But *I'm unsafe, or I can't keep myself safe* isn't doing that. It's setting how much estimated danger is tolerable before the system calls it an error. It's a set point, not a prior, and it produces a different presentation and needs a different intervention. Under the standard model all five sit in the same box.

And the loop explains why a wrong belief survives contact with reality. This is the part the left-to-right chain can't show, because a chain ends at behavior and reality never gets to answer.

Work an example. The prior is *I'm not worthy of love*.

The event is a partner who's quiet at dinner. The estimator, running on that prior, reads quiet as withdrawal, and withdrawal as the front edge of leaving. The set point is being securely connected. The gap is large, so the emotion is large, somewhere between fear and shame. The controller has high gain and no damping, so the reaction is to press. What's wrong. Are we okay. You'd tell me, right. Three hours of it.

The result is a partner who was tired at the start of the evening and is now irritated and genuinely pulling back.

That result is the next event.

Nobody in that loop did anything irrational. Every step follows correctly from the step before it. The prior was wrong when the evening started and accurate by the time it ended, because the loop went and made it accurate.

That's what a maladaptive core belief is, mechanically. Not a wrong idea sitting in storage waiting to be corrected. A prior running a loop that manufactures its own evidence. Which is also why arguing with the belief so often does nothing: you're disputing a conclusion the system re-derives from fresh data every week.

DON'T say: "You're being irrational." SAY: "I think your instrument is reading off."

The anxious person usually isn't over-reacting to their situation. They're reacting accurately to an estimate of their situation that happens to be wrong.

Does the assessment do both jobs?

Here's the question I keep coming back to, and I don't have an answer to it.

The clean version of Figure 5.2 gives the cognitive assessment one job, which is to produce an estimate. But look at what an assessment actually delivers in a session. It tells you what's happening and how much it matters, in the same breath. *He's pulling away* is an estimate. *He's pulling away and this is the end of everything* is an estimate with a gain setting attached to it.

Appraisal research has said something close to this for decades. The appraisal is held to determine the intensity of the emotion and not only its content (Smith & Ellsworth, 1985). If that's right, the assessment box is doing two mechanically different jobs, and the standard model is treating them as one.

Which opens something the diagram doesn't show. A maladaptive core belief may be able to enter the loop in three different places.

As a prior. *I'm defective* biases what the estimator concludes from thin evidence. Ambiguous feedback gets read as criticism.

As a set point. *I'm unsafe, or I can't keep myself safe* changes what counts as an error at all. Less estimated danger is tolerable, so the same situation opens a gap that wouldn't open in somebody else.

As a gain. *I'm not worthy of love* may not distort the reading whatsoever. It may set how hard the system acts once an error exists, because it reads the stakes as total. The estimate is accurate. The response is maximum.

Same five beliefs. Three mechanisms. Three different treatments, which is the argument of this entire document arriving from a direction I hadn't planned on.

I can't settle it from here and as far as I can tell nobody has. What's worth saying is that it's separable in principle. A belief operating as a prior predicts biased interpretation of *ambiguous* evidence and roughly normal handling of unambiguous evidence. A belief operating as a gain predicts accurate interpretation across the board and an outsized response to both. Hand the same person ambiguous and unambiguous versions of the same situation, and the two accounts come apart.

That's a study, and it's a small one.

The vocabulary, side by side

CBT term	Control term	What the translation adds
Activating event	Sensory input	Nothing. Same thing.
Core belief, schema	Prior	Little on its own. Vocabulary substitution unless precision comes with it.
Automatic thought	The prediction the estimator emits	Separates the sentence from the estimate driving the feeling
Conviction rating, 0 to 10	Where the estimate currently sits	Says nothing about whether it can move. That's a second parameter.
Emotion	Error signal = Emotion	Gives the assessment-to-emotion arrow an actual mechanism

Coping behavior	Corrective action	Ties intensity to gain rather than to the belief
Safety behavior	A correction that removes the signal without testing the estimate	Explains why comfort can be worse here than distress
Behavioral experiment	Prediction-error protocol	Explains precisely how one fails
Exposure	Forced error generation	Moves the active ingredient off habituation
Response prevention	Damping	Explains why the first week gets worse

Most of that right-hand column is a rename. Saying so up front is cheaper than being caught at it later. Three of the rows aren't.

Three places the translation pays

Conviction is not movability. In the cognitive model a belief has content and strength, and strength gets, for example, a number out of a hundred. In the predictive account, the confidence a prediction is held with is a separate parameter from what the prediction says, and it can move without the content moving at all. Calling a schema a prior and stopping there is vocabulary substitution. Starbucks calling a small a tall doesn't put more coffee in the cup.

Two clients both say people think they're stupid. Both rate it 7/10. In the first, the content is wrong and contrary evidence gets in, so restructuring moves the estimate. In the second, the content may be partly accurate somewhere, and the prediction is held so tightly that nothing moves it. Restructure that one and you change the words while the parameter sits untouched. The client agrees with you and feels exactly the same, and both of you leave the session confused about why.

Aside... a similar concept: "Affirmations only work if you believe them" – Jim Rohn

I read conviction ratings for years as though they told me how hard a belief would be to shift. They don't. They tell you where the estimate is sitting. Whether it can be moved is a different question, and it's the one you're actually being paid to answer.

Expectancy violation and prediction error are the same quantity. Craske, Treanor, Conway, Zbozinek and Vervliet (2014) argue that exposure works by building a new inhibitory association rather than erasing the old fear, and that the target is the mismatch between what the client expected and what happened, not the reduction of fear inside the session. Read that in control terms and it isn't an analogy, it's the same quantity under a different name. Craske says the therapeutic work is done by the size of the gap between what somebody expected and what actually happened. This model says the estimate updates in proportion to the size of the prediction error. Those are two descriptions of one

number. She got there from exposure research and I got here from control engineering, and we are both pointing at the gap.

The control account doesn't improve on Craske. It agrees with her, which is the most useful thing it does in this entire section, because her version already has evidence behind it.

DON'T ask: "How anxious are you now compared with when we started?" SAY: "What did you predict would happen? What actually happened?"

The same reading sharpens why a behavioral experiment fails, and there are only three ways. The prediction was never stated precisely enough to be violated, because "it will go badly" can't generate a clean error. The prior was held at such high precision that the discrepancy got absorbed as a fluke. Or a covert safety behavior meant the experiment tested a different prediction from the one written on the form. All three are prevented by the same step, which is writing the prediction down before anything happens. Skip it and you get the Home Depot outcome. One trip for the part you thought you needed, two more for the parts the job actually required.

"I know it's irrational but I still feel it." Stott (2007) called this rational-emotional dissociation, and clients produce it in close to those exact words. The predictive reading is specific. Verbal work can move the high-level prior, the one that generates sentences, and leave the precision on lower-level predictions exactly where it was. The explicit model updated. The model actually running the interoceptive estimate did not. Both are running. Only one of them talks.

Which fits something worth stating plainly: people decide emotionally and then assemble the logical case afterward. A Vulcan couldn't run a real business. Every clinician reading this has watched somebody build an airtight argument for a conclusion they had reached before they sat down.

~~I am afraid to be wrong in front of these people.~~ "I just want to make sure we're all aligned before I commit."

The argument isn't the cause. It's the transcript.

Four intervention families instead of two

Separating the estimator from the controller splits the standard toolkit into four targets, and the reason to bother is that two of them present identically.

Target	What it changes	Interventions
Prior, content	The expectation being emitted	Restructuring, thought records, psychoeducation, formulation, meaning-making
Observer gain, precision	How tightly the prediction is held, independent of content	Behavioral experiments, exposure designed for expectancy violation,

Controller gain	Intensity of the response, without touching the belief	attention training, tolerating uncertainty Arousal regulation, paced breathing, applied relaxation, distress tolerance skills
Damping	Whether corrections converge or oscillate	Delay, urge surfing, response prevention, scheduled worry time, ritual postponement

The rows aren't clean and I'm not going to pretend they are. A behavioral experiment moves precision and usually content with it. Mindfulness reallocates attention and lowers arousal at the same time. The question the table answers is what an intervention is mainly for.

Rows two and three are where it matters. A person with an accurate belief and a controller that overshoots, and a person with an inaccurate belief held at high precision, both arrive distressed, both describe an intense reaction, and both may rate conviction at 80. Restructuring is close to useless for the first and central for the second. Arousal work is central for the first and a distraction for the second. Standard assessment doesn't separate them, because standard assessment asks what somebody believes and how strongly, and both of those questions land on the same row.

DON'T stop at: "What do you believe, and how convinced are you?" ALSO ASK: "What would it take to change your mind?" (precision). "What did you do about it?" (controller gain). "How many times did you do it?" (damping).

All three are answerable inside an ordinary intake, and they point at different treatments.

One warning falls straight out of the table. Breathing retraining sits in the controller row, and in panic it can quietly become a safety behavior, a correction that removes the error signal and blocks the test. Salkovskis, Clark and Hackmann (1991) treated panic successfully with cognitive therapy and no breathing retraining at all, which is worth remembering. Lowering controller gain helps when the estimate is accurate. It hurts when it becomes the thing preventing the estimate from being corrected.

What this doesn't do

It doesn't add a technique. Every intervention in that table already existed, most with better evidence behind it than anything in this document. What changes is the selection rule. This is an effort to improve the model (or maybe just add more detail).

And the dismantling data should keep anyone honest about how much that's worth. Ougrin (2011) pooled 20 randomized trials, 1,308 participants, comparing cognitive therapy against exposure, and found no significant difference in panic, PTSD or OCD, with cognitive therapy ahead only in social phobia. Öst, Havnen, Hansen and Kvale (2015) found the same

inside OCD specifically: exposure and response prevention versus cognitive therapy came out at an effect size of 0.07, which is nothing.

If two routes reach the same outcome, a model explaining the mechanism of one route has not thereby explained the disorder's treatment. Whatever is doing the work is reachable from both directions, and the model I've built predicts a separation that decades of outcome data have not produced.

The clean answer is that the trials never sorted people by which parameter was broken, so the two groups washed out against each other. That answer is also exactly the kind of unfalsifiable patch Section 9 says should make you suspicious. It's testable, which is the only thing that redeems it, and until somebody runs it, CBT works better than this model can currently explain.

6. What it explains in the clinic

The chain the model proposes:

learning history → prior → prediction → the estimate the system regulates against → error signal → corrective action → the action prevents disconfirmation → the prior is reinforced

THE REINFORCING LOOP

the action prevents the very evidence that would correct the belief



Figure 6.1. The reinforcing loop.

Most of that is standard cognitive therapy in different words, and saying so is cheaper than being caught at it. Beck's schemas are priors. Safety behaviors preventing disconfirmation is Salkovskis, tested experimentally, and control theory contributed nothing to it. The overlap isn't a weakness in the argument. It's the reason to trust the argument, because a mechanical account that contradicted decades of outcome data would be the account with the problem.

What the layer adds is a separation the cognitive model runs together.

DON'T say: "You're being irrational." SAY: "I think your instrument is reading off."

The anxious person usually isn't over-reacting to their situation. They're reacting accurately to an estimate of their situation that happens to be wrong.

The four questions

Is the belief inaccurate, or accurate and acted on too hard? Restructuring addresses the second and does little for the first. The first needs work on response intensity: arousal regulation, distress tolerance, deliberate delay.

The person who says their boss is unhappy with them, and the boss is in fact unhappy with them. Nothing to restructure. The estimate is correct. What's disproportionate is the response, which is three days of rumination and a resignation letter drafted at midnight. Restructuring here tells an accurate person they're wrong, and they will correctly conclude you're not listening.

Is the belief wrong, or held with too much confidence? Content and confidence come apart. A moderately accurate belief held with total certainty behaves worse than a somewhat inaccurate one held loosely, because certainty blocks updating. Behavioral experiments work on confidence. That's why they beat argument.

Is the person failing to update, or being prevented from updating? Safety behavior stops the error signal from ever being generated. The prior isn't stubborn. It's starved.

A phobia is a great example case. Say I'm afraid of heights. The prediction is that height means danger, the error is large (emotion=fear!), and the corrective action is to get away from the edge. The action works, in the narrow sense that the fear drops. It also guarantees I never collect the evidence that would revise the estimate, which is standing near an edge and not falling (it's not unsafe). Every successful escape confirms the model rather than providing new data to revise the estimate. Thirty years of that and the prior is untouched, not because it's stubborn but because it has never once been given new data that revises the estimate.

Flooding and systematic desensitization work by forcing the data through. Stay near the edge long enough, without the escape, and the sensory return finally contradicts the prediction. The estimate moves, the prediction resets, and the error the controller sees gets smaller on its own. You didn't argue anyone out of anything. You restored the evidence stream the avoidance was blocking.

What environment taught this setting, and does it still exist? Usually answerable, usually informative, and it changes the conversation from correction to recognition.

The person who reads a two-word text as a catastrophe. Ask where a two-word message used to mean something was badly wrong. Frequently there's an immediate answer. The setting was correct in that scenario. The current sender (new environment) is just terse.

Rumination, worry and compulsion as oscillation

Rumination, worry, and compulsions all look like oscillation. A loop that can't settle, cycling at a characteristic frequency, each pass generating a correction that doesn't resolve the error.

Under that reading, "not acting" is therapeutic in a specific mechanical sense. You're adding damping.

Take the person who checks the lock. The urge arrives, the correction is immediate, the urge returns in ten minutes, and the cycle runs all evening. Standard response prevention asks them to feel the urge and not check. In control terms, you've inserted a delay between the error and the correction, which is the definition of damping. The first few cycles get worse, because an under-damped loop always overshoots before it settles. Then the oscillation loses amplitude and the intervals stretch out.

The same logic covers the worry spiral where you agree to postpone the worrying to a set hour, and the argument where the rule is to wait a day before replying. In every case the intervention isn't insight and it isn't a better answer. It's slowing the loop down enough to let it settle.

This also explains why telling someone to stop worrying doesn't work, it just raises the gain, because now failing to stop is a second error to correct. Adding damping works. Adding pressure adds gain.

Depression, read as several faults rather than one

The control reading makes a prediction that the diagnostic category doesn't: depression should turn out to be mechanically distinct faults that happen to share a surface.

A set point set too high produces someone who never registers success, because nothing clears the bar. A controller with the gain turned down produces someone who registers the error and cannot mount the correction, which is what anhedonia and psychomotor slowing look like from outside. An estimator that has stopped updating produces someone whose predictions no longer move when the world improves. Those are three different faults. They present similarly and they should respond to different things.

Status: this is the model's clearest testable claim in the clinic, and nobody has run it.

7. What it explains outside the clinic

This part surprised me. The same four parameters describe a person deciding under pressure, a manager running a team, and a whole company steering itself. Not by analogy either. Same mechanics, because they're all regulators with sensors, targets and lag.

The lag problem, which is most of it

Go back to the person with the thermostat dial. They aren't stupid, and they aren't over-emotional. They're acting on a system whose response arrives later than their expectation of it, and they don't know that yet.

That is the single most common failure I see in businesses.

A marketing change goes in on Monday. Results land in four to six weeks. The owner checks on Wednesday, sees nothing, and changes it again. Then again the following week. By the time the first change would have shown up, three more have been layered on top and no result can be attributed to anything. The campaign didn't fail. The loop was run faster than the system's response time.

The fix isn't patience as a virtue, it's damping as an engineering decision. **Find out what the system's response time actually is, then set your review interval longer than that.** A change reviewed on a shorter cycle than its own feedback delay can't be evaluated. It can only be replaced.

The two dials at work

The separation between estimate and response is at least as useful in an office as in a session, because the two failures look identical from the outside and get the same unhelpful advice.

An accurate read with too much gain. A leader who correctly notices a real problem and responds at a magnitude the problem doesn't warrant. The estimate is right. Telling them they're overreacting to nothing is both wrong and expensive, because they aren't reacting to nothing. What needs to come down is the amplitude, not the perception. Reducing controller gain doesn't mean the response stops. Keith Richards is 82 with arthritis and books a residency instead of a world tour. The output continues. The amplitude comes down.

An inaccurate read held with high confidence. Someone certain about a market, a competitor or a person, on thin evidence, who doesn't update when contradicted. Arousal work does nothing here. What moves it is a test with a written-down prediction.

Same complaint from a colleague — "he overreacts" — two different repairs. Ask which one you're looking at before you intervene.

Micromanagement is a gain and damping problem

A manager who checks work every two hours is running a control loop at a frequency far above the rate at which the work actually produces signal. Every check generates an error, every error generates a correction, and the corrections arrive faster than the work can absorb them. The employee's output starts oscillating, which the manager reads as unreliability, which raises the checking rate.

The loop is the problem, not either person. Lengthen the interval and the oscillation damps out on its own. That's hard to do because it means acting less at the exact moment the system looks worst... which is the same thing that makes response prevention hard for a client.

KPIs are set points, and most of them are wrong

An organization steering by numbers is a control system with all the same failure modes.

The set point is arbitrary. A number picked to be motivating rather than to describe a defended state. The system will defend it anyway. **The sensor measures a proxy.** Plasma osmolality is not thirst, and monthly revenue is not business health. Both get defended as though they were the thing. **The delay is unacknowledged.** Reporting lag means every correction is aimed at where the business was, not where it is. **The gain is too high.** One bad month triggers a response sized for a bad year. **There's no damping.** Strategy changes every quarter, and no strategy runs long enough to produce evidence about itself.

A company with those five settings looks chaotic, and it usually gets described in terms of culture and personality. It's a loop with the parameters set wrong. Parameters you can adjust. Culture you mostly can't.

The reinforcing loop, in a business

The clinical loop has an exact commercial twin, and it's worth stating because it's expensive.

A belief that a market segment won't buy produces a prediction, which produces reduced effort with that segment, which produces poor results, which confirms the belief. The action prevented the evidence. Nobody was irrational at any point, and the belief is now backed by data the belief itself generated.

The test is the same as it is in a session: state the prediction in advance, run the experiment without the safety behavior, and check the result against what you wrote down rather than against how it felt.

Decisions under pressure

Three things fall out that are worth carrying into a hard decision.

Separate what you think is happening from how hard you're going to act on it. Say them as two sentences. Most bad decisions under pressure are an accurate read with a disproportionate response, and they get defended by pointing at the accuracy of the read.

Ask what your review interval is before you start. If you don't set it in advance you'll set it under stress, and under stress it will be too short.

Ask what would change your mind, and write it down before you have the answer. That converts a belief into a test. It's the only reliable way to tell whether you're updating on evidence or defending a prior.

None of that is new advice. What the model adds is why it works, and that matters at the moment you least want to follow it.

8. What the evidence actually supports

Anyone using this framework should know how much of it is established and how much is proposal. The two are not close to equal.

What is well supported. People track environmental volatility and adjust how fast they learn from it. Learning rates rise when the world becomes unpredictable and fall when it stabilizes (Behrens, Woolrich, Walton & Rushworth, 2007). Anxious people fail to make that adjustment, learning at a similar rate whether the world is changing or not (Browning, Behrens, Jocham, O'Reilly & Bishop, 2015). That's the empirical backbone of the observer-gain half, and it's real work with real data.

What is well supported and already known. Exposure works by violating expectations rather than by wearing fear down (Craske, Treanor, Conway, Zbozinek & Vervliet, 2014). Safety behaviors maintain anxiety by preventing disconfirmation (Salkovskis, Clark, Hackmann, Wells & Gelder, 1999). Behavioral experiments work by testing a prediction written down in advance (Bennett-Levy et al., 2004). None of that needs this model. The model agrees with it, which is the honest reason to take the model seriously and not a reason to credit it with anything.

What has a long literature that psychology mostly ignored. Control-theoretic accounts of behavior are not new. Powers set out Perceptual Control Theory in 1973. Carver and Scheier brought control thinking into personality and health psychology in 1982. Method of Levels is a therapy built explicitly on it. The evidence base there is modest but it exists, and any claim of novelty has to be made against it rather than around it.

What is reinterpretation rather than result. The mapping of CBT onto control terms, the reading of rumination as oscillation, the reading of response prevention as damping. Each is coherent. None is a finding. Calling a schema a prior explains nothing extra until precision comes with it.

What is untested. The central claim. Nobody has shown that observer gain and controller gain dissociate in real people, that they predict differential treatment response, or that the two-gain framing outperforms a conviction rating. That study is designable and cheap, and it has not been run.

A cautionary note from the harder case. This model was first worked out on autism and repetitive behavior, and the evidence audit there was sobering. The largest synthesis of Bayesian accounts of autism, 83 studies, reports that a slight majority find no group difference at all (Angeletos Chrysaitis & Seriès, 2023), while the second largest reaches the opposite conclusion (Cannon, O'Brien, Bungert & Sinha, 2021). The flagship computational finding failed a preregistered test using the identical model (Plank et al., 2026). The mismatch negativity, described everywhere as the neural signature of aberrant prediction error, pools to no significant effect and reverses direction with age (Schwartz, Shinn-

Cunningham & Tager-Flusberg, 2018; Sapey-Triomphe et al., 2025). The companion documents give that audit in full. The relevant lesson here is that a predictive-processing story can be repeated confidently for a decade on an evidence base that does not hold, and this document is a predictive-processing story.

9. Where the model falls apart

It's flexible enough to fit anything. Free parameters for set point, two gains, damping, delay and arbitration threshold can accommodate almost any behavior. That's the same objection this document levels at the literature it reviews, and adding a control layer makes it worse, not better.

The defense against that objection is not argument but falsifiability. A framework that can accommodate any possible observation makes no empirical claim, because a claim is defined by what it rules out. The criterion is Popper's (2002/1935) and it applies here without modification: the content of a hypothesis lies in the observations it forbids. A model carrying this many free parameters therefore has to specify those observations in advance, before the data are in, or it is a post-hoc interpretive scheme rather than a scientific proposal.

Here are four. **The model forbids all of them:**

- **Observer and controller differences that never dissociate.** If belief accuracy and response intensity always move together across people, there aren't two dials. There's one, and this whole document is a long way of describing it.
- **Damping interventions adding nothing over content interventions.** If inserting delay between error and correction, response prevention, urge surfing, scheduled worry, produces no benefit beyond what restructuring produces, then the control reading of those techniques is decoration.
- **Gain manipulations leaving repetitive behavior untouched.** Amplify the error a person is shown and repetitive corrective behavior should rise. If it doesn't move, the oscillation account is wrong.
- **Volatility history failing to predict observer gain.** If people who grew up in genuinely unpredictable environments do not show higher learning rates from surprising evidence, the learning theory in Section 4 is wrong at its foundation.

If those results come in, the model is wrong and either needs adjustment or just throw it out.

The identifiability problem survives. In a Gaussian update only the ratio of sensory to prior precision enters the posterior. "They trusted their senses too much" and "they held their expectation too loosely" are the same claim from a different point of view for most data.

The separation may not be real. Friston argues there's no separate controller to bolt an estimator onto, because the goal is a prior inside the model rather than a cost outside it (Friston, 2011; Adams, Shipp & Friston, 2013). Kennaway (2025) argues from the other

direction that the two frameworks are fundamentally incompatible. If either is right, the central move here is an engineering convenience projected onto a system that doesn't work that way.

It says almost nothing about meaning. Control systems regulate variables. Much of the work is narrative, relationship, identity, and grief. This describes a regulatory layer. It isn't a theory of persons.

And it says nothing about several ordinary feelings. Amusement, awe, aesthetic pleasure and curiosity have no obvious regulated variable, as Section 2 admits. A theory of psychology that can't account for why something is funny is not a theory of psychology yet.

Nobody has tested any of it. Not the two-gain dissociation, not the oscillation reading, not the clinical chain. This is a framework proposal.

10. Where this goes

Four studies would settle most of it, and none of them is expensive.

The dissociation study. Measure belief accuracy and response intensity separately in the same people, across situations. If they separate, the two dials are real. If they don't, they aren't. Everything else in this document depends on that result.

The treatment-matching study. Sort people by which parameter looks mis-set, then randomize to cognitive versus arousal-focused intervention. The model predicts an interaction. Fifty years of dismantling trials have not produced one, and the model's best answer is that nobody sorted first. That answer is only respectable if somebody now sorts.

The precision-versus-conviction study. Does a measure of how movable a belief is predict treatment response better than a rating of how strongly it's held? That's a straightforward comparison and it would tell a practitioner something they can use on Monday.

The gain-manipulation study. Amplify displayed error and see whether corrective behavior scales. That one is a bench experiment.

Three things don't require waiting.

Ask the two questions separately. How confident is this belief, and how hard is this person acting on it.

Set the review interval before you start, and make it longer than the system's response time.

And when a pattern looks maladaptive, ask what environment would have made it correct.

A closing note

Here's where I'm supposed to hand you the framework that ties it together and tell you it changes everything, but that wouldn't help.

Ask yourself if this seems to fit, and if you can explain it in a useful way. Is it the predictive system (learned expectation, schema) or the controller (the reaction to the input is too high)?

I wanted to write this down and invite disagreement, comment and independent opinions. Perhaps this can be useful.

The companion documents apply this model to autism and repetitive behavior, including published first-person accounts, the full treatment-evidence audit, and the phase-locking extension.

Appendix A. How this differs from *Psycho-Cybernetics*

Anyone who has spent time in the self-help section will have noticed that a plastic surgeon got to this vocabulary in 1960 and sold more than thirty million copies with it. Maxwell Maltz's *Psycho-Cybernetics* borrowed Wiener's engineering language, applied it to the self, and became one of the most influential business and sports psychology books ever written. It deserves a straight answer rather than a sneer, so here it is.

What Maltz got right. His central claim is that behavior organizes itself around a defended reference, and that the reference sets the ceiling rather than effort does. That's a set point, described in a mass-market paperback a decade before control theory reached clinical psychology in any serious way. He also saw that the system is goal-seeking and self-correcting rather than driven, which is the same insight Powers formalized in 1973. The observation was sound. Coaches have been getting mileage out of it for sixty years for a reason.

What's missing is the estimator. Maltz has a servo-mechanism and a target and essentially nothing between the world and the target. There is no separate machinery that takes sensory evidence, weighs it against a prior, and produces an estimate that the set point then tracks. Which is precisely the hole this document exists to fill. Without an estimator you cannot distinguish a person whose reading of the situation is wrong from a person whose reading is right and whose response is too large, and that distinction is the entire clinical payload of Section 5. Maltz collapses both into self-image.

Three further differences follow from that one.

One master set point instead of many. Maltz runs everything through a single global reference, the self-image, and treats it as the thing that determines outcomes across domains. The account here has many parallel governors, each defending its own variable, competing for one body. That's why a person can be professionally confident and socially

terrified at the same time without any contradiction, which a single-self-image model has to explain away.

Set points move on imagination. Maltz's mechanism of change is mental rehearsal: picture the outcome vividly enough and the servo will steer toward it. Under this account the estimate moves on evidence weighted by precision, which is why visualizing yourself calm at the edge of a roof does close to nothing and standing at the edge of a roof does a great deal. The distinction isn't academic. It is the difference between a technique with fifty years of outcome trials behind it and a technique without.

There is no evidence base. Maltz worked from case observation, uncontrolled, with no falsification criteria and no way to be wrong. Section 9 of this document lists four ways the model proposed here could be shown false. *Psycho-Cybernetics* offers none, and a system that cannot fail a test is not making a claim about the world.

The twenty-one days

Maltz is also the origin of the most repeated false number in behavior change, and the story is worth telling accurately because it is the same failure this whole document keeps finding in the research literature.

What Maltz actually wrote is that it requires a minimum of about twenty-one days for an old mental image to dissolve and a new one to jell. He was describing how long his surgical patients took to stop seeing their old face in the mirror, and how long amputees took to stop feeling a limb that was gone. A careful observation about psychological adjustment to a changed body, hedged twice, in the same sentence.

What got repeated dropped "minimum," dropped "about," and swapped "mental image" for "habit." A clinical field note became a law of behavior change, and it is now printed on wall calendars.

The actual figure comes from Lally, van Jaarsveld, Potts and Wardle (2010), who tracked 96 people adopting a new daily behavior and measured the time to automaticity. The median was 66 days. The range ran from 18 to 254. There is no single number, which is the finding, and any program built on a fixed one is built on a misquotation of a surgeon's observation about faces.

Maltz was more careful than the people who cite him. That pattern shows up in Section 8 of this document, and on nearly every page of the companion literature review. The damage is almost never done by the person who made the original claim. It is worth noticing that the most durable damage in a field usually isn't done by the person who made the original claim.

Key sources

Bloomer, B. F., et al. (2025). Postural sway dynamics in adults across the autism spectrum. *Molecular Autism*, 16(1), 44. <https://doi.org/10.1186/s13229-025-00676-y>

- Kapp, S. K., et al. (2019). "People should be allowed to do what they like": Autistic adults' views and experiences of stimming. *Autism*, 23(7), 1782–1792. <https://doi.org/10.1177/1362361319829628>
- Adams, R. A., Shipp, S., & Friston, K. J. (2013). Predictions not commands: Active inference in the motor system. *Brain Structure and Function*, 218(3), 611–643. <https://doi.org/10.1007/s00429-012-0475-5>
- Beker, S., Foxe, J. J., & Molholm, S. (2021). Oscillatory entrainment mechanisms and anticipatory predictive processes in children with autism spectrum disorder. *Journal of Neurophysiology*, 126(5), 1783–1798. <https://doi.org/10.1152/jn.00329.2021>
- Crompton, C. J., Ropar, D., Evans-Williams, C. V., Flynn, E. G., & Fletcher-Watson, S. (2020). Autistic peer-to-peer information transfer is highly effective. *Autism*, 24(7), 1704–1712. <https://doi.org/10.1177/1362361320919286>
- Friston, K. (2011). What is optimal about motor control? *Neuron*, 72(3), 488–498. <https://doi.org/10.1016/j.neuron.2011.10.018>
- Friston, K., Daunizeau, J., Kilner, J., & Kiebel, S. J. (2010). Action and behavior: A free-energy formulation. *Biological Cybernetics*, 102(3), 227–260. <https://doi.org/10.1007/s00422-010-0364-z>
- Kennaway, R. (2025). Perceptual control theory and the free energy principle: A comparison. *Current Opinion in Behavioral Sciences*, 65, 101575. <https://doi.org/10.1016/j.cobeha.2025.101575>
- Lidstone, D. E., et al. (2020). Children with autism spectrum disorder show impairments during dynamic versus static grip-force tracking. *Autism Research*, 13(12), 2177–2189. <https://doi.org/10.1002/aur.2370>
- Bhat, A. N. (2020). Is motor impairment in autism spectrum disorder distinct from developmental coordination disorder? A report from the SPARK study. *Physical Therapy*, 100(4), 633–644. <https://doi.org/10.1093/ptj/pzz190>
- Doumas, M., McKenna, R., & Murphy, B. (2016). Postural control deficits in autism spectrum disorder: The role of sensory integration. *Journal of Autism and Developmental Disorders*, 46(3), 853–861. <https://doi.org/10.1007/s10803-015-2621-4>
- Fournier, K. A., Amano, S., Radonovich, K. J., Bleser, T. M., & Hass, C. J. (2014). Decreased dynamical complexity during quiet stance in children with autism spectrum disorders. *Gait & Posture*, 39(1), 420–423. <https://doi.org/10.1016/j.gaitpost.2013.08.016>
- Haswell, C. C., Izawa, J., Dowell, L. R., Mostofsky, S. H., & Shadmehr, R. (2009). Representation of internal models of action in the autistic brain. *Nature Neuroscience*, 12(8), 970–972. <https://doi.org/10.1038/nn.2356>
- Izawa, J., Pekny, S. E., Marko, M. K., Haswell, C. C., Shadmehr, R., & Mostofsky, S. H. (2012). Motor learning relies on integrated sensory inputs in ADHD, but over-selectively on

proprioception in autism spectrum conditions. *Autism Research*, 5(2), 124–136.
<https://doi.org/10.1002/aur.1222>

Licari, M. K., et al. (2020). Prevalence of motor difficulties in autism spectrum disorder: Analysis of a population-based cohort. *Autism Research*, 13(2), 298–306.
<https://doi.org/10.1002/aur.2230>

Lim, Y. H., et al. (2020). Effect of visual information on postural control in children with autism spectrum disorder. *Journal of Autism and Developmental Disorders*, 50(9), 3320–3325. <https://doi.org/10.1007/s10803-019-04182-y>

Marko, M. K., Crocetti, D., Hulst, T., Donchin, O., Shadmehr, R., & Mostofsky, S. H. (2015). Behavioural and neural basis of anomalous motor learning in children with autism. *Brain*, 138(3), 784–797. <https://doi.org/10.1093/brain/awu394>

Mosconi, M. W., Mohanty, S., Greene, R. K., Cook, E. H., Vaillancourt, D. E., & Sweeney, J. A. (2015). Feedforward and feedback motor control abnormalities implicate cerebellar dysfunctions in autism spectrum disorder. *Journal of Neuroscience*, 35(5), 2015–2025. <https://doi.org/10.1523/JNEUROSCI.2731-14.2015>

Shafer, R. L., et al. (2025). Visual feedback and motor memory contributions to sustained motor control deficits in autism spectrum disorder across childhood and into adulthood. *Journal of Neurodevelopmental Disorders*, 17(1), 26. <https://doi.org/10.1186/s11689-025-09607-7>

Stania, M., et al. (2023). Modulation of center-of-pressure signal in children on the autism spectrum: A case-control study. *Gait & Posture*, 103, 67–72.
<https://doi.org/10.1016/j.gaitpost.2023.04.018>

Surgent, O. J., Walczak, M., Zarzycki, O., Ausderau, K., & Travers, B. G. (2021). IQ and sensory symptom severity best predict motor ability in children with and without autism spectrum disorder. *Journal of Autism and Developmental Disorders*, 51(1), 243–254.
<https://doi.org/10.1007/s10803-020-04536-x>

Pettinato, F., et al. (2024). Dynamical complexity of postural control system in autism spectrum disorder. *Journal of NeuroEngineering and Rehabilitation*, 21(1), 225.
<https://doi.org/10.1186/s12984-024-01520-9>

Pezzulo, G., Rigoli, F., & Friston, K. (2015). Active inference, homeostatic regulation and adaptive behavioral control. *Progress in Neurobiology*, 134, 17–35.
<https://doi.org/10.1016/j.pneurobio.2015.09.001>

Salkovskis, P. M., et al. (1999). An experimental investigation of the role of safety-seeking behaviours in the maintenance of panic disorder with agoraphobia. *Behaviour Research and Therapy*, 37(6), 559–574. [https://doi.org/10.1016/S0005-7967\(98\)00153-3](https://doi.org/10.1016/S0005-7967(98)00153-3)

Stephan, K. E., et al. (2016). Allostatic self-efficacy: A metacognitive theory of dyshomeostasis-induced fatigue and depression. *Frontiers in Human Neuroscience*, 10, 550.
<https://doi.org/10.3389/fnhum.2016.00550>

- Sterling, P. (2012). Allostasis: A model of predictive regulation. *Physiology & Behavior*, 106(1), 5–15. <https://doi.org/10.1016/j.physbeh.2011.06.004>
- Bennett-Levy, J., Butler, G., Fennell, M., Hackmann, A., Mueller, M., & Westbrook, D. (Eds.). (2004). *Oxford guide to behavioural experiments in cognitive therapy*. Oxford University Press. <https://doi.org/10.1093/med:psych/9780198529163.001.0001>
- Craske, M. G., Treanor, M., Conway, C. C., Zbozinek, T., & Vervliet, B. (2014). Maximizing exposure therapy: An inhibitory learning approach. *Behaviour Research and Therapy*, 58, 10–23. <https://doi.org/10.1016/j.brat.2014.04.006>
- Lally, P., van Jaarsveld, C. H. M., Potts, H. W. W., & Wardle, J. (2010). How are habits formed: Modelling habit formation in the real world. *European Journal of Social Psychology*, 40(6), 998–1009. <https://doi.org/10.1002/ejsp.674>
- Maltz, M. (1960). *Psycho-cybernetics*. Prentice-Hall.
- Ougrin, D. (2011). Efficacy of exposure versus cognitive therapy in anxiety disorders: Systematic review and meta-analysis. *BMC Psychiatry*, 11, 200. <https://doi.org/10.1186/1471-244X-11-200>
- Salkovskis, P. M., Clark, D. M., & Hackmann, A. (1991). Treatment of panic attacks using cognitive therapy without exposure or breathing retraining. *Behaviour Research and Therapy*, 29(2), 161–166. [https://doi.org/10.1016/0005-7967\(91\)90044-4](https://doi.org/10.1016/0005-7967(91)90044-4)
- Stott, R. (2007). When head and heart do not agree: A theoretical and clinical analysis of rational-emotional dissociation (RED) in cognitive therapy. *Journal of Cognitive Psychotherapy*, 21(1), 37–50. <https://doi.org/10.1891/088983907780493313>
- Öst, L.-G., Havnen, A., Hansen, B., & Kvale, G. (2015). Cognitive behavioral treatments of obsessive-compulsive disorder: A systematic review and meta-analysis of studies published 1993–2014. *Clinical Psychology Review*, 40, 156–169. <https://doi.org/10.1016/j.cpr.2015.06.003>
- Smith, C. A., & Ellsworth, P. C. (1985). Patterns of cognitive appraisal in emotion. *Journal of Personality and Social Psychology*, 48(4), 813–838. <https://doi.org/10.1037/0022-3514.48.4.813>
- Angeletos Chrysaitis, N., & Seriès, P. (2023). 10 years of Bayesian theories of autism: A comprehensive review. *Neuroscience & Biobehavioral Reviews*, 145, 105022. <https://doi.org/10.1016/j.neubiorev.2022.105022>
- Cannon, J., O'Brien, A. M., Bungert, L., & Sinha, P. (2021). Prediction in autism spectrum disorder: A systematic review of empirical evidence. *Autism Research*, 14(4), 604–630. <https://doi.org/10.1002/aur.2482>
- Cui, K., Lin, X., Gao, R., Jing, S., Luo, F., & Wang, J. (2026). A systematic review and meta-analysis of empirical evidence for the simple Bayesian model of autism. *Neuropsychology Review*, 36(2), 184–195. <https://doi.org/10.1007/s11065-025-09672-8>

Lawson, R. P., Mathys, C., & Rees, G. (2017). Adults with autism overestimate the volatility of the sensory environment. *Nature Neuroscience*, 20(9), 1293–1299.
<https://doi.org/10.1038/nn.4615>

Plank, I. S., Pior, A., Yurova, A., Nowak, J., Shi, Z., & Falter-Wagner, C. M. (2026). Predicting emotional valence in autism: A preregistered study in the Bayesian Brain framework. *Molecular Autism*, 17(1), 32. <https://doi.org/10.1186/s13229-026-00730-3>

Popper, K. R. (2002). *The logic of scientific discovery*. Routledge. (Original work published 1935)

Sapey-Triomphe, L.-A., Bouet, R., Mattout, J., Sonié, S., Schmitz, C., & Lecaiguard, F. (2025). Systematic review and meta-analysis of mismatch negativity in autism: Insights into predictive mechanisms. *Autism Research*, 18(12), 2431–2450.
<https://doi.org/10.1002/aur.70131>

Schwartz, S., Shinn-Cunningham, B., & Tager-Flusberg, H. (2018). Meta-analysis and systematic review of the literature characterizing auditory mismatch negativity in individuals with autism. *Neuroscience & Biobehavioral Reviews*, 87, 106–117.
<https://doi.org/10.1016/j.neubiorev.2018.01.008>

Beck, A. T. (1967). *Depression: Clinical, experimental, and theoretical aspects*. Harper & Row.

Beck, A. T., Rush, A. J., Shaw, B. F., & Emery, G. (1979). *Cognitive therapy of depression*. Guilford Press.

Quiller-Couch, A. (1916). *On the art of writing: Lectures delivered in the University of Cambridge, 1913–1914*. Cambridge University Press.
<https://www.gutenberg.org/ebooks/17470>

Slime Mold Time Mold. (2025). *The mind in the wheel* [Blog series].
<https://slimemoldtimemold.com/2025/02/06/the-mind-in-the-wheel-prologue-everybody-wants-a-rock/>