

**Beyond Learned Helplessness:
A Control-Response Fit Model of Agency, Persistence, and Adaptive Disengagement**

Kevin L. Michel

Saint Lucia Centre for Psychological Science

Working Paper

Version 1.0 | August 2026

Copyright © 2026 Kevin L. Michel

<https://doi.org/10.5281/zenodo.21863785>

Abstract

Perceived control is associated with persistence, affective resilience, and mental health, but more perceived control and more persistence are not uniformly adaptive. Withdrawal can preserve a false expectation of helplessness when effective action is available, while repeated direct effort can waste resources or obscure collective and alternative routes when personal leverage is low. This theory-guided integrative review proposes the Control-Response Fit (CRF) model. CRF separates three questions that are often blended: whether inferred leverage is calibrated to effective leverage, whether a selected response fits feasible access, values, costs, risks, uncertainty, and alternatives, and whether the response improves a prespecified outcome for a specified beneficiary and time horizon. Fit is defined prospectively as low policy regret, not retrospectively as whatever produced a favorable outcome. The model represents leverage by agency mode, distinguishes who acts from what is regulated, treats personality as a source of probabilistic priors and switching thresholds, and allows action to reshape future affordances. Eight falsifiable propositions and a five-study program specify experimental, computational, longitudinal, cross-cultural, and intervention tests. CRF's proposed contribution is a dynamic architecture joining causal inference, feasible capacity, response selection, structural constraint, and affordance construction. Resilience, on this view, is the capacity to locate leverage, act through a fitting route, and redirect effort without globalizing the loss of one path.

Keywords: perceived control, learned helplessness, agency, coping flexibility, policy fit

Beyond Learned Helplessness: A Control-Response Fit Model of Agency, Persistence, and Adaptive Disengagement

The capacity to act despite adversity is often described as an unqualified strength. People who expect their behavior to matter tend to explore, persist, recover after setbacks, and report better mental health. Low perceived control, by contrast, is associated with depression and anxiety, and exposure to uncontrollable aversive events can disrupt later learning and action (Gallagher et al., 2014; Maier & Seligman, 2016; Msetfi et al., 2024). These findings support an appealing practical message: agency protects, while helplessness harms.

That message is useful but incomplete. An organism can fail by withdrawing from a controllable problem, but it can also fail by persisting with a response that has no causal leverage. A student who stops studying after one poor grade may abandon real control. A caregiver who treats an irreversible diagnosis as a problem to defeat may exhaust resources that could instead support connection, comfort, or a different valued goal. A worker facing discrimination may accurately perceive that individual effort alone cannot alter an institutional contingency; collective action, an advocate, exit, or policy change may be the relevant forms of agency. In each case, the adaptiveness of a response depends on what the environment permits and on whether the person can detect, learn, and use that structure.

The distinction matters because several neighboring literatures have developed different pieces of the same problem. Learned-helplessness research asks what follows exposure to uncontrollability. Personality research examines generalized expectations such as locus of control, self-efficacy, optimism, and neuroticism. Coping research asks whether strategies fit appraised demands. Goal-adjustment research shows that disengaging from an unattainable goal and reengaging elsewhere can protect well-being. Cultural psychology distinguishes direct,

individual influence from accommodation and relational forms of agency. Social-cognitive theory recognizes personal, proxy, and collective agency. Neuroscience and computational work examine how controllability is inferred and represented. These literatures converge on a proposition that is more precise than the simple valorization of control: adaptation depends on matching a revisable response to a sufficiently accurate model of causal opportunity.

This article develops that proposition as the Control-Response Fit (CRF) model. The model does not redefine passivity as adaptive, nor does it deny the benefits of perceived control. Instead, it makes two complementary errors visible. Under-engagement occurs when a person behaves as though useful leverage is absent when an effective and feasible route is available. Over-engagement occurs when a policy continues despite low effective leverage, disproportionate cost, or a superior alternative. Flexible agency occupies the space between them. It requires learning what can be changed, choosing who should act and what should be regulated, and revising course when evidence, goals, or feasible options change.

Approach to the Evidence

This article is a theory-guided integrative narrative review, not a systematic review or meta-analysis. The synthesis draws on five bodies of peer-reviewed work: experimental controllability and learned helplessness; personality and control beliefs; coping flexibility and goal adjustment; human and animal neuroscience and computational modeling; and cultural or structural accounts of agency. Sources available through August 8, 2026 were located through targeted searches of bibliographic and publisher platforms and through citation tracing from foundational papers, recent reviews, and meta-analyses. Primary experiments, meta-analyses, construct-clarification articles, and recent translational studies received priority. Because the search and screening process was not designed to estimate literature prevalence or an exhaustive

effect, absence from this review should not be interpreted as evidence that a finding or theory does not exist. The purpose is narrower: to integrate established distinctions, expose unresolved logical problems, and derive a falsifiable personality-process model.

What the CRF Model Adds

The CRF model shares important premises with several mature theories. Bayesian accounts formalize how agents infer whether actions cause outcomes under uncertainty (Huys & Dayan, 2009; Ligneul et al., 2022). Coping-flexibility and regulatory-flexibility theories emphasize context sensitivity, a repertoire of strategies, and responsiveness to feedback (Bonanno & Burton, 2013; Cheng et al., 2014). The goodness-of-fit hypothesis links coping strategy to appraised controllability (Forsythe & Compas, 1987). Life-span control theory and goal-adjustment research connect opportunity and attainability to engagement, disengagement, and reengagement (Heckhausen et al., 2010; Wrosch et al., 2003). Social-cognitive theory distinguishes personal, proxy, and collective agency (Bandura, 2001). CRF therefore does not claim that situation-strategy fit, flexible switching, or plural agency modes are individually new.

Its proposed contribution is their joint architecture. CRF represents control as mode-specific, separates environmental causal leverage from an actor's feasible access and capacity, treats inferred leverage as a distribution with uncertainty, and evaluates policy choice prospectively in light of values, costs, risks, alternatives, beneficiary, and time horizon. It further separates belief calibration, policy fit, and adaptation, and allows personal, proxy, and collective action to change the future action set. This last feature, affordance construction, matters because current low individual leverage can coexist with a rational contribution to coalition formation or institutional change.

The empirical case for strategy-situation fit is not settled. Haines et al. (2016) reported that daily-life emotion-regulation fit was associated with well-being, but a well-powered preregistered direct replication found no evidence for the focal fit hypothesis (Willroth et al., 2025). In a preliminary 7-day correlational diary study of 75 undergraduates, the predicted interaction appeared for emotion-focused coping but not for problem-focused coping (Leslie-Miller et al., 2025). CRF treats these mixed results as a design challenge rather than as confirmation: leverage, access, response, values, and outcomes must be measured independently, and the fit rule must be specified before results are known. Table 1 summarizes the model's relationship to these neighboring approaches.

Table 1

CRF in Relation to Neighboring Theories

Theory family	Primary contribution	Boundary relative to CRF	CRF extension to test
Bayesian controllability	Inference over action-sensitive versus action-insensitive causal models	Usually emphasizes individual task control and action selection	Mode-specific leverage, feasible access, and social or structural affordance change
Coping and regulatory flexibility	Context sensitivity, strategy repertoire, feedback-responsive switching	Fit is often based on appraisal or observed association	Prospectively scored policy regret using independently measured affordances and costs
Life-span control and goal adjustment	Opportunity-sensitive engagement, disengagement, and reengagement	Less explicit about trial-level causal inference and uncertainty	Dynamic belief-policy updating and mode switching
Social-cognitive agency	Personal, proxy, and collective efficacy and action	Efficacy beliefs do not establish environmental leverage or access	Separate belief, causal leverage, feasible capacity, and policy choice
CRF joint architecture	Calibration, feasible policy selection, adaptation, and affordance construction	Requires demanding measurement and prospective decision rules	Joint prediction that mode-matched policy and access add value beyond beliefs or persistence alone

The Construct Problem: What Kind of Control?

More than 100 constructs have been described using the language of control (Skinner, 1996). Treating them as interchangeable produces jingle-jangle errors: the same label can refer to

different mechanisms, and different labels can refer to overlapping mechanisms. The CRF model begins by separating the following questions.

Objective causal leverage, often called objective controllability in laboratory paradigms, asks whether a specified response changes the probability, timing, intensity, or value of a specified outcome. It is a property of an indexed action-outcome relation within a context, actor level, and time horizon, not a global property of a person or situation. A context may offer leverage over process but not outcome, over one time scale but not another, or only through coordinated action. *Effective leverage* is the subset of that causal opportunity the focal actor can feasibly access and execute.

Perceived controllability asks whether the person believes that available responses influence outcomes. It is an estimate, not simply an illusion or a personality trait. It may be accurate, pessimistically biased, optimistically biased, or uncertain. Experimental evidence shows substantial variation in perceived control even among people assigned to the same objective condition (Meier et al., 2025).

Desired control asks how much influence a person wants. It is distinct from a normative expectation about how much control one ought to have and from a perceived likelihood that control will be available. The same perceived control may be sufficient when desired control is modest and distressing when it falls far below what is wanted. In one cross-sectional U.S. sample of 942 adults, an independently rated control-discrepancy item moderated nonlinear associations among perceived control, desired control, and concurrent elevated depressive-symptom classification (Davis et al., 2026). The study did not validate a computed perceived-minus-desired gap or establish temporal order and causation. CRF therefore recommends modeling

perceived control, desired control, and any separately appraised discrepancy jointly rather than treating a raw difference score as a mechanism.

Feasible access and capacity ask whether the focal actor can reach and execute a causally effective response now. A pathway may be controllable in principle but unavailable because of skill, authority, money, safety, disability access, coordination, or institutional permission. This construct concerns actual opportunity, not confidence in capability.

Self-efficacy asks whether a person believes they can execute a particular response (Bandura, 1977). A person may correctly perceive that an outcome is controllable in principle yet doubt their capability to perform the controlling action. Conversely, high confidence in performance cannot create a causal pathway that the environment does not provide.

Attribution asks how outcomes are explained. Internality, stability, and globality versus specificity influence whether a failure is expected to recur and whether it is generalized beyond its original domain (Abramson et al., 1978). Attribution is an interpretation of evidence, not the evidence itself.

Locus of enactment asks who acts. Personal agency operates through one's own action, proxy agency recruits another person or institution with relevant power, and collective agency operates through coordinated action (Bandura, 2001). These loci may form hybrid or sequential pipelines. For example, an individual may gather evidence, recruit an advocate, and contribute to a coalition.

Target of regulation asks what is changed. Primary control targets external conditions. Secondary control targets the person's attention, interpretation, emotion, or commitment to improve person-environment fit when external change is unavailable or not worth its cost (Morling & Evered, 2006). It is not a fourth locus parallel to personal, proxy, and collective

action. Relational agency refers here to intentional influence exercised through interdependent roles and relationships; it can involve proxy or collective enactment and can target either external conditions or internal accommodation.

Response policy asks what the person actually does: explore, persist, increase effort, change strategy, recruit assistance, coordinate with others, accept, disengage, or reengage elsewhere. None is synonymous with a control belief. The same belief can produce different responses depending on values, costs, fatigue, social norms, and available alternatives.

Finally, *outcome* asks what counts as adaptation. Immediate task success, later learning, affect, physiological stress, goal progress, symptom burden, and social impact are not interchangeable. A response can improve mood while reducing performance, preserve resources now while delaying long-term change, or benefit an individual while leaving a harmful structure intact. Strong tests must declare the outcome and time scale in advance. Table 2 summarizes the distinctions.

Table 2

Constructs That Must Remain Distinct in a Science of Agency

Construct	Core question	Illustrative indicator	Not equivalent to
Objective causal leverage	Can a specified response change a specified outcome?	Experimentally specified action-outcome contingency	Choice, effort, success rate, or subjective control
Feasible access and capacity	Can this actor reach and execute that response now?	Authority, skill, resources, safety, accessibility, and coordination	Self-efficacy belief or leverage in principle
Inferred leverage	What distribution of leverage does the actor infer?	Trial-level estimate plus confidence or posterior uncertainty	Objective leverage or optimism
Desired control	How much influence does the actor want?	Desired-control rating modeled jointly with perceived control	Entitlement, expected likelihood, or a raw discrepancy mechanism
Self-efficacy	Does the actor believe they can execute the response?	Domain- and action-specific capability belief	Actual capacity or outcome controllability

Construct	Core question	Illustrative indicator	Not equivalent to
Attribution	Why did the outcome occur, and will it recur?	Stability, globality, internality, and specificity	Contingency learning
Locus of enactment	Who acts?	Personal, proxy, collective, or hybrid sequence	Target of regulation
Target of regulation	What is changed?	External conditions or internal accommodation	Locus of enactment
Response policy	What response is selected and maintained?	Explore, persist, switch, recruit, coordinate, accept, disengage, or reengage	A stable trait or belief
Adaptation criterion	What outcome matters, for whom, and when?	Performance, health, resource use, relationship quality, or social change	Calibration or fit by definition

From Learned Helplessness to Controllability Learning

Classic experiments with dogs demonstrated that exposure to inescapable shock could impair later escape and avoidance responding even when escape became possible (Overmier & Seligman, 1967). The early theory proposed that organisms learn that responses and outcomes are independent, producing motivational, cognitive, and emotional deficits (Maier & Seligman, 1976). Human experiments extended the phenomenon across aversive noise and problem-solving tasks, while also exposing variability in generalization and the importance of locus of control (Hiroto & Seligman, 1975). The attributional reformulation later argued that the consequences of uncontrollability depend on how causes are explained: global and stable attributions should support broader and more persistent helplessness than specific and unstable attributions (Abramson et al., 1978).

A prominent contemporary rodent-circuit account revises an important part of the early story. In canonical studies conducted largely with male rats, passivity and heightened fear after prolonged uncontrollable stress appear to arise partly from subcortical stress-responsive processes, whereas detecting control recruits medial prefrontal circuitry that inhibits downstream stress effects and creates a memory of controllability that can protect against later challenges (Amat et al., 2005, 2006; Maier & Seligman, 2016). In this preparation, it may be more accurate

to say that control is learned than that helplessness itself is learned. This revision preserves the importance of contingency while rejecting a simple narrative in which an organism first forms an explicit belief in futility and then becomes passive.

The circuit claim must not be overtranslated. Female-rat studies have not shown the same behavioral immunization or prelimbic-to-dorsal-raphé pattern as the canonical male-rat work, and more recent evidence points to sex-selective circuitry rather than a simple presence-versus-absence account (Baratta et al., 2018; McNulty et al., 2023). A 2026 brief report found that escapable and inescapable stress produced different peripheral corticosterone responses in female rats without corresponding differences in sampled brain tissue, illustrating that behavioral, endocrine, and circuit indices can dissociate (Levy et al., 2026). Human studies have identified condition-related BOLD activity in ventromedial prefrontal, striatal, insular, and other networks, but no human experiment has established the specific medial-prefrontal-to-dorsal-raphé serotonergic mechanism demonstrated in rodents (Baratta et al., 2023; Meine et al., 2021). Human control can also be symbolic, delayed, social, collective, and verbally instructed, making one-to-one homology unlikely.

Human paradigms nonetheless support several parts of the controllability account. Perceived control can increase persistence after setbacks and alter neural responses to failure (Bhanji & Delgado, 2014). Under acute stress, the relation between stress and persistence depends on perceived control (Bhanji et al., 2016). In a small fixed-sequence fMRI study of 22 participants, avoidance attempts declined across uncontrollable trials and rebounded when controllable feedback was subsequently introduced; ventromedial prefrontal activity covaried with the rebound, although recovery could not be separated cleanly from order or time (Wang & Delgado, 2021). Lower post-task perceived control was associated with poorer subsequent

working-memory change and altered network measures in another experiment, whereas randomized controllability assignment did not affect working memory (Wanke & Schwabe, 2020). A translational triadic paradigm using controllable, yoked-uncontrollable, and no-stress conditions initially struggled to separate controllability from stress, then produced differences in perceived control, helplessness, exhaustion, and escape behavior after design refinement (Meine et al., 2020). Such work shows both the feasibility and fragility of human induction.

Controllability is also not uniformly analgesic or affectively beneficial. Across behavioral and fMRI pain experiments with a combined sample of 113, control increased the precision of pain expectations relative to matched predictability. This made lower-intensity stimulation feel less painful but made higher-intensity stimulation slightly more painful (Habermann & Büchel, 2025). The result directly illustrates why the consequence of control depends on the outcome distribution and evaluative criterion rather than on control level alone.

Recent results further challenge a unitary helplessness response. In a randomized study with a final sample of 168 healthy young adults, below the planned sample of 179, assigned objective control was associated with less reported helplessness. Exploratory post-randomization perceived-control trajectory classes were associated with helplessness and state-depression scores beyond assigned condition; roughly half of participants in the objectively uncontrollable condition belonged to rising or medium perceived-control classes. A group-by-gender negative-affect interaction, observed after excluding extreme outliers, warrants replication because the study was not powered for that interaction. In this study, a particular brief autobiographical induction failed to change its self-efficacy manipulation check or affective outcomes, and objective control did not directly reduce state-depression scores (Meier et al., 2025). These findings are not evidence that inaccurate control beliefs are universally protective. They show

that objective assignment, experienced contingency, perceived control, affect, and clinical inference are separable.

Alternative mechanisms reinforce that conclusion. Learned-helplessness tasks can confound uncontrollability with low reward prevalence and the cost of exploration. Teodorescu and Erev (2014) found conditions in which reward prevalence predicted subsequent exploration better than reported control. A Bayesian formulation likewise treats behavioral control as causal inference under uncertainty: agents compare action-sensitive and action-insensitive models and generalize learned priors to new environments (Huys & Dayan, 2009). In social settings, people can use prospective, model-based computation to estimate how current actions will influence another person's future behavior, with ventromedial prefrontal activity tracking projected values (Na et al., 2021). Thus, apparent resignation can arise from at least four processes: an inference of low causal control, low expected value of exploration, high response cost, or a strong prior that control will not generalize. A credible personality theory must distinguish them.

A 2026 study extended the perceived-control account across tasks: among 116 healthy young men, latent classes defined by perceived control during an uncontrollable aversive task differed in helplessness, psychosomatic symptoms, locus of control, cortisol response, and posterior-insula activation during a separate psychosocial stressor (Meier et al., 2026). Because class membership was observational, the sample excluded women, and the analyses cannot establish that perceived control caused the cross-task differences, the findings are better read as a promising individual-differences signal than as intervention evidence.

Experimental helplessness should also not be treated as a laboratory reproduction of major depressive disorder. Acute passivity after an aversive manipulation is a state effect in a narrow task. Depression is heterogeneous, recurrent, temporally extended, and affected by

biological, relational, developmental, and structural factors. A registered report protocol proposed testing whether anhedonia covaries with a human learned-helplessness manipulation, but no publisher-hosted Stage 2 result was located for this review (Song & Vilares, 2021). The protocol therefore documents an intended test, not an observed effect. Controllability may be a transdiagnostic mechanism without being a complete model of any diagnosis.

The Control-Response Fit Model

The unit of analysis in CRF is a decision episode defined by a focal outcome, actor or agency mode, response, context, beneficiary, and time horizon. Control is not a scalar property of a person or an entire situation. It is an indexed causal relation. The model has three recursive layers: estimating effective leverage, selecting a feasible response policy, and monitoring evidence and consequences to decide whether to maintain, switch, or reshape the action set.

Layer 1: Estimating Effective Leverage

For agency mode (m), response (a), outcome (o), context (c), and horizon (h), causal leverage can be represented as the difference between the expected outcome under that response and the expected outcome under a specified comparison response:

$$C_{m,a,o,c,h} = \mathbb{E}[Y_o \mid do(m, a), c, h] - \mathbb{E}[Y_o \mid do(a_0), c, h].$$

This notation is a discipline for specification, not a claim that every real-world effect can be identified without assumptions. Leverage may concern the probability, timing, intensity, or value of an outcome. It can be positive, negligible, negative, delayed, or conditional on other actors. Personal, proxy, and collective leverage must be indexed separately because low personal leverage does not imply that no effective route exists.

Potential causal leverage is distinct from actual access and capacity. Let $A_{m,a,t}$ represent whether the focal person or coalition can safely reach and execute a response at time t . Effective

leverage, $L_{m,a,t} = f(C_{m,a,t}, A_{m,a,t})$, is therefore constrained by both what would work if executed and whether execution is feasible. A legal appeal may have causal force but be inaccessible without information, money, standing, language access, or counsel. Actual capacity remains distinct from self-efficacy, which is a belief about capacity.

People rarely know either causal leverage or feasibility with certainty. They maintain an inferred distribution, $\hat{\mathbf{L}}_t$, based on prior experience, current action-outcome evidence, reward prevalence, delay, volatility, instructions, social information, and observed institutional behavior. Personality and control history are hypothesized to influence priors, attention, learning rate, generalization, and switching thresholds. Culture can shape which responses are visible, valued, or legitimate. Structure constrains available actions, access, risks, and the consequences of acting.

Layer 2: Selecting a Response Policy

A response policy maps beliefs and constraints onto behavior. The feasible policy set can include bounded exploration, direct persistence, strategy change, help-seeking, proxy recruitment, coalition contribution, collective action, acceptance, exit, disengagement, and reengagement. Policies can be hybrid or sequential. An individual may first gather evidence, then recruit an advocate, then contribute to a coalition. When uncertainty is substantial, exploration may be valuable because action generates information even before a stable control belief forms.

Policy selection depends on more than leverage. It includes the value of the focal outcome, response cost, risk, delay, uncertainty, available alternatives, ethical constraints, beneficiary, and time horizon. Let $U(\pi)$ denote the expected net value of a feasible policy π under these prespecified elements. If π^* is the best feasible policy, or one member of an approximately equivalent set, policy regret is

$$R_t(\pi) = U_t(\pi^*) - U_t(\pi).$$

Low regret indicates better policy fit; high regret indicates misfit. This does not assume that people consciously maximize a single numerical utility. It specifies how researchers must score fit without circularity. In experiments, the benchmark can be derived from manipulated contingencies, access, payoffs, and costs. In field research, the benchmark may require externally documented affordances, multiple plausible decision rules, and sensitivity analyses. It cannot be inferred retrospectively from the same outcome later used to declare the response adaptive. The criterion should also respect declared values, safety, rights, and obligations; the policy with the greatest short-term personal payoff is not automatically the appropriate one.

Layer 3: Monitoring, Switching, and Affordance Construction

Monitoring asks whether prediction errors, accumulating costs, changed goals, or new alternatives justify a policy change. The process is bidirectional: beliefs guide action, but action produces evidence. Exploration can therefore precede confident belief revision, and a policy can change without a leverage estimate changing if costs, values, risks, or alternatives have changed. Flexibility is not switching frequency. It is evidence-responsive revision that reduces prospective regret under change.

Monitoring must include stopping rules. Persistence can be extended by sunk costs, identity investment, shame, intermittent reinforcement, or a social demand never to quit. Goal-adjustment research supplies a complementary mechanism. Across 31 samples, disengagement and especially reengagement capacities were associated with quality of life, although patterns varied by outcome and depression-risk status (Barlow et al., 2020). A synthesis of 235 studies likewise concluded that goal adjustment is multidimensional and context dependent, while judging much of the evidence to be low or moderate because cross-sectional designs,

heterogeneity, and publication bias remain common (Riddell et al., 2026). Life-span theory similarly links engagement to favorable opportunity and disengagement to declining attainability (Heckhausen et al., 2010). These findings motivate, but do not yet causally establish, the prediction that closing one path is most adaptive when it enables investment in another valued path.

Responses can also alter future affordances. Skill building can increase capacity; help-seeking can open proxy access; organizing can transform low individual marginal influence into coalition leverage; and institutional reform can change the causal structure itself. CRF calls this *affordance construction*. It prevents the model from treating the environment as fixed or classifying collective participation as irrational merely because one person's immediate marginal effect is small. Analysis should distinguish an individual's direct action, the individual's contribution to coalition formation, and coalition-level action.

Calibration, Policy Fit, and Adaptation

CRF separates three evaluative questions. *Calibration* concerns the divergence between inferred and effective leverage. A preregistered calibration error may be expressed as

$$E_C(t) = D(\hat{\mathbf{L}}_t, \mathbf{L}_t),$$

where the divergence function and commensurate scale are specified in advance and uncertainty is retained. *Policy fit* concerns regret relative to a feasible benchmark. *Adaptation* concerns prespecified consequences such as performance, affect, health, resource preservation, relationship quality, or collective change for a named beneficiary and horizon. These constructs can dissociate. A biased belief can accidentally yield a low-regret policy. An accurate belief can be followed by a poor policy because of habit, coercion, shame, or switching cost. A sound

policy can also fail by chance. CRF predicts that calibration often supports low-regret policy and that low regret often supports selected outcomes, but it does not define one from another.

Under-engagement and over-engagement are complementary forms of possible misfit, not diagnoses based on activity level. Under-engagement occurs when a feasible higher-value route is prematurely omitted. Over-engagement occurs when continued investment has lower prospective value than switching, redirecting, or stopping. Table 3 shows why behavior alone is insufficient for classification.

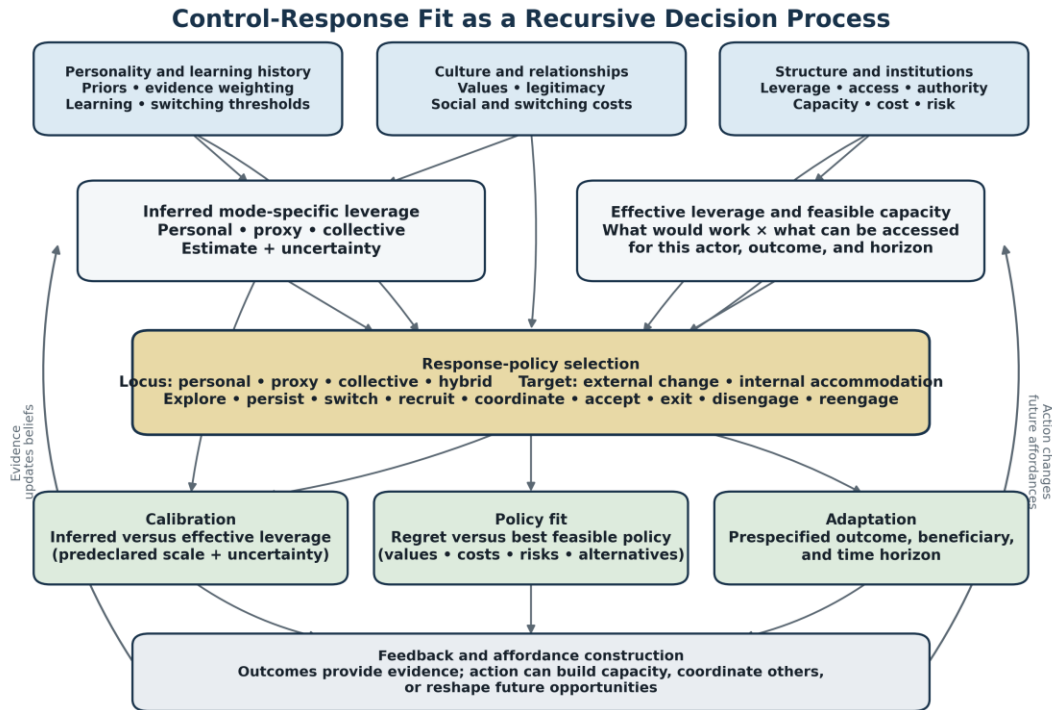
Table 3

Diagnostic Possibilities for Direct Personal Engagement

Effective direct leverage	Direct engagement	Possible interpretations	Evidence needed before judging fit
High	High	Fitting direct agency; excessive cost; ethically unacceptable route	Capacity, cost, value, risk, alternatives, and outcome horizon
High	Low	Premature withdrawal; low outcome value; high cost; absent access; superior alternative	Belief and uncertainty, actual capacity, values, alternatives, and prior sampling
Low	High	Costly persistence; bounded exploration; expressive or moral action; coalition contribution; affordance construction	Information value, stopping rule, beneficiary, coalition mechanism, and long-term effects
Low	Low	Fitting accommodation; safe exit; resource conservation; resignation generalized beyond the focal context	Active choice, reengagement, affect, generalization, safety, and available routes

Note. “High” and “low” refer only to effective direct personal leverage and direct engagement. A cell is not a diagnosis; policy fit also depends on feasible access, costs, values, risks, alternatives, beneficiary, and time horizon.

Figure 1 summarizes the architecture: context shapes inferred leverage, feasible capacity, and policy thresholds; policies generate outcomes and evidence; and personal, proxy, or collective action can sometimes construct new affordances.

Figure 1*The Control-Response Fit Model*

Note. Calibration, policy fit, and adaptation are evaluated separately. Feedback from action and outcomes may update inferred leverage, alter capacity, or change the environment itself.

A Personality-Process Account of Agency

Personality science is most useful here when it explains patterned variation without turning context-sensitive behavior into a fixed label. The CRF model locates individual differences in four processes: control priors, evidence weighting, generalization, and response thresholds.

Control Priors

Before sufficient evidence is available, people rely on expectations acquired from repeated experience. Rotter's locus-of-control framework describes generalized expectancies that reinforcement is contingent on one's behavior rather than chance or powerful others (Rotter, 1966). Bandura's self-efficacy theory places greater emphasis on capability beliefs tied to action (Bandura, 1977). Dispositional optimism concerns generalized expectations of favorable outcomes, which may or may not depend on personal action. These constructs overlap in their motivational consequences but answer different causal questions.

A favorable prior can sustain exploration long enough to discover control. In a preregistered computational study of 279 adults, a Bayesian model with a task-invariant prior best explained trial-level costly active-avoidance choices across differently framed gain and loss tasks (Granwald et al., 2025). The prior showed moderate-to-good reliability and a weak association with positive affect, but the preregistered trait-anxiety association was absent; the finding supports a shared behavioral prior, not a clinical mechanism. A favorable prior can also generate a control premium: people may prefer options that confer choice or perceived control even when expected material value is equivalent, with corticostriatal and ventromedial prefrontal BOLD signals tracking some of that subjective value (Ly et al., 2019; Wang & Delgado, 2019; Wang et al., 2021). Yet a strong prior becomes maladaptive when it prevents recognition of low

control or makes every bad outcome evidence of personal failure. The model therefore predicts curvilinear or context-moderated benefits rather than a universal main effect of internality, efficacy, or optimism.

Evidence Weighting and Updating

People differ in how quickly they infer a contingency, how much weight they give negative versus positive outcomes, and how readily they revise an expectation. In a small exploratory study of 54 adults, neuroticism was associated with lower subjective controllability and greater affective disruption during acute social-evaluative stress, although some findings did not survive multiple-comparison correction and physiological associations were not uniformly aligned (Xin et al., 2017). CRF further hypothesizes that anxiety may raise the cost of exploratory error, conscientious persistence may extend sampling but delay exit, and optimism may protect against premature collapse but become insensitive to stable disconfirmation. These proposed trait-to-process mappings require direct parameter tests; they are not established facts.

Generalization

A central danger in helplessness is not merely learning that one response failed. It is generalizing that inference beyond the response, context, time, or self for which it was valid. Attribution theory explains one route: stable and global interpretations widen the expected boundary of failure (Abramson et al., 1978). Computational accounts explain another: a strong low-control prior learned in one environment can bias inference in the next (Huys & Dayan, 2009). Personality research should therefore study generalization gradients, including how causal structure, available action, domain, social identity, and context cues determine transfer, rather than infer a global trait from one task.

Response Thresholds and Switching Costs

Even with the same estimated control, people may require different amounts of evidence before acting, persisting, seeking help, or quitting. These thresholds reflect reward sensitivity, effort valuation, shame, identity, available alternatives, and cultural norms. A person can be highly persistent because control seems probable, because withdrawal threatens identity, because intermittent success sustains hope, or because no safe alternative exists. Only the first mechanism directly indexes perceived controllability. Designs that measure only duration of effort cannot distinguish them.

Culture, Structure, and the Distribution of Control

Control is not distributed evenly. Wealth, health, citizenship, gender, race, disability, organizational rank, and institutional design can alter the actions available to a person, the costs of using them, and the probability that authorities will respond. A theory that measures only internal mastery risks interpreting an accurate perception of constraint as a cognitive deficit. This is both a scientific error and an ethical one.

Evidence from the COVID-19 pandemic illustrates the distinction. A systematic review identified 20 survey-based observational studies, 16 of which entered a meta-analysis of 24 nested observations. Lower sense of control was associated with greater depression, with a pooled absolute correlation of $|r| = .41$. Measures of perceived external constraints and overall control were more strongly associated with depression than measures of internal mastery, although the evidence was cross-sectional and heterogeneous (Msetfi et al., 2024). The result does not show that constraints cause depression or that changing beliefs would remove the association. It shows that perceived external and internal control are empirically distinguishable and that the environment cannot be treated as background noise.

Longitudinal research similarly indicates that perceived constraints can predict health change more strongly than mastery (Infurna & Mayer, 2015). Yet perceived constraints remain perceptions. Stronger designs should independently measure institutional rules, resource access, decision authority, network leverage, and action–outcome contingencies. Socioeconomic status or group membership is not a sufficient proxy for actual control, because people in the same category can face different affordances and because the same resource can have different causal value across contexts.

Culture shapes which forms of agency are legitimate, visible, and effective. A meta-analysis across 152 independent samples and 18 cultural regions found that external locus of control was associated with depression and anxiety at about $r = .30$. The external-control-anxiety association was weaker in more collectivistic regions, suggesting that the psychological meaning of externality depends partly on cultural emphasis on individual agentic goals (Cheng et al., 2013). This does not license country-level stereotyping. National indicators are ecological summaries, within-culture variation is large, and collectivism is not the absence of agency.

Secondary control makes the point more directly. Adjusting the self, accepting a constraint, or preserving relational harmony may be an active attempt to achieve person–environment fit rather than a passive failure to change the world (Morling & Evered, 2006). Likewise, relational and collective agency can produce outcomes that an individual-agency intervention cannot. A large field experiment in rural Niger found that a culturally grounded relational-agency intervention improved economic outcomes over 12 months, whereas a personal-agency intervention more closely aligned with Western assumptions changed intrapersonal outcomes without the same economic effect (Thomas et al., 2025). One setting

cannot establish a universal cultural rule, but it demonstrates why efficacy must be judged against the social mechanism through which change actually occurs.

The CRF model therefore treats structure and culture as causal inputs with several pathways. Structure shapes mode-specific causal leverage, feasible action sets, access, cost, and risk. Culture shapes response value, legitimacy, interpretation, and switching cost. Personality and learning history may shape priors, attention, updating, and response thresholds within those contexts. None can substitute for the others. When personal control is low because a decision is institutionally closed, raising personal confidence alone may increase frustration or self-blame. The fitting intervention may instead involve access, authority, accompaniment, collective efficacy, legal protection, or a safe exit. Psychological support remains valuable, but it should not be asked to simulate a material opportunity that does not exist.

Collective control also requires explicit levels of analysis. A person's immediate marginal influence on an outcome may be small even when coordinated action has substantial leverage. CRF therefore distinguishes direct individual action, individual contribution to recruitment or coordination, and coalition-level action. A strike vote, for example, may have little direct effect as one isolated act while still being a causally necessary contribution to a threshold process. Evidence for fit must specify the mechanism and level rather than compare individual effort with a collective outcome as if they were the same unit.

Propositions of the Control-Response Fit Model

The following propositions contrast CRF with simpler belief-only, effort-only, reward-only, and context-insensitive accounts. Each requires independent measurement or manipulation of the elements used to score fit.

Proposition 1: Prospective Fit Will Outpredict Belief or Effort Level

Policies with lower preregistered regret will predict performance, resource use, and later adjustment beyond the main effects of perceived control, self-efficacy, optimism, and response intensity. This prediction should hold when leverage, feasibility, costs, values, and alternatives are measured independently of the outcome used to assess adaptation. Failure of prospective fit to improve cross-validated prediction over simpler reward-, cost-, and belief-based models would weaken CRF.

Proposition 2: Calibration Errors Will Produce Distinct Forms of Regret

When feasible leverage exceeds inferred leverage, people should explore too little, omit attainable responses, and withdraw prematurely. When inferred leverage exceeds feasible leverage, people should persist longer in low-value policies and incur greater cost. Frustration and self-blame should increase primarily when overestimation is combined with repeated failure, high desired control, and internal responsibility attribution. These predictions require leverage and control estimates on commensurate scales, with uncertainty modeled explicitly.

Proposition 3: Personality Will Shape Identifiable Process Parameters

Task-specific efficacy, optimism, neuroticism or anxiety, conscientiousness, and prior control history are hypothesized to predict separable parameters such as initial leverage priors, evidence weighting, exploration, and switching thresholds. Their consequences should depend on the causal environment. Optimism may support discovery when informative rewards are sparse but delay exit from a stably null contingency. Claims about trait mechanisms require parameter recovery and out-of-sample prediction. Uniform trait effects that cannot be linked to identifiable CRF parameters would favor a broad response-tendency account.

Proposition 4: Generalization Will Follow Perceived and Structural Similarity

Control learning should alter initial beliefs in a new setting in proportion to perceived similarity and actual similarity in causal structure, available actions, domain, and relevant social actors. These dimensions should be manipulated or measured separately. Stable and global attributions should broaden transfer, whereas specific attributions should constrain it. Transfer that remains global after controlling for affect, fatigue, demand, and similarity would challenge the proposed generalization mechanism.

Proposition 5: Flexibility Will Reduce Regret Under Contingency Change

Evidence-responsive policy switching should matter most when leverage, costs, or feasible agency routes change. In reversal tasks, belief updating and policy change should be temporally coordinated, although exploration may appropriately precede confident belief revision because action generates information. Switching frequency alone is not flexibility. People who switch in response to diagnostic evidence should show lower subsequent regret than people who switch indiscriminately or persist despite stable disconfirmation.

Proposition 6: Disengagement Will Help Most When It Supports Valued Reengagement

When continued goal pursuit has high prospective regret, disengagement should reduce intrusive effort and resource loss. Gains in meaning, positive affect, and functioning should be larger when a valued and attainable alternative is available. When feasible leverage remains high and costs are proportionate, induced disengagement should impair goal progress. The predicted attainability by disengagement by alternative-availability interaction distinguishes adaptive redirection from generalized withdrawal.

Proposition 7: Response Mode Will Track the Location of Feasible Leverage

Personal action should increase when personal leverage and capacity are high. Proxy action should increase when an accessible expert or institution controls the relevant decision. Collective action should increase when coordinated behavior has leverage and contribution or coordination costs remain acceptable. Accommodation should be favored when external change is presently infeasible and internal adjustment does not violate safety or core values. Mode-specific profiles should predict policy allocation better than a single internal-locus score after content overlap is controlled.

Proposition 8: Affordance and Belief Change Will Have Complementary Effects

Affordance-enhancing interventions should expand feasible action and improve reachable outcomes. Belief or mastery interventions should improve exploration and use of available routes when those routes exist. Their combination should produce the most durable gains when new opportunities must be detected, trusted, and enacted. In persistently closed environments, successful calibration should promote redirection or alternative agency rather than unsupported persistence. If independently validated affordance changes add no predictive or causal value across settings with different access, the structural component of CRF would require revision.

A Research Program for Personality Science

The CRF model is designed for cumulative testing across experiments, longitudinal observation, and culturally diverse settings. No single study can establish all layers. A coordinated program can.

Study 1: Controllability, Reward Prevalence, and Reversal

A preregistered experiment should manipulate causal leverage, reward prevalence, and response cost as independently as the task permits. Counterfactual or yoked schedules should

equate marginal reward prevalence and information across control conditions so that participants do not obtain systematically different evidence merely because their policies differ. A later reversal would change leverage while holding other task features stable. Control judgments should be probed sparsely or in randomized probe schedules because repeated ratings can alter learning.

Computational comparison of action-sensitive and action-insensitive models could estimate inferred leverage and arbitration between instrumental and Pavlovian policies (Dorfman & Gershman, 2019; Ligneul et al., 2022). Simulation, model recovery, parameter recovery, and posterior predictive checks should precede substantive parameter interpretation. Confirmatory outcomes should be limited to a small set, such as information-seeking, policy allocation, switching latency, and resource expenditure; affect can be a registered secondary outcome. The critical analysis would compare CRF's preregistered policy-regret prediction out of sample with reward-only, cost-only, belief-only, and effort-only models.

A small set of theoretically directional personality predictors should be measured before the task. For example, task-specific efficacy may predict the prior, anxiety may predict uncertainty or error cost, and conscientious persistence may predict a switching threshold. The study should be powered for selected interactions, use multilevel trial models, and avoid treating a significant trait moderator as proof of a latent computational parameter.

Study 2: Generalization Across Domains

A second preregistered study should test transfer across counterbalanced task families. Participants would first encounter verifiable control, yoked uncontrollability, or an information-matched comparison and then enter a novel task. Surface similarity and causal-structure

similarity should be manipulated separately, rather than bundled into one post hoc rating. The new task's leverage should also be independently manipulated.

The design would test whether prior experience shifts an initial prior, learning rate, generalization boundary, or policy at a given belief. Affect, stress, fatigue, expectations about the experiment, and demand should be assessed as rival explanations. Perceived similarity is a potential mediator or process measure, not a randomized moderator. A delayed follow-up would distinguish immediate carryover from durable generalization, and claims of inoculation or helplessness should be reserved for preregistered transfer outcomes rather than same-task performance.

Study 3: Intensive Longitudinal Control-Response Fit

Laboratory tasks simplify control but cannot represent the layered contingencies of daily life. A prospective event-contingent study could enroll focal stressors at onset, before outcomes are known. At baseline, participants would estimate personal, proxy, and collective leverage with uncertainty, report actual access and capacity, desired control, values, risks, and intended policy. Later assessments would capture enacted responses, costs, new information, reengagement, and outcomes.

Independent ratings of narratives should be described as externally coded affordance indicators, not objective control. They should be supplemented where possible by administrative rules, decision authority, deadlines, network position, third-party evidence, or known causal mechanisms. Within-person models could test whether prospectively scored lower-regret policies predict later outcomes relative to a person's usual pattern. Lagged or state-space models should allow bidirectional belief-policy coupling and should separate contemporaneous mood from

temporal prediction. Sensitivity analyses should vary the feasible benchmark rather than assume one rater can identify the uniquely correct real-world policy.

Study 4: Cross-Cultural and Structural Tests

A multisite study should sample settings that differ in relational interdependence and institutional affordances without treating nationality as the mechanism. Measures should be translated and co-developed with local participants where needed. Metric, scalar, partial, or approximate measurement invariance should be evaluated according to the comparison being made; perfect invariance should not be treated as an all-or-none gate. Personal, proxy, collective, relational, and accommodative processes should be measured with parallel items, while actual access, authority, network leverage, and institutional response are documented separately.

The strongest test would experimentally vary agency route and affordance where ethically feasible. Clusters, sites, or networks should be assigned when the intervention operates above the person, with enough units for cross-level inference and a plan for spillover. Conditions might include individual problem solving, relational or collective coordination, an affordance-enhancing resource, and their combination. Outcomes should include observed behavior, access, and institutional response, not only self-report. This design can reveal whether a program fails because beliefs did not change, a feasible route was absent, collective coordination did not occur, or the selected policy did not use the available route.

Study 5: Adaptive Intervention

An eventual intervention should begin with an uncertainty-aware diagnostic phase rather than a positivity exercise. Participants would learn to separate wanted from inferred control, run low-cost tests where safe, identify personal, proxy, and collective routes, set stopping rules, and

pair disengagement with a reengagement plan. Ground truth will often be incomplete, so module assignment should use explicit confidence thresholds and allow a safe “uncertain” pathway.

A sequential multiple assignment randomized trial could compare candidate modules: interpretable mastery experience for underestimation where leverage exists; monitoring and stopping rules for costly over-engagement; advocacy or resource navigation for restricted access; and acceptance plus values-based reengagement for outcomes unlikely to change. A nested factorial component should directly test Proposition 6 by varying evidence of attainability, disengagement support, and the availability or value of an alternative goal. Conditions should have equal contact, mediators should precede outcomes, module fidelity should be assessed, and safety triage should take priority wherever coercion, violence, severe illness, or legal jeopardy is possible. The model predicts improvement through better calibration, lower policy regret, and fitting switching, not through a uniform increase in perceived control.

Table 4

Proposition-to-Study Test Matrix

Proposition	Primary study	Critical design element	Primary discriminating result or falsifier
1. Prospective fit	1 and 3	Independent leverage, access, cost, belief, and policy measures	Regret improves out-of-sample prediction beyond simpler models; repeated null incremental value weakens CRF
2. Calibration errors	1	Commensurate estimate and leverage scales with uncertainty and cost	Underestimation predicts omitted opportunity; overestimation predicts cost; self-blame requires internal attribution and failed effort
3. Personality parameters	1	Selected directional trait-parameter hypotheses and recovery tests	Recoverable, environment-dependent parameters; invariant undirected trait effects favor simpler tendencies
4. Generalization	2	Separate structural and surface similarity; counterbalanced tasks	Transfer follows preregistered similarity structure beyond affect, fatigue, and demand
5. Flexibility	1 and 3	Reversal or natural change with intensive temporal measurement	Evidence-responsive switching lowers later regret; raw switching frequency does not suffice

Proposition	Primary study	Critical design element	Primary discriminating result or falsifier
6. Disengagement and reengagement	5	Attainability by disengagement by valued-alternative manipulation	Reengagement amplifies benefit when pursuit is unattainable; disengagement harms progress when leverage remains high
7. Location of leverage	3 and 4	Parallel personal, proxy, and collective leverage and efficacy measures	Mode-matched profiles predict allocation beyond general internal control
8. Affordance and belief	4 and 5	Factorial affordance and belief or mastery change	Affordance contributes independently and jointly where access limits behavior; durable null affordance effects weaken structural claims

Measurement and Open Science Requirements

A precise theory can still fail through imprecise measurement. Control scales should not be selected solely because they report high internal consistency. Reliability does not establish structural validity, invariance, or correspondence to an action–outcome mechanism (Flake & Fried, 2020). Broad locus-of-control or mastery scores should be supplemented with domain-specific efficacy, repeated perceived-control judgments, and behavioral indicators. Perceived constraints should not be reverse-coded and relabeled as mastery.

Objective control requires a documented causal structure. Choice alone is not control; a participant can choose among options that yield identical outcomes. Predictability alone is not control; an uncontrollable event can be perfectly predicted. Success rate alone is not control; rewards can be frequent but response-independent. Effort and motor output should be matched across yoked conditions or modeled explicitly. Manipulation checks should assess whether participants detected the contingency, but post-randomization perception should not replace the randomized estimate of objective control.

Response policy also requires disaggregation. Persistence, exploration, strategy switching, help-seeking, coordinated action, acceptance, and disengagement should be coded separately. Low activity cannot reveal whether someone accepted, froze, conserved resources, or

redirected effort. Outcomes should be registered by level and time: task performance, affect, physiological response, later learning, symptoms, and real-world functioning.

The empirical program should preregister primary interactions and computational comparisons, report all conditions and exclusions, and archive materials, deidentified data where ethically possible, and analysis code. Multiverse or specification-curve analyses would be useful when fit can be operationalized in several defensible ways. Cross-cultural studies should test measurement invariance and avoid post hoc explanations based on country labels. Null effects and failed manipulations are theoretically informative. For example, the particular brief self-efficacy induction used by Meier et al. (2025) failed to change its manipulation check or affective outcomes; that result constrains the induction, not every possible mastery or efficacy intervention.

Implications for Intervention and Practice

The model supports hope, but it changes the object of hope. The aim is not to persuade people that every obstacle is controllable. It is to increase their ability to locate leverage, test it safely, and preserve agency across changing conditions.

For under-engagement in a controllable domain, CRF identifies an interpretable mastery experience as a leading intervention candidate. The task would need to be difficult enough to provide diagnostic evidence, structured so that the person's action genuinely affects the outcome, and followed by feedback that identifies the effective contingency. Generic praise or positive self-statements provide less direct causal evidence because they can be dismissed or can conflict with experience. Prior control can protect against later uncontrollable stress in some animal preparations, but durable, cross-domain inoculation in humans remains to be demonstrated.

For over-engagement, a candidate intervention would preserve caring and ambition while clarifying the contingency, making costs visible, identifying stopping rules, and expanding the response repertoire. Questions such as “What evidence would tell us this strategy is not working?” and “Who else has leverage?” can reframe stopping from a verdict on the self into a decision about a policy. The model predicts that benefits of disengagement will be larger when one route is linked to another valued and attainable direction; this remains a causal proposition to test.

Where personal control is low but social control exists, intervention should strengthen help-seeking, advocacy, coalition, or collective efficacy. Where danger, coercion, or institutional exclusion makes action costly, safety and access take precedence over persistence training. Where an outcome is irreversible, acceptance can be an expression of agency when it permits choice about attention, meaning, relationship, or the next goal.

The clinical implications remain provisional. Perceived control is associated with anxiety and depression, but association does not prove that increasing control beliefs will treat either disorder. Acute controllability tasks in healthy participants cannot establish a clinical mechanism. Intervention research should test whether changes in inferred controllability and response fit precede symptom change, whether effects generalize to daily life, and whether structural supports moderate benefit. Individual psychological work and environmental change are complementary, not competing, levels of intervention.

Boundary Conditions and Limitations

First, objective causal leverage is difficult to identify outside the laboratory. Real goals contain multiple outcomes, actors, time scales, and values. Researchers may incorrectly label a situation uncontrollable because one response fails while another, collective or delayed response

remains possible. The solution is not to abandon the construct but to specify the response, comparison, outcome, actor level, horizon, access conditions, and cost.

Second, fit can become tautological if the “right” response is defined after observing a good outcome. CRF studies must specify the causal affordance and predicted response in advance. When the optimal policy is uncertain, the model predicts bounded exploration and updating rather than a single strategy.

Third, mild positive illusions may sometimes protect affect or sustain search. The model does not require perfect perceptual accuracy at every moment. It predicts that bias becomes costly when it systematically produces policy misfit under consequential, repeated feedback. The optimum may therefore be a region of functional calibration rather than zero error.

Fourth, the model integrates literatures with unequal evidence. Causal circuitry for controllability is strongest in male rodents; human neural findings are convergent but do not establish the same circuit. Human behavioral paradigms remain heterogeneous, frequently underpowered for individual-difference interactions, and vulnerable to differences in effort, reward rate, predictability, and demand. Findings about sex in rodents cannot be mapped directly onto human gender, and observed human gender interactions require replication.

Fifth, many personality and well-being findings are correlational and depend on self-report. Perceived control can influence mood, mood can influence perceived control, and both can reflect third variables or actual constraint. Intensive longitudinal and experimental work is needed before describing control beliefs as causal resilience factors.

Finally, no general model can decide the value of every goal. Some controllable goals are harmful, and some costly collective struggles remain worthy despite low short-term success probability. CRF is a model of functional alignment, not a moral instruction to optimize comfort.

Values, justice, and obligations determine which outcomes deserve pursuit; the model asks how agency can be exercised honestly within that choice.

Conclusion

The enduring insight of learned helplessness is that a history of uncontrollability can change later behavior even after opportunity returns. The enduring insight of positive psychology is that people can learn patterns that support hope, persistence, and flourishing. A contemporary synthesis must add a third insight: more control belief and more persistence are not always better.

Agency is calibrated when people can distinguish an unavailable route from an unavailable future. It persists where action has leverage, explores when leverage is uncertain, recruits others when control is distributed, and releases a goal when continued direct effort cannot serve it. It remains hopeful without requiring denial of constraint.

The Control-Response Fit model makes this position testable. It predicts different costs for under- and overestimation, locates personality in priors and updating, treats culture and structure as causal features of control, and defines resilience partly as the ability to switch modes without globalizing failure. The practical promise is neither passivity nor relentless effort. It is wiser action: discovering what can be changed, changing it where possible, and carrying agency forward when the form of change must change.

References

- Abramson, L. Y., Seligman, M. E. P., & Teasdale, J. D. (1978). Learned helplessness in humans: Critique and reformulation. *Journal of Abnormal Psychology*, 87(1), 49–74.
<https://doi.org/10.1037/0021-843X.87.1.49>
- Amat, J., Baratta, M. V., Paul, E., Bland, S. T., Watkins, L. R., & Maier, S. F. (2005). Medial prefrontal cortex determines how stressor controllability affects behavior and dorsal raphe nucleus. *Nature Neuroscience*, 8(3), 365–371. <https://doi.org/10.1038/nn1399>
- Amat, J., Paul, E., Zarza, C., Watkins, L. R., & Maier, S. F. (2006). Previous experience with behavioral control over stress blocks the behavioral and dorsal raphe nucleus activating effects of later uncontrollable stress: Role of the ventral medial prefrontal cortex. *Journal of Neuroscience*, 26(51), 13264–13272. <https://doi.org/10.1523/JNEUROSCI.3630-06.2006>
- Bandura, A. (1977). Self-efficacy: Toward a unifying theory of behavioral change. *Psychological Review*, 84(2), 191–215. <https://doi.org/10.1037/0033-295X.84.2.191>
- Bandura, A. (2001). Social cognitive theory: An agentic perspective. *Annual Review of Psychology*, 52, 1–26. <https://doi.org/10.1146/annurev.psych.52.1.1>
- Baratta, M. V., Leslie, N. R., Fallon, I. P., Dolzani, S. D., Chun, L. E., Tamalunas, A. M., Watkins, L. R., & Maier, S. F. (2018). Behavioural and neural sequelae of stressor exposure are not modulated by controllability in females. *European Journal of Neuroscience*, 47(8), 959–967. <https://doi.org/10.1111/ejn.13833>
- Baratta, M. V., Seligman, M. E. P., & Maier, S. F. (2023). From helplessness to controllability: Toward a neuroscience of resilience. *Frontiers in Psychiatry*, 14, Article 1170417. <https://doi.org/10.3389/fpsyt.2023.1170417>

- Barlow, M. A., Wrosch, C., & McGrath, J. J. (2020). Goal adjustment capacities and quality of life: A meta-analytic review. *Journal of Personality*, 88(2), 307–323.
<https://doi.org/10.1111/jopy.12492>
- Bhanji, J. P., & Delgado, M. R. (2014). Perceived control influences neural responses to setbacks and promotes persistence. *Neuron*, 83(6), 1369–1375.
<https://doi.org/10.1016/j.neuron.2014.08.012>
- Bhanji, J. P., Kim, E. S., & Delgado, M. R. (2016). Perceived control alters the effect of acute stress on persistence. *Journal of Experimental Psychology: General*, 145(3), 356–365.
<https://doi.org/10.1037/xge0000137>
- Bonanno, G. A., & Burton, C. L. (2013). Regulatory flexibility: An individual differences perspective on coping and emotion regulation. *Perspectives on Psychological Science*, 8(6), 591–612. <https://doi.org/10.1177/1745691613504116>
- Cheng, C., Cheung, S.-F., Chio, J. H.-M., & Chan, M.-P. S. (2013). Cultural meaning of perceived control: A meta-analysis of locus of control and psychological symptoms across 18 cultural regions. *Psychological Bulletin*, 139(1), 152–188.
<https://doi.org/10.1037/a0028596>
- Cheng, C., Lau, H.-P. B., & Chan, M.-P. S. (2014). Coping flexibility and psychological adjustment to stressful life changes: A meta-analytic review. *Psychological Bulletin*, 140(6), 1582–1607. <https://doi.org/10.1037/a0037913>
- Davis, C. J., Spearman, S. R., & Burrow, A. L. (2026). Perceived, desired, and discrepant control: Initial cross-sectional evidence for a metacognitive calibration framework for depression risk. *Cognitive Therapy and Research*. Advance online publication.
<https://doi.org/10.1007/s10608-026-10725-2>

- Dorfman, H. M., & Gershman, S. J. (2019). Controllability governs the balance between Pavlovian and instrumental action selection. *Nature Communications*, 10, Article 5826. <https://doi.org/10.1038/s41467-019-13737-7>
- Flake, J. K., & Fried, E. I. (2020). Measurement schmeasurement: Questionable measurement practices and how to avoid them. *Advances in Methods and Practices in Psychological Science*, 3(4), 456–465. <https://doi.org/10.1177/2515245920952393>
- Forsythe, C. J., & Compas, B. E. (1987). Interaction of cognitive appraisals of stressful events and coping: Testing the goodness of fit hypothesis. *Cognitive Therapy and Research*, 11(4), 473–485. <https://doi.org/10.1007/BF01175357>
- Gallagher, M. W., Bentley, K. H., & Barlow, D. H. (2014). Perceived control and vulnerability to anxiety disorders: A meta-analytic review. *Cognitive Therapy and Research*, 38(6), 571–584. <https://doi.org/10.1007/s10608-014-9624-x>
- Granwald, T., Dayan, P., Lengyel, M., & Guitart-Masip, M. (2025). A task-invariant prior explains trial-by-trial active avoidance behaviour across gain and loss tasks. *Communications Psychology*, 3, Article 82. <https://doi.org/10.1038/s44271-025-00254-1>
- Habermann, M., & Büchel, C. (2025). Controllability changes pain perception by increasing the precision of expectations. *Nature Communications*, 16, Article 10113. <https://doi.org/10.1038/s41467-025-66038-7>
- Haines, S. J., Gleeson, J., Kuppens, P., Hollenstein, T., Ciarrochi, J., Labuschagne, I., Grace, C., & Koval, P. (2016). The wisdom to know the difference: Strategy-situation fit in emotion regulation in daily life is associated with well-being. *Psychological Science*, 27(12), 1651–1659. <https://doi.org/10.1177/0956797616669086>

- Heckhausen, J., Wrosch, C., & Schulz, R. (2010). A motivational theory of life-span development. *Psychological Review*, 117(1), 32–60. <https://doi.org/10.1037/a0017668>
- Hiroto, D. S., & Seligman, M. E. P. (1975). Generality of learned helplessness in man. *Journal of Personality and Social Psychology*, 31(2), 311–327. <https://doi.org/10.1037/h0076270>
- Huys, Q. J. M., & Dayan, P. (2009). A Bayesian formulation of behavioral control. *Cognition*, 113(3), 314–328. <https://doi.org/10.1016/j.cognition.2009.01.008>
- Infurna, F. J., & Mayer, A. (2015). The effects of constraints and mastery on mental and physical health: Conceptual and methodological considerations. *Psychology and Aging*, 30(2), 432–448. <https://doi.org/10.1037/a0039050>
- Leslie-Miller, C. J., Joormann, J., & Quinn, M. E. (2025). Coping flexibility: Match between coping strategy and perceived stressor controllability predicts depressed mood. *Affective Science*, 6(1), 94–103. <https://doi.org/10.1007/s42761-024-00275-9>
- Levy, E. S., Frank, M. G., Maier, S. F., & Baratta, M. V. (2026). Coping with stress produces differential systemic and brain corticosterone patterns in female rats. *Frontiers in Behavioral Neuroscience*, 20, Article 1760267. <https://doi.org/10.3389/fnbeh.2026.1760267>
- Ligneul, R., Mainen, Z. F., Ly, V., & Cools, R. (2022). Stress-sensitive inference of task controllability. *Nature Human Behaviour*, 6(6), 812–822. <https://doi.org/10.1038/s41562-022-01306-w>
- Ly, V., Wang, K. S., Bhanji, J., & Delgado, M. R. (2019). A reward-based framework of perceived control. *Frontiers in Neuroscience*, 13, Article 65. <https://doi.org/10.3389/fnins.2019.00065>

Maier, S. F., & Seligman, M. E. P. (1976). Learned helplessness: Theory and evidence. *Journal of Experimental Psychology: General*, 105(1), 3–46. [https://doi.org/10.1037/0096-](https://doi.org/10.1037/0096-3445.105.1.3)

[3445.105.1.3](https://doi.org/10.1037/0096-3445.105.1.3)

Maier, S. F., & Seligman, M. E. P. (2016). Learned helplessness at fifty: Insights from neuroscience. *Psychological Review*, 123(4), 349–367.

<https://doi.org/10.1037/rev0000033>

McNulty, C. J., Fallon, I. P., Amat, J., Sanchez, R. J., Leslie, N. R., Root, D. H., Maier, S. F., & Baratta, M. V. (2023). Elevated prefrontal dopamine interferes with the stress-buffering properties of behavioral control in female rats. *Neuropsychopharmacology*, 48, 498–507.

<https://doi.org/10.1038/s41386-022-01443-w>

Meier, J., Kollmann, B., Meine, L. E., Meyer, B., Yuen, K., Stork, M., Tüscher, O., & Wessa, M. (2026). Perceived control as a resilience factor: Associations with neural, physiological and affective stress responses and mental health. *Translational Psychiatry*, 16, Article 39.

<https://doi.org/10.1038/s41398-025-03786-6>

Meier, J., Meine, L. E., Schüler, K., & Wessa, M. (2025). Distinct trajectories of perceived control over aversive stimulation predict affective reactions to stressors over and above objective control. *Scientific Reports*, 15, Article 35009. [https://doi.org/10.1038/s41598-](https://doi.org/10.1038/s41598-025-19958-9)

[025-19958-9](https://doi.org/10.1038/s41598-025-19958-9)

Meine, L. E., Meier, J., Meyer, B., & Wessa, M. (2021). Don't stress, it's under control: Neural correlates of stressor controllability in humans. *NeuroImage*, 245, Article 118701.

<https://doi.org/10.1016/j.neuroimage.2021.118701>

- Meine, L. E., Schüler, K., Richter-Levin, G., Scholz, V., & Wessa, M. (2020). A translational paradigm to study the effects of uncontrollable stress in humans. *International Journal of Molecular Sciences*, 21(17), Article 6010. <https://doi.org/10.3390/ijms21176010>
- Morling, B., & Evered, S. (2006). Secondary control reviewed and defined. *Psychological Bulletin*, 132(2), 269–296. <https://doi.org/10.1037/0033-2909.132.2.269>
- Msetfi, R. M., Kornbrot, D. E., & Halbrook, Y. J. (2024). The association between the sense of control and depression during the COVID-19 pandemic: A systematic review and meta-analysis. *Frontiers in Psychiatry*, 15, Article 1323306. <https://doi.org/10.3389/fpsyt.2024.1323306>
- Na, S., Chung, D., Hula, A., Perl, O., Jung, J., Heflin, M., Blackmore, S., Fiore, V. G., Dayan, P., & Gu, X. (2021). Humans use forward thinking to exploit social controllability. *eLife*, 10, Article e64983. <https://doi.org/10.7554/eLife.64983>
- Overmier, J. B., & Seligman, M. E. P. (1967). Effects of inescapable shock upon subsequent escape and avoidance responding. *Journal of Comparative and Physiological Psychology*, 63(1), 28–33. <https://doi.org/10.1037/h0024166>
- Riddell, H., Sedikides, C., Sivaramakrishnan, H., Wan, P., Maltagliati, S., Jackson, B., Thøgersen-Ntoumani, C., Gucciardi, D. F., & Ntoumanis, N. (2026). A meta-analytic review and conceptual model of the antecedents and outcomes of goal adjustment in response to striving difficulties. *Nature Human Behaviour*, 10, 317–332. <https://doi.org/10.1038/s41562-025-02312-4>
- Rotter, J. B. (1966). Generalized expectancies for internal versus external control of reinforcement. *Psychological Monographs: General and Applied*, 80(1), 1–28. <https://doi.org/10.1037/h0092976>

- Skinner, E. A. (1996). A guide to constructs of control. *Journal of Personality and Social Psychology*, 71(3), 549–570. <https://doi.org/10.1037/0022-3514.71.3.549>
- Song, X., & Vilares, I. (2021). Assessing the relationship between the human learned helplessness depression model and anhedonia [Registered report protocol]. *PLOS ONE*, 16(3), Article e0249056. <https://doi.org/10.1371/journal.pone.0249056>
- Teodorescu, K., & Erev, I. (2014). Learned helplessness and learned prevalence: Exploring the causal relations among perceived controllability, reward prevalence, and exploration. *Psychological Science*, 25(10), 1861–1869. <https://doi.org/10.1177/0956797614543022>
- Thomas, C. C., Premand, P., Bossuroy, T., Sambo, S. A., Markus, H. R., & Walton, G. M. (2025). How culturally wise psychological interventions can help reduce poverty. *Proceedings of the National Academy of Sciences*, 122(46), Article e2505694122. <https://doi.org/10.1073/pnas.2505694122>
- Wang, K. S., & Delgado, M. R. (2019). Corticostriatal circuits encode the subjective value of perceived control. *Cerebral Cortex*, 29(12), 5049–5060. <https://doi.org/10.1093/cercor/bhz045>
- Wang, K. S., & Delgado, M. R. (2021). The protective effects of perceived control during repeated exposure to aversive stimuli. *Frontiers in Neuroscience*, 15, Article 625816. <https://doi.org/10.3389/fnins.2021.625816>
- Wang, K. S., Yang, Y.-Y., & Delgado, M. R. (2021). How perception of control shapes decision making. *Current Opinion in Behavioral Sciences*, 41, 85–91. <https://doi.org/10.1016/j.cobeha.2021.04.003>

- Wanke, N., & Schwabe, L. (2020). Subjective uncontrollability over aversive events reduces working memory performance and related large-scale network interactions. *Cerebral Cortex*, 30(5), 3116–3129. <https://doi.org/10.1093/cercor/bhz298>
- Willroth, E. C., Young, G., Ford, B. Q., Troy, A. S., Swierzewicz, D., & Mauss, I. B. (2025). Preregistered direct replication and extension of “The wisdom to know the difference: Strategy-situation fit in emotion regulation in daily life is associated with well-being.” *Psychological Science*, 36(5), 367–383. <https://doi.org/10.1177/09567976251335567>
- Wrosch, C., Scheier, M. F., Miller, G. E., Schulz, R., & Carver, C. S. (2003). Adaptive self-regulation of unattainable goals: Goal disengagement, goal reengagement, and subjective well-being. *Personality and Social Psychology Bulletin*, 29(12), 1494–1508. <https://doi.org/10.1177/0146167203256921>
- Xin, Y., Wu, J., Yao, Z., Guan, Q., Aleman, A., & Luo, Y. (2017). The relationship between personality and the response to acute psychological stress. *Scientific Reports*, 7, Article 16906. <https://doi.org/10.1038/s41598-017-17053-2>